

Sergey Subbotin
Editor

Computer Modeling and Intelligent Systems

Proceedings of The Fifth International Workshop
CMIS-2022



Zaporizhzhia, Ukraine
May 12, 2022

Subbotin S., (Ed.): The Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022). Zaporizhzhia, Ukraine, May 12, 2022, CEUR-WS.org, online

This volume represents the proceedings of the Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022), held in Zaporizhzhia, Ukraine, in May, 2022. It comprises 25 contributed papers that were carefully peer-reviewed and selected from 53 submissions.

Copyright © 2022 for the individual papers by the papers' authors.
Copying permitted only for private and academic purposes.
This volume is published and copyrighted by its editors.

Preface

It is my pleasure to present you the proceedings of The Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022), held in Zaporizhzhia, Ukraine on 12 of May 2022.

The CMIS workshop was held in 2022 for the fifth time at the National University "Zaporizhzhia Polytechnic" (<http://www.zp.edu.ua>), which is the largest and oldest regional center for science and higher education in the Zaporizhzhia region of South-Eastern Ukraine.

The Fifth International Workshop on Computer Modeling and Intelligent Systems is devoted to the 20-th Anniversaries of the foundation of the Faculty of Computer Science and Technology and the Department of Software Tools of the National University "Zaporizhzhia Polytechnic".

The Faculty of Computer Science and Technology traces its history back to the Faculty of Radio Engineering, where computer specialties were opened in the 1990s (1991 - "Computer Software and Automated Systems", 1992 - "Computer Systems and Networks", 2000 - "Specialized Computer Systems"). A new Faculty of Informatics and Computer Technology was established in 2002, and headed by the Dean, Associate Professor, Ph.D M. Kasian. The new faculty contained three new departments: Department of Software Tools, Department of Computer Systems and Networks, and Department of Computational Mathematics. In 2004 the specialty "Systems Analysis and Management" was opened, and the Department of Computational Mathematics was renamed as the Department of Systems Analysis and Computational Mathematics. Today, the faculty trains bachelors and masters in four specialties: 121 "Software Engineering", 122 "Computer Science", 123 "Computer Engineering" and 124 "Systems Analysis". Doctors of philosophy are prepared in two specialties: 122 "Computer Science", and 124 "Systems Analysis". Over the years of existence of the faculty's specialties, more than three thousand graduates of IT specialties have graduated.

The Department of Software Tools is the biggest IT department in Zaporizhzhia region and the biggest department in the University by the number of students. The history of the Software Tools Department begin in 1991, when, the specialty "Software for Computer Engineering and Automated Systems" was opened for the first time in the Zaporizhzhia region on the basis of the Department of Design and Production of Radio Electronic Equipment at the Vlas Chubar Zaporizhzhia Machine-Building Institute (later Zaporizhzhia State Technical University, then Zaporizhzhia National Technical University, now is the National University "Zaporizhzhia Polytechnic"). The founders of the specialty and IT education at University in the 1990s were Associate Professors V. Kryshchuk, S. Kornienko, M. Vasylega, V. Vysochyn, V. Dubrovin, V. Lebid', V. Onyshchenko, O. Petryshchev, G. Polyansky, V. Rysikov, S. Serdyuk, E. Tkachev, Senior Lecturers I. Levada, T. Onikienko, Assistant Professors O. Pozdnyakov, O. Stepanenko, L. Romanova (Deinega), S. Mylynchuk, L. Skachko and others. The employees of other departments were also involved in teaching at the department: Associate Professor V. Pinchuk, Associate Professor V. Tomashevs'ky, Senior Lecturers B. Soldatov, A. Markin and others. The specialty

was licensed and subsequently accredited thanks to the efforts of the teaching staff. In 1996, the first graduation of specialists with the qualification of "Software Engineer" took place.

Then there was a change in the name of the specialty: it became known as "Automated Systems Software" (ASS). On the basis of the ASS specialty, in 2000 the specialization "Software in business, management and entrepreneurship" was introduced. In 2002, a new Department of Software Tools was created, which became the basis for the specialty ASS and became part of the newly formed Faculty of Informatics and Computer Engineering (now it is the Faculty of Computer Science and Technologies). From the beginning of its creation and until 2012, the department was headed by Associate Professor, Ph.D A. Prytula. In 2004, a new specialty "Information Technologies of Design" (ITD) was opened at the department. In 2006-2009 the Department of Software Tools firstly in the University participated in Tempus project for joint Master Degrees in Software Engineering with Polish, Czech and German universities. This was a start of further massive work on internationalization at the department and in the University. In the 2000s, the department experienced a rapid expansion of the teaching staff: representatives of production (Senior Teacher A. Vasylenko) and young teachers (A. Parkhomenko, H. Nelasa, G. Shylo, G. Tabunshchyyk, S. Subbotin, N. Myronova, M. Andreev, M. Kalinina, Zh. Kamins'ka, Ye. Fedorchenko, O. Kachan, A. Oliynyk). The teachers from other departments was also involved to the work of the Department of Software Tools (Assoc. Prof., Ph.D K. Kasian, Assoc. Prof., Ph.D R. Kudermetov and others). The technical staff (M. Andreev, M. Kalinina, L. Zonova, Ya. Zhornyk, O. Kavins'ka, and O. Ryashchenko) helped in ensuring the department's educational process.

From 2011 to 2016, the department was headed by Professor, Ph.D V. Dubrovin. There was a change in the name of the specialty ASS: it became known as "Software of Systems" (SS). In the 2010s, the new teachers (O. Oliynyk, Ye. Gofman, S. Zaitsev, V. Liovkin, T. Fedoronchak, T. Kolpakova, T. Zaiko) starting work at the department. In this period the technical staff of the department was extended by new people (K. Ludvikevich, O. Goloviznin, V. Ryabets). In 2012, the specialties "Software Engineering" (SE) and "Artificial Intelligence Systems" (AIS) were opened at the department. Two new laboratories have been created at the department: the laboratory "CAD / CAM / CAE" (2012) and the laboratory "Embedded systems and remote engineering" (2015).

In 2014, for the first time, an employee of the department (S. Subbotin) defended his doctoral habilitation dissertation (the scientific consultant is Prof. D. Piza).

In 2015, according to the new list of specialties, a transition was made to training students in the specialties 121 "Software Engineering" (combines the specialties of SE and SS) and 122 "Computer Science" (combines the specialties of ITD and AIS).

In 2015, the team of students of the Software Tools Department took the first absolute place in the final table of the semi-finals of the world of the ACM Student Programming Contest in South-East Europe and the first absolute place in the final table of the final of the Ukrainian Championship in programming ACM in the major league.

Since 2016, the department has been headed by a former graduate of the department, Doctor habilitated of Science, Professor S. Subbotin. In autumn 2016, the Department of Software Tools expanded the licensed recruitment of Master's students to 150 people in total. In 2016, in the spring, the department licensed postgraduate studies for the preparation of Doctors of Philosophy (Ph.D) in the specialty 122 "Computer Science". In May 2016, a team of students of the department represented Ukraine in the finals of the International ACM Student Programming Contest in Thailand. In the 2018-19 academic year, the department successfully passed the accreditation of bachelor's and master's programs in both specialties 121 "Software Engineering" and 122 "Computer Science".

Each year from 2016 to 2021 the Department of Software Tools became a leader on the recruitment of applicants at the university (both in terms of the number of applications and the number of recruited students). In the 2021-2022 academic year, the number of students of the department reached 850 people, which is the largest number of students in the history of the department, and the department itself became the largest in terms of the number of students in the university and in the region.

For the period 2016-2021 the department has twice increased the number of lecture audiences assigned to it. The new laboratories of "Intelligent Computing, Virtual and Augmented Reality" and "Software Engineering" have been created at the department. The number of computer classes and laboratories has increased by 1.75 times. All classrooms and laboratories of the department were repaired and equipped with new furniture, air conditioners and multimedia projectors.

In 2021, two applicants for the Department of Software Tools (A. Oliinyk, M. Polyakov) has defended their dissertations for doctoral habilitation (the scientific consultant is S. Subbotin).

More than two thousand graduates of full-time, part-time and postgraduate forms of education at the levels of "bachelor", "specialist" (engineer), and "master" have been prepared on specialties of the Department of Software Tools over the years of existence in the University. They work at the enterprises of the Zaporizhzhia region, Ukraine and abroad.

Today, among the department's staff and external staff involved to the department's work there are three Doctors habilitated of Sciences (S. Subbotin, A. Oliinyk, M. Polyakov) and 21 Ph.Ds (O. Gladkova, T. Golub, Ye. Gofman, V. Dubrovin, T. Zaiko, I. Zeleneva, A. Kazurova, T. Kaplienko, M. Kasian, T. Kolpakova, S. Kornienko, M. Kotsur, V. Liovin, A. Parkhomenko, S. Serdyuk, S. Skrupsky, O. Stepanenko, G. Tabunshchik, E. Tereshchenko, T. Fedoronchak, O. Shytikova). Four employees have the academic title of Professor (S. Subbotin, A. Oliinyk, V. Dubrovin, G. Tabunshchik), 16 employees have the academic title of Associate Professor (T. Zayko, I. Zeleneva, A. Kazurova, T. Kaplienko, M. Kasian, T. Kolpakova, S. Kornienko, M. Kotsur, V. Liovin, A. Parkhomenko, M. Poliakov, S. Serdyuk, S. Skrupsky, A. Stepanenko, E. Tereshchenko, T. Fedoronchak). There are three Senior Lecturers (L. Deinega, O. Kachan, Ye. Fedorchenko), seven Assistant Professors (Ya. Zalubovsky, A. Timenko, D. Kavrin, Zh. Kamins'ka, O. Korotky, S. Leoshchenko, A. Tulenkov), three Heads of laboratories (M. Andreev, M. Kalinina, T. Kashina), three Senior Laboratory Assistants (A. Belova, V. Kostyuk, Ye. Vinnyk),

two Laboratory Assistants (S. Liovkina, E. Polyansky), and one Engineer (R. Shloma) working at the Department of Software Tools. The postgraduate students are also involved in the educational process at the department.

The department has two multimedia lecture rooms and seven modern computer laboratories and classrooms equipped with unique valuable equipment obtained through international and national educational and scientific projects. The material and technical base, educational, methodological and information support are updated annually.

The teachers of the department carry out significant educational, methodological and research work. The department carried out research work under the guidance of Professor S. Subbotin, Professor V. Dubrovin, Professor G. Tabunshchyk, Associate Professor V. Liovin, and others. The main areas of scientific activity of the department staff are intelligent decision support systems and mathematical modeling in the problems of technical and biomedical diagnostics, in particular, based on artificial neural and neuro-fuzzy networks, using digital signal and image processing methods, methods optimization. Programming of intelligent embedded and robotic systems has become a new area of research in recent years. During the period of existence of the specialties of the department, its employees, graduate students and students took part in the implementation of more than thirty research projects.

The training of scientific and pedagogical personnel is an integral part of the work of the department, in which much attention is paid to the training of the young generation of scientists. Since 2016, a full-time postgraduate course has been operating at the Department of Software Tools for the preparation of Doctors of Philosophy (Ph.D) in specialty 122 "Computer Science" (Guarantor of the specialty is S. Subbotin). In previous years, Post-graduate students and applicants of the department prepared dissertations in the following specialties: 05.13.06 - information technology, 05.13.03 - systems and processes of control, 05.13.23 - systems and means of artificial intelligence, etc. During the existence of the department since 2002, the alumni of specialties, staff, post-graduate students and applicants of the Department of Software Tools defended four doctoral habilitation dissertations (S. Subbotin, A. Oliinyk, M. Polyakov, G. Shylo) and eighteen Ph.D thesis (G. Tabunshchyk, S. Subbotin, T. Kiprich, A. Stepanenko, A. Oliinyk, O. Oliinyk, O. Shitikova, T. Zaiko, V. Liovin, T. Fedoronchak (Yur), Ye. Gofman, G. Nelasa, T. Kaplienko, A. Sogorin, N. Myronova, T. Kolpakova, K. Vodolazkina, Yu. Tverdokhlib).

For the scientific work "Methods and algorithms for the synthesis of neural network diagnostic and predictive models" S. Subbotin was awarded by the 2004 Prize of the President of Ukraine for young scientists by Decree of the President of Ukraine No. 1451/2004 dated 9.12.2004 on the basis of the submission of the Committee on State Prizes of Ukraine in the field of science and technology. For the work "Progressive technologies and software for modeling, optimization and intelligent automation of the manufacturing and operation of new generation aircraft engine parts" in 2010 S. Subbotin and A. Oliinyk was awarded by the Prize of the Parliament of Ukraine in the field of fundamental and applied research and scientific and technical developments (Resolution of the Parliament of Ukraine No. 2876-VI of December 23, 2010). S. Subbotin (in 2001), G. Tabunshchyk (in 2006), A. Oliinyk (in 2009), and O. Oli-

inyk (in 2010) were awarded by stipendiums from the Cabinet of Ministers of Ukraine.

Teachers, graduate students and students of the department take part in international, all-Ukrainian and university scientific and scientific-practical seminars, conferences, and schools. During the existence of the department since 2002, the teachers of the department have prepared and published more than thirty scientific monographs and educational textbooks, published more than a thousand articles in scientific journals and papers at conferences and workshops, received dozens of Ukrainian patents and certificates of copyright registration for computer programs.

The department annually hosts an international scientific workshop "Computer Modeling and Intelligent Systems", which papers are indexed by the Scopus scientometric database, and provides the editorial work of the scientific journal "Radio Electronics, Computer Science, Control", indexed by the Web of Science scientometric database and included in the lists of scientific professional periodicals of Ukraine (the highest category "A") and Poland.

Students of the department engaged in scientific work, under the guidance of teachers, take part and win prizes at university, regional and all-Ukrainian Olympiads and competitions of scientific works, make presentations at seminars and conferences. In particular, at the All-Ukrainian competition for the best scientific work in the nomination "Informatics" (2006, Sevastopol), the student An. Oliinyk (supervisor S. Subbotin) won. Students A. Oliinyk, Ye. Kuznetsov, O. Oliinyk, D. Nezhevenko, C. Boichenko, G. Ivanov, A. Larionov, and D. Panin have been winners of different years of regional competitions of the Zaporizhzhia Regional State Administration for gifted youth and are awarded with prizes and diplomas (supervisor is S. Subbotin). Student I. Lymarev have been the winner of the regional (2018, supervisor is S. Subbotin) and two all-Ukrainian competitions (2018, supervisor is V. Dubrovin; 2019, supervisor is S. Subbotin) for the best student research work.

The department is doing a lot of work to prepare university students for participation in programming contests, including the ACM / ICPC World Command Programming Championship, which is the largest student programming competition. Classes with students in different years were conducted by G. Tabunshchik, S. Subbotin, G. Nelasa, N Myronova, O. Kachan, L. Deynega, Zh. Kamins'ka. The most outstanding results are: the victory in the II stage (Southern region) and reaching the semi-final of the world student programming ACM-ICPC Olympiad, the Final of Ukraine in the Major league in 2013, 3rd place in the Final of Ukraine in 2012, 1-3 places at the regional stage (annually in different years), III place at the KPI-OPEN Olympiad in 2014.

Young scientists of the department are actively involved in the work of the Council of Young Scientists and Specialists of the University. At present, the Council is headed by post-graduate student S. Leoshchenko, earlier the council was headed by S. Subbotin.

The department pays considerable attention to cooperation with leading foreign and Ukrainian partners: Vimeo (Livestream), SoftServe, Noosphere, Freshcode, Computools, QATestLab, Light IT, Group BWT, Extrums, Infocom Ltd, PowerCode, CHI Software, InCode, Kozak Group, Squro, Natife, Apriorit, Multinetworks srl, MPA

Group srl, AP Consulting, Ukrosoft, Open Geeks Lab, CookieDev, S-Pro, Absolute Web, DoneIT, Drum`N`Code, Koitechs, Logix ITS, Lumighost, Revolut ltd., Rezet, Uran, Urich, Upwork, Behance, Boosta, Zakaz.ua, 4DPipeline, Affiliate Dragons, Albedo, Alvakos, BestWebSoft, Binance, Coinmarket, CubeX, CUDEV, DataXDev, Digital Sight, DnC, Evernetica, GBKSOFIT, Koitechs, Magora Systems , Mobidev, Neosight, Phonexa, Pitch, Qwerty Software, SquadMind, Svitla Systems, Visartech, Volta, Weetteam, WhiteBit, X-Torus, Carlsberg Ukraine, ZalK, DTEK, Iskra, Zaporizhstal, Motor Sich, Privat Bank, UkrGrafIT, etc., on the basis of which students of the department undergo internships and alumni of the department are employed. A wide variety of enterprises where students train and graduates work indicates a high demand in the labor market for specialists in the specialties of the Software Tools Department and the opportunity for students of the department to easily find work in the region, in the country, and worldwide.

The Department of Software Tools together with the Public Organization "Ecosens" and the Public Organization "Foundation for Educational and Scientific Initiatives" are working to attract schoolchildren of the city of Zaporizhzhia to scientific work and further education at the university within the framework of the "SmartKids School of Robotics and Programming". Teachers of the department actively participate in the "Night of Science" and "Open Day" of the university.

The department participates in educational projects of the international programs Tempus and Erasmus+ of the European Union: "Cross-domain competencies for healthy and safe work in the 21st century" (Ref. No. 619034-EPP-1-2020-1-UA-EPPKA2-CBHE-JP, 2021-2023), "Innovative Multidisciplinary Curriculum in Artificial Implants for Bio-Engineering Bsc/Msc Degrees" (Ref. No. 586114-EPP-1-2017-1-ES-EPPKA2-CBHE-JP, 2017-2020), "Internet of Things: Emerging Curriculum for Industry and Human Applications" (Ref. No. 573818-EPP-1-2016-1-UK-EPPKA2-CBHE-JP, 2016-2019), "Centers of Excellence for Young Researchers" (Ref. No. 544137-tempus-1-2013-1-sk-tempus-jphes, 2013-2017), "Development of Embedded System Courses with implementation of Innovative Virtual approaches for Integration of Research, Education and Production in UA, GE, AM" (Ref. No. 544091-tempus-1-2013-1-be-tempus-jpcr, 2013-2016), "Industrial Cooperation and Creative Engineering Education based on Remote Engineering and Virtual Instrumentation" (Ref. No. 530278-tempus-1 -2012-1-de-tempus-jphes, 2012-2015), "Practice-Oriented Master Programs in Engineering in Russia, Uzbekistan and Ukraine" (Ref. No. 510920-tempus-1-2010-1-de-tempus-jpcr , 2010-2013), "EU-UA Master Degree in Software Engineering" (Ref. No. JEP_26182_2005, 2006-2009), DAAD project "Virtual Master Cooperation Data Science", as well as Erasmus+ KA1 academic mobility (Belgium, Germany, Romania, Spain, Poland, Czech Republic, Austria) and DAAD (Germany). This make it possible to improve the qualifications of the teachers of the department, who underwent internships at leading universities in Europe and the world, to prepare curricula and courses in accordance with world requirements, to receive the latest computer laboratory equipment, as well as licensed software.

The main purpose of the CMIS-2022 workshop is a discussion of the recent research results in all areas of Computer Sciences and Information Technologies.

The workshop is soliciting literature review, survey and research papers including, whilst not limited to, the following areas of interest:

- Artificial Intelligence (Pattern Recognition, Machine Learning, Computational Intelligence (Neural Networks, Fuzzy and Hybrid Systems, Intelligent Optimization and Search), Knowledge-based Systems);
- Computer Science Methods for Modeling;
- Info-Communication Technologies for Data Analysis (Data Mining, Decision Support, Signal and Image Processing) and Applications for Engineering, Medicine and Education.

There CMIS-2022 workshop countries of PC members and authors of submitted papers are: Belgium, Estonia, Germany, India, Israel, Kazakhstan, Netherlands, Poland, Portugal, Slovakia, Spain, Ukraine, United Kingdom, and United States of America. The number of countries of CMIS-2022 authors and PC members is 14 (instead of 21 in 2021, and 12 in 2020).

The language of CMIS-2022 Workshop is English. The workshop took the form of oral presentations of peer-reviewed regular papers. The video of presentations is available at <https://www.facebook.com/groups/cmishorkshop/>.

The CMIS-2022 Organizing Committee have received 53 paper submissions (instead of 122 submissions in 2021, 177 submissions in 2020, and 130 submissions in 2019), out of which 25 were accepted for presentation as a regular papers (instead of 44 in 2021, 83 in 2020, and 81 in 2019) in the result of reviewing process, where the editor, program committee members, external and peer reviewers sought to ensure a high scientific level, as well as a high-quality presentation of the selected papers.

The paper acceptance ratio in 2022 is 47% (instead of 39% in 2021, 47% in 2020, and 62% in 2019). These papers were published in this volume of CMIS-2022 proceedings.

The workshop would not have been possible without the support of many people. I would like to thank the Workshop Program Committee Members, organization staff, and external reviewers for their excellent and tireless work. I sincerely wish that all attendees benefited scientifically from the workshop and wish them every success in their research. It is the humble wish of the workshop organizers that the professional dialogue among the researchers, scientists, and engineers continues beyond the event and that the friendships and collaborations forged will linger and prosper for many years to come.

May, 2022

Sergey Subbotin, Dr. Sc., Prof.,
CMIS-2022 PC Chair and Proceedings Editor
Head of the Department of Software Tools,
National University "Zaporizhzhia Polytechnic"

Programme Committee

Chair

Prof. **Sergey Subbotin**, National University "Zaporizhzhia Polytechnic", Ukraine

Members

Prof. **Gustavo Alves**, University of Porto, Portugal

Prof. **Yevgeniy Bodyanskiy**, Kharkiv National University of Radio Electronics, Ukraine

Prof. **Karsten Henke**, Ilmenau University of Technology, Germany

Prof. **Igor Korobiichuk**, Warsaw University of Technology, Poland

Prof. **Iurii Krak**, Taras Shevchenko National University of Kyiv, Ukraine

Prof. **Vitaly Levashenko**, University of Zilina, Slovak Republic

Prof. **David Luengo**, Universidad Politecnica de Madrid, Spain

Prof. **Andrii Oliinyk**, National University "Zaporizhzhia Polytechnic", Ukraine

Prof. **José Otegi**, University of the Basque Country, Spain

Prof. **Nataliya Pankratova**, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Ukraine

Prof. **Mykhailo Poliakov**, National University "Zaporizhzhia Polytechnic", Ukraine

Prof. **Anatoliy Sachenko**, Western-Ukrainean National University, Ukraine

Prof. **Natalya Shakhovska**, Lviv Polytechnic National University, Ukraine

Prof. **Galyna Tabunshchyk**, National University "Zaporizhzhia Polytechnic", Ukraine

Prof. **Carsten Wolff**, Dortmund University of Applied Sciences and Arts, Germany

Prof. **Elena Zaitseva**, University of Zilina, Slovak Republic

Dr. **Peter Arras**, Katholieke Universiteit Leuven, Belgium

Dr. **Ivan Izonin**, Lviv Polytechnic National University, Ukraine

Dr. **Anzhelika Parkhomenko**, National University "Zaporizhzhia Polytechnic", Ukraine

Dr. **Alexei Sharpanskykh**, Delft University of Technology, Netherlands

Dr. **Thomas (Tom) Trigano**, Shamoon College of Engineering, Israel

Dr. **Heinz-Dietrich Wuttke**, Ilmenau University of Technology, Germany

Organizers



National University "Zaporizhzhia Polytechnic",
Zaporizhzhia, Ukraine

<http://www.zp.edu.ua>



Department of Software Tools
of the National University "Zaporizhzhia Polytechnic",
Zaporizhzhia, Ukraine

<http://pz.zntu.net>

Local Organizational Committee

Prof. **Sergey Subbotin**, National University "Zaporizhzhia Polytechnic", Ukraine

MSc **Serhii Leoshchenko**, National University "Zaporizhzhia Polytechnic", Ukraine

Maxim Andreev, National University "Zaporizhzhia Polytechnic", Ukraine

Partners

Co-funded by the
Erasmus+ Programme
of the European Union



WORK4CE

International project "Cross-domain competences for healthy and safe work in the 21st century"(WORK4CE) of Erasmus+ Programme co-funded by the European Union (Ref. No. 619034-EPP-1-2020-1-UA-EPPKA2-CBHE-JP)

DAAD



ViMaCs

Virtual Master Cooperation
Data Science

International project «Virtual Master Cooperation Data Science (ViMaCs)» of DAAD

Sponsors



<https://freshcodeit.com>
<https://freshcodeit.jobs>
<https://jobs.dou.ua/companies/freshcode/>
<https://freshcode.training/>
<https://www.instagram.com/freshcodeit/>
<https://www.facebook.com/freshcodeit/>



<https://groupbwt.com/>
<https://jobs.dou.ua/companies/groupbwt/>

Reviewers

Name and email	Number of reviews
Andrii Kopp <kopp93@gmail.com>	3
Andriy Goncharenko <andygoncharenco@yahoo.com>	3
Andriy Palamar <palamar.andrij@gmail.com>	3
Artem Sokolov <radiosquid@gmail.com>	3
Bohdan Sus <bnsuse@gmail.com>	3
Dmitriy Klyushin <dokmed5@gmail.com>	3
Dmytro Orlovskyi <orlovskyi.dm@gmail.com>	3
Dmytro Ostrovka <ostrovkadk@gmail.com>	1
Emil Faure <e.faure@chdtu.edu.ua>	3
Eugene Fedorov <fedorovee75@ukr.net>	3
Gennadii Martynenko <gmartynenko@ukr.net>	6
Halyna Osukhivska <osukhivska@tntu.edu.ua>	3
Hlib Horban <gleb.gorban@gmail.com>	3
Igor Povkhan <igor.povkhan@uzhnu.edu.ua>	3
Iryna Myronets <i.myronets@chdtu.edu.ua>	3
Iryna Pidkurkova <i.v.pidkurkova@nlu.edu.ua>	3
Iryna Svyd <iryna.svyd@nure.ua>	3
Ivan Obod <ivan.obod@nure.ua>	3
Karina Melnyk <karina.v.melnyk@gmail.com>	3
Lesia Mochurad <lesiamochurad@gmail.com>	3
Maxim Davidovsky <m.davidovsky@gmail.com>	3
Mykhailo Dvoretzkyi <m.dvoretzkyi@gmail.com>	3
Mykhaylo Palamar <palamar.m.i@gmail.com>	3
Nataliya Boyko <nataliya.i.boyko@lpnu.ua>	3
Nickolay Rudnichenko <nickolay.rud@gmail.com>	2

Name and email	Number of reviews
Oksana Mulesa <Oksana.Mulesa@uzhnu.edu.ua>	4
Oksana Pichugina <oksanapichugina1@gmail.com>	2
Oleksandr Bauzha <asb@univ.kiev.ua>	3
Oleksandr Hres <o.hres@chnu.edu.ua>	2
Oleksandr Vorgul <oleksandr.vorgul@nure.ua>	1
Olena Melnyk <olena.melnyk@uzhnu.edu.ua>	3
Olga Nechyporenko <olne@ukr.net>	3
Olha Matsyi <olga.matsiy@gmail.com>	4
Pavlo Radiuk <radiukpavlo@gmail.com>	3
Pavlo Skladannyi <p.skladannyi@kubg.edu.ua>	3
Sergey D. Bushuyev <SBushuyev@ukr.net>	3
Sergiy Zagorodnyuk <szagorodniuk@gmail.com>	3
Serhii Choporov <s.choporoff@znu.edu.ua>	3
Serhii Leoshchenko <sergleo.zntu@gmail.com>	1
Serhii Vladov <ser26101968@gmail.com>	3
Svitlana Gadetska <svgadetska@ukr.net>	3
Taras Chaikivskyi <Taras.V.Chaikivskyi@lpnu.ua>	2
Taras Lendiuk <tl@wunu.edu.ua>	3
Tetiana Korobeinikova <tetianakorobeinikova@gmail.com>	3
Tetiana Shmelova <shmelova@ukr.net>	3
Vitalii Larin <vjlarin@gmail.com>	3
Vladimir Semenov <semenov.volodya@gmail.com>	3
Vladyslav Kuznetsov <kuznetsow.wlad@gmail.com>	3
Volodymyr Gorokhovatskyi <gorohovatsky.vl@gmail.com>	3
Volodymyr Rusyn <rusyn_v@ukr.net>	2
Yevhen Davydenko <davydenko@chmnu.edu.ua>	3
Yurii Kryvenchuk <yurkokryvenchuk@gmail.com>	3

The problem of convergence of classifiers construction procedure in the schemes of logical and algorithmic classification trees

Igor Povkhan^a, Oksana Mulesa^a, Olena Melnyk^a, Yuriy Bilak^a, Volodymyr Polishchuk^a

^a *Uzhhorod National University, Zankoveckoy str., 89B, Uzhhorod, 88000, Ukraine*

Abstract

The paper considers the problem of convergence in the procedure of classifier schemes synthesis by methods of logical and algorithmic classification trees. It suggests the upper evaluation of complexity for the scheme of algorithms tree in the problem of approximation of real data array by a set of generalized features with a fixed criterion of termination of the branching procedure at the stage of the classification tree construction. This approach provides the required accuracy of the model, evaluates its complexity, reduces the number of branches, and achieves the required performance features. The paper represents the evaluation of convergence for the procedure of recognition schemes construction is logical and algorithmic classification trees on conditions of weak and robust separation of primary initial sampling classes.

Keywords

logical tree, algorithmic tree, classifier, pattern recognition, feature, initial sample.

1. Introduction

Problems united by pattern recognition are various and currently widely spread in both the economic and social content of the human activity, leading to the necessity of building and investigating the mathematical models of the relevant systems [1-5]. A universal approach to their solution is still missing, while plenty of general theories and approaches help solve different problem types (classes). However, their practical application varies by high sensitivity to the specific parameters of the problem itself or subject area of application [6]. Various theoretical results derive from exceptional cases and subproblems, pointing out that the bottleneck of successful real recognition systems requires performing a considerable amount of computation and focusing on powerful hardware tools. The classification trees (solution trees) have fixed a significant part of the above shortcomings. It enables to work effectively with problems of arbitrary scale data (where the information is set in natural form) [7-10].

Today there are various relevant approaches to constructing classification systems in the form of logical trees and algorithmic classifications (LCT/ACT) [11-13]. Moreover, the interest in recognition methods using LCT is caused by several valuable properties. From one side, the complexity of the class of recognition functions (RF) in the form of LCT models, under certain conditions, does not exceed the complexity of the class of linear recognition functions (the simplest of the known). On the other side, RF in the form of classification trees allows distinguishing in the process of classification both causal links (and unambiguously take them into account in the future) and factors of chance or uncertainty. Additionally, it considers both functional and stochastic links between properties and

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL igor.povkhan@uzhnu.edu.ua (I. Povkhan); oksana.mulesa@uzhnu.edu.ua (O. Mulesa); olena.melnyk@uzhnu.edu.ua (O. Melnyk)
yuriy.bilak@uzhnu.edu.ua (Y. Bilak); volodymyr.polishchuk@uzhnu.edu.ua (V. Polishchuk)
ORCID: 0000-0002-1681-3466 (I. Povkhan); 0000-0002-6117-5846 (O. Mulesa); 0000-0001-7340-8451 (O. Melnyk);
0000-0001-5989-1643 (Y. Bilak); 0000-0003-4586-1333 (V. Polishchuk);



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

behavior of the entire system and the process of missing interactions classification between environmental objects, different animal species, and people (except the objects, information about which is transmitted genetically (hereditary) or other). It takes place according to the so-called logical decision tree [14].

In most forecasting and classification problems, which use unstructured data (such as sets of discrete patterns or text arrays), an artificial neural network (of selected type) outperforms all other types of algorithms or decision tree frameworks. Otherwise (in the case of structured high-volume discrete data arrays), the methods and algorithms of the decision tree concept are vastly preferred [15-17]. Practically, LCT algorithms and construction methods usually give structurally complex logical trees at the output (in terms of the number of vertices, number of branches belonging to the class of irregular trees), which are unevenly filled with data and have different numbers of branches. Such complex tree-like structures are difficult to perceive for external analysis due to a massive number of nodes (vertices) and a massive number of stepwise partitions of the primary initial sample (IS) containing a minimum number of objects (maybe even single objects in the worst case). In practice, let us review the fundamental question regarding methods of classification trees (classification models) – the question of convergence of procedure for constructing a classification tree (methods of classification trees), structures LCT/ACT.

2. Formal problem statement

Let it be the IS in the following form:

$$(x_1, f_R(x_1)), \dots, (x_m, f_R(x_m)) \quad (1)$$

Consider $x_i \in G, (i = 1, \dots, m)$, where m – is the quantity of objects with the IS, $f_R(x_i)$ – is a finite value function that designates the partition of R for the set G into classes (patterns) $H_0, \dots, H_{k-1}$. Ratio $f_R(x_i) = l$ ($l = 1, \dots, k-1$) means $x_i \in H_l$, $x_i = \{x_{i_1}, \dots, x_{i_n}\}$, x_{i_j} – value j – kind features for object x_i , ($j = 1, \dots, n$), n – quantity of features in IS.

As a result, IS is a population (or the sequence) of the set, where each set is a population of indicated values and functional values [18]. Additionally, the set of indicated values is a particular image, where the value of the function relates it to the particular pattern. The problem is to build a structure LCT/ACT – L based on the array of primary IS base of the type (1) and determine the value of its structural parameters p (that is $F(L(p, x_i), f_R(x_i)) \rightarrow opt$).

3. Literature review

All basic approaches in the theory of recognition have their advantages and disadvantages and form a single toolkit for solving practical problems of artificial intelligence theory — especially the integrally elaborated classical algebraic approach developed by Y.I. Zhuravlyov [19]. This direction of recognition theory development is connected with the classification algorithms model construction and the optimized quality recognition algorithm selection. The research focuses on the actual concept of decision trees (LCT structures). In particular, The classification scheme, which is given by an arbitrary approach and method and algorithm of the classification tree, has a tree-like logical structure [1,3,20]. Moreover, the structure of the logical tree consists of vertices (features), grouped by circles and built (selected) at a certain step (stage) of classification tree model construction [21]. The main peculiarity of tree-like recognition systems is that the importance of individual features (groups of features or sets of features) is determined relative to the function that defines the division of objects into classes [22].

In the research [21], a generation scheme for the classification tree structure is suggested based on the step-by-step selection of elementary attributes, the disadvantage of which is the high dependence of the model complexity on the effectiveness of the final minimization, tree cutting procedure. In

researches [23-25], the modular scheme of classifiers construction has been offered in the form of classification tree structures, which circumvents the limitations of traditional decision tree methods. The research [24] proposes an efficient scheme for generating generalized features based on constructing sets of geometric objects. However, the disadvantage is the limitation of the initial training structure of the sample and the lack of practical application versatility. In addition, challenges of LDK models structural complexity estimating the minimization stage evaluated in research [20].

Thus, the research [22] shows that the resulting classification rule has a tree-like logical structure built by an arbitrary method or the indication branched selectional algorithm. In this case, the question of qualitative branching criterion selection comes first. The logical tree consists of vertices grouped on tiers and received at a particular step of recognition tree construction [25]. Thereby, it challenges the effective structure minimization for the constructed model of the classification tree. A significant problem that arises from article [23] is the synthesis of recognition trees, which the actual tree of algorithms will represent. An essential area of LDK structures research remains the issues of the generation of decision trees for the case of uninformative features [14] and the actual issue of classification trees theory. It questions the possibility of constructing all logical tree variants, which correspond to primary IS, and selecting minimum depth or structural complexity (number of tiers) classification tree [26-30].

4. Convergence of synthesis of logical and algorithmic classification tree structure

Consider that at every step during building a logical tree (some LCT model), only one selected elementary feature is picked from the set of the fixed features $(\varphi_1, \dots, \varphi_n)$. Then on $n - \text{th}$ step of the tree classification construction procedure, the LCT scheme represents by itself a predicate p_n (generalized feature, which is built from the set of elementary features) [23, 30], that is the most effective approximation of the initial IS of general view (1) (applicable for the ACT structure case).

In particular, p_n represents some tree-like scheme (classification tree), which consists of n vertices. The structure of the predicate p_n includes only n elementary features (attributes of the IS discrete object) from the initial set.

The sequence of predicates $p_1, \dots, p_j$ (generalized features) coincides with the primary IS of the (1) kind. In case it starts from some Q , the following condition fulfills:

$$f_{Q+m} = f_R(x_i), (i = 1, \dots, m), (m \geq 0).$$

Let us denote with φ_n some elementary features selected (fixed) on $n - \text{th}$ step in the scheme of the LCT model construction. The feature φ_n corresponds to some fixed path $r_1, r_2, \dots$, which ends by the given attribute (the vertex of classification tree – LCT model). For example, on Fig. 1 there is the LCT in which the vertex φ_2 (feature) corresponds to the path $\{0\}$, and the vertex φ_5 corresponds to the path $\{0,1\}$.

The path that corresponds to the elementary sign φ_n as indicated, let us denote as T_n , and with D_n let us denote the set of those pairs $(x_i, f_R(x_i))$ of the primary IS of the general view (1), where objects w_i belong to the path T_n . For example, for the LCT structure (Fig. 1), let $\varphi_n = \varphi_4$, then the path T_n looks as $\{1,0\}$.

In such a case some object w_i belongs to the path $\{1,0\}$, if the following conditions are fulfilled: $\varphi_1(w_i) = 1$ and $\varphi_1(w_i) = 0$.

We consider further, that the elementary feature φ_n weakly separates the set D_n , if in D_n there are such pairs $(x_i, f_R(x_i))$ and $(x_j, f_R(x_j))$, that $\varphi_n(x_i) = 0$ and $\varphi_n(x_j) = 1$ (that is $\varphi_n(x_i) \neq \varphi_n(x_j)$).

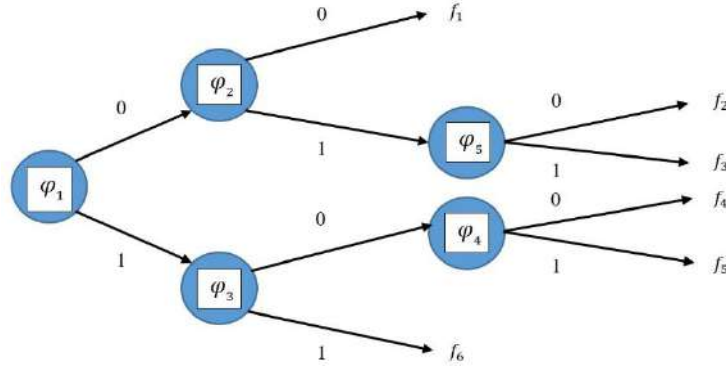


Figure 1: LCT structure with elementary features as vertices.

In the final power of the scheme of classification tree method (models LCT/ACT), we call the number of all finite vertices (certain sheets) of the scheme. For example, for the LCT on Fig. 1 the power is 6.

Obviously, the final power of the classification tree method scheme is also equal to the number of all final paths in it. It is clear that with induction of n , it can be easily proved that the final power of each of the above schemes p_n (predicates) is equal to $n + 1$. Indeed, the fact that the final power p_1 which includes only one feature or algorithm (cases with LCT/ACT), equals 2, is obvious.

Let the final power of the scheme p_n equals to $n + 1$. Let us calculate the final power p_{n+1} . It is clear that this scheme is based on the scheme p_n , when in some final vertex, a new vertex is successively added (feature, algorithm) with the number $n + 1$. Obviously, when adding this feature (algorithm) to the scheme p_n one end vertex disappears and two new end vertices are added. Therefore, we can conclude that the number of all end vertices of the scheme p_n equals to $n + 2$.

Let us assume that on each n -th step of the classification tree construction procedure (of the LCT model) the set D_n is weakly separated by some feature φ_n . Next, we consider the scheme p_n . According to this scheme, it has $n + 1$ of the final paths. Due to the fact that D_n at each step is weakly separated, every such path contains at least one pair of the primary IS of general view (1). It is also obvious that different end paths in p_n do not have common pairs from the sample (1). Therefore, we can conclude that the scheme (predicate) p_n divides IS (based on the basic branching criterion introduced by the current method of classification tree) in $n + 1$ non-empty parts (subsets) that do not intersect. Since in the primary IS in total there are m training pairs, the scheme p_{m-1} (or a predicate with a smaller number) completely separates the primary IS, that is p_{m-1} fully recognizes the sample.

Thereby, if on each n -th step the selected elementary feature φ_n weakly separates the set D_n , then in this case the LCT construction process coincides with the primary IS and ends in no more than in $m - 1$ steps, where m - the number of all training pairs of primary IS.

Let us notice that the condition of weak separation of classes for primary IS is relatively weak - therefore, it provides low convergence of the construction procedure for the classification tree. Thus, it is essential to consider the convergence of the process under more vital conditions. Therefore, we will assume that we are dealing with a case where IS contains information about two classes (patterns)

H_0 and H_1 , and IS itself has a deterministic nature. Let n_j – is a number of training pairs $(x_i, f_R(x_i))$ within primary IS, which satisfy the ratio $f_R(x_i) = j, (j = 0, 1)$, and for simplification and certainty, we put it as $n_0 \geq n_1$.

Having fixed $f_R(x) \equiv 0$, we obtain some generalized feature (scheme) f_0 , which approximates (in whole or in part) the primary IS. Obviously, in this case (that is in a situation where the choice of any elementary feature has not yet been made φ_n), generalized feature (scheme) f_0 is the best approximation of primary IS. Further the value n_1 we call the unconditional number of errors in the primary IS.

Let in the first step of building a classification tree some elementary features has been selected (arbitrarily) φ_1 – and this feature will break the initial sample into two parts (subsets) H_0 and H_1 , where H_j – is the set of all pairs $(x_i, f_R(x_i))$ of primary IS, for which the ratio is satisfied $f_1(x_i) = j, (j = 0, 1)$.

Let n_m^j – the set of all pairs $(x_i, f_R(x_i))$ from sample $H_j, (j = 0, 1)$, for which the relation is fulfilled $f_R(x_i) = m, (m = 0, 1)$. Feature φ_1 can be considered a generalized feature f_1 (scheme), which is built on the first step of the LCT construction process.

Let us enter the value $\rho = \max(n_0^0, n_1^0) + \max(n_0^1, n_1^1)$, which is the number of correct answers (classifications) that are realized by a generalized feature f_1 , and accordingly, the value n_0 is the number of correct answers (classifications), which are realized by a generalized feature f_0 .

By the number of correct answers we mean the number of those training pairs $(x_i, f_R(x_i))$ in the initial training sample of the type (1), for which the relation of equality is fulfilled $f_R(x_i) = f_1(x_i)$.

Since $n_0^0 + n_0^1 = n_0$ and $n_1^0 + n_1^1 = n_1$, then we will have the following:

$$\rho = \max(n_0^0, n_1^0) + \max(n_0^1, n_1^1) \geq n_0. \quad (2)$$

Therefore, when choosing a feature φ_1 the number of correct answers at least does not decrease. The number of errors given by the generalized algorithm f_1 , will be equal to:

$$m - n = n_1 - (\rho - n_0) \leq n_1. \quad (3)$$

Let us note that (3) follows from (2). Let us enter the value $\lambda_1 = \frac{n_1}{m - \rho}$ and call it the quality of elementary feature φ_1 relating to primary IS, similarly we determine λ_n of feature φ_n relating to primary IS ($n = 1, 2, 3, \dots$).

With the power of some constructed generalized feature (GF) or set GF (for a fixed step of ACT scheme) we will call the number of training pairs $(x_i, f_R(x_i))$ of primary IS that look like (1), which is approximated (correctly classified) by this generalized feature (sequence of generalized features).

It is important for ACT schemes that at step-by-step division of IS on two samples H_0 and H_1 (and so on) part of the sample will be completely covered by the current classification algorithm (generalized feature or their set) – that is, we will have a case of strong separation of array IS classes. Therefore, we can assume that the complexity of the final ACT scheme (the total number of steps to build a tree) largely depends on the procedure of initial evaluation and selection of a set of independent classification algorithms a_i , their initial parameters, parameters of set GF f_i , which they generate for each step of the ACT scheme.

Then, for the ACT scheme, it is essential to consider the general complexity of the procedure for constructing a classification tree in the condition of weak separation of the primary IS classes. A single GF is generated with the power of one unit per verticle of the tree. Under high separation

conditions, when a problem and practical feasibility have not set the number limit on GF and its power, it is possible to build them.

In the first stage, we consider a case of weak class division with restrictions on the GF sets being built by the ACT scheme. Let us note that the procedure for constructing an algorithmic tree has certain features in terms of step-by-step approximation of the primary IS by sequence GF. Let at every step of the construction of some ACT model, select for work one fixed classification algorithm from a set of selected algorithms $(a_1, a_2, \dots, a_n)$. Moreover, the classification tree can be built by one algorithm and sequence GF, which he is generating.

Therefore, after fulfilling the n steps of the classification tree constructing procedure the ACT structure represents some scheme s_n (generalized second-order feature, which is built from a set of GF synthesized by classification algorithms), which is the most effective approximation of the primary IS of general view (1) a set of independent classification algorithms and their GF. In particular s_n will represent some tree-like scheme (GFT structure), which consists of n vertices, that is, in the design of the scheme s_n enters only n classification algorithms (GF – conditioned that for each step of the tree constructing procedure no more than one generalized feature of the minimum power per unit is generated) from the initial set.

5. Experiments and results

In the next stage for the LCT structure let us make an assumption – quality λ_n of elementary features φ_n relative to the array of primary IS not less than some number y , where $y > 1$.

Let us analyze the complexity of the procedure for classification tree construction under this condition ($y > 1$), to do this, let us estimate the number of steps for which this process (procedure) will implement full recognition of the initial training sample array.

To be certain, let us consider the following scheme of classification tree construction (Fig. 2).

Let n_1 – an unconditional number of errors within primary IS. Elementary feature φ_1^1 separated IS in two samples: H_0 and H_1 . Let h_0 and h_1 accordingly is an unconditional number of errors in the samples H_0 and H_1 . Feature φ_1^2 separate the set H_0 in two sets H_{00} and H_{01} . Let h_{00} and h_{01} - is an unconditional number of errors in the samples H_{00} and H_{01} . Similarly, we define sets H_{10}, H_{11} and quantities h_{10} and h_{11} for the elementary feature φ_2^3 .

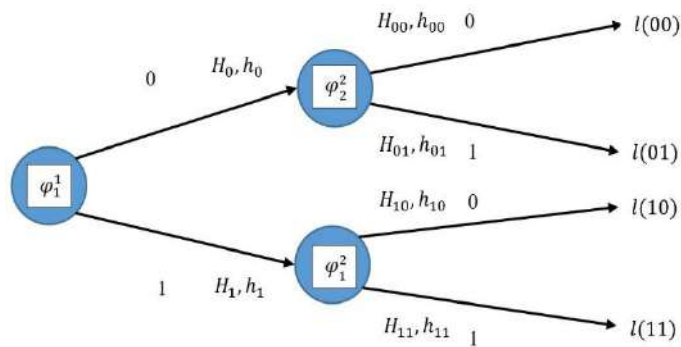


Figure 2: Scheme of division into subsets in the structure of the classification tree.

From the initial condition ($y > 1$) it follows:

$$\begin{cases} h_0 + h_1 \leq \frac{1}{y} * n_1 \\ h_{00} + h_{01} \leq \frac{1}{y} * h_0 \\ h_{10} + h_{11} \leq \frac{1}{y} * h_1 \end{cases} \quad (4)$$

From (4) we receive the following:

$$h_{00} + h_{01} + h_{10} + h_{11} \leq \frac{1}{y^2} * n_1. \quad (5)$$

Let us make the following assumptions in this regard: $h_0 \geq 1$, $h_1 \geq 1$, $h_{00} \geq 1$, $h_{01} \geq 1$, $h_{10} \geq 1$ and $h_{11} \geq 1$. From here we will have the following:

$$2^1 \leq \frac{1}{y} * n_1, \quad 2^2 \leq \frac{1}{y^2} * n_1. \quad (6)$$

Similarly for the set of features $\varphi_1^i, \varphi_2^i, \dots$, located on i -th tier of the logical tree, we will have:

$$2^i \leq \frac{1}{y^i} * n_1 \text{ or } (2y)^i \leq n_1. \quad (7)$$

Hence we can conclude that the process of constructing a classification tree will continue until there the structure of will have m tiers (levels), where m has the following form:

$$m = R\left(\frac{\log_2 n_1}{1 + \log_2 y}\right). \quad (8)$$

Under $R(x)$ means rounding the number x to the nearest integer, which exceeds x . For example $Q(1.2) = 2, Q(3.7) = 4, Q(4.1) = 5$.

Hence the classification tree, which has m full tiers (that is, the case when on i -th tier there are vertices), has vertices – thus recognizing the primary IS with condition that ($y > 1$) by means of full LCT takes place no more than in steps, where m is calculated using an expression (8).

For the case of structure ACT, we can conclude that the sequence of constructed schemes $s_1, s_2, \dots, s_j$ (generalized features of the second order) coincides with the primary IS with the view (1), not more than in M steps (where M – total power of primary IS), even if at each step only one GF is generated, while the power of each is not more than one unit.

Some classification algorithm that will be selected (fixed) on n -s th step in the procedure of building the ACT model (to generate the appropriate GF), let us denote with a_n , and it is clear, that this algorithm a_n corresponds to some scheme s_n , which consists of algorithms $a_1, a_2, \dots, a_{n-1}$ and ends with this attribute (vertice of classification tree - ACT model). For example, on Fig. 3 some ACT model is shown, in which a fixed scheme s_2 (vertice of the classification tree under construction) corresponds to the sequence of steps (schemes) $\{s_1\}$, and scheme s_M – sequential path $\{s_1, s_2, \dots, s_{M-1}\}$.

Therefore, for the ACT model it is possible to make the following conclusion: scheme s_n (in the structure of the classification tree) divides IS on n non-empty parts (subsets) that do not intersect, and because in the primary IS in total there are M training pairs, that is why the scheme s_M will completely divide (approximates) the primary IS (that fully recognizes the sample at the condition that it generates one GF at each step with the power of one unit). So, if on every n -th step of the ACT construction scheme the generated GF (selected by classification algorithm a_n) weakly separates the set of primary IS, in this case, the process of classification tree construction coincides with the primary IS and finishes not more than in M steps, where M – is the quantity of all training

pairs of primary IS. In the following stage of research, it is important to consider the case of strong separation of classes in primary IS, when there are no set restrictions on algorithms a_i regarding generation GF (power of constructed GF is limited only by the practical possibility of the classification algorithm itself a_i and structural parameters of IS). Let with $P(f_j)$ we denote the total power (approximation ability) of corresponding GF f_j , ($1 \leq j \leq s$), where s – is the quality of GF in constructed ACT scheme.

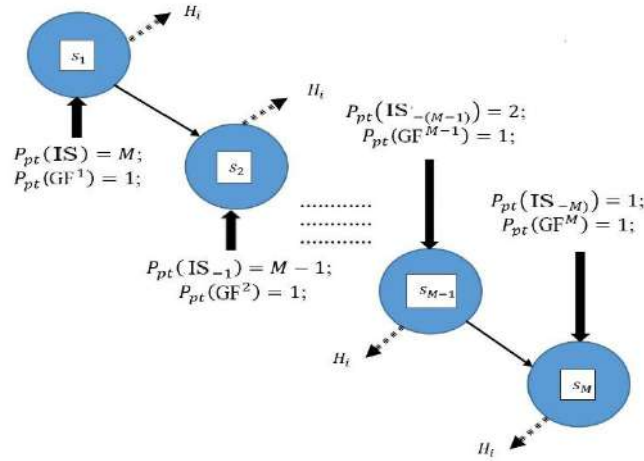


Figure 3: Example of ACT structure with GF as vertices.

Further on some step $r, (1 \leq r \leq M)$ of the ACT scheme, a sequence of generalized features is constructed $f_1, \dots, f_r$ with their corresponding values $P(f_i) = z_i$, where $(1 \leq z \leq M), (1 \leq i \leq r)$, M – general power IS, and among them there are values $z^{\max}$ and $z^{\min}$, which are for them, respectively, the maximum and minimum (relatively to the current step of ACT scheme). Then in this case the ACT scheme (model) will be constructed in t steps, where the value t is determined by the ratio (9).

$$t \leq 2 * \frac{P_{pt}(IS)}{z^{\max} + z^{\min}} = \frac{2M}{z^{\max} + z^{\min}}. \quad (9)$$

We notice that in the case when by the condition of the practical problem the power restrictions of synthesized GF are imposed on the ACT scheme under construction (not exceeding the appropriate value P) – classification tree scheme (ACT model) will be built in t steps, where the value t is determined by the ratio (10).

$$t \leq \frac{M}{P}. \quad (10)$$

At strict restrictions of the ACT scheme on one generated GF (where according to the condition $P(f_j) = 1, (1 \leq i \leq t)$, that is in case of weak classes separation of the current problem, the classification tree scheme (model ACT) will be built in t steps, where the value $t \leq M$.

In the next stage, for simplification, let us take the classification problems for which sets of structures LCT/ACT have been built from researches [19-24,30]. The initial parameters of these practical problems are presented in the Table 1. So in IS, the information regarding separation into two classes has been represented. At the examination stage, the constructed classification system should effectively recognize objects of unknown classification concerning these two classes. Let us note that the training and test sample was automatically checked for correctness at the initial stage. (searching and deleting the identical objects of different affiliations - errors of the first and second type).

Here (in Table 2) is represented the evaluation of the constructed structures of classification trees (LCT/ACT) of problems from Table 1. To assess the quality of constructed classifiers (classification schemes) we used an integrated feature of the quality of the classification tree Q_{Main} from reser [30].

Convergence of the structure synthesis LCT/ACT procedure is estimated on the basis of quantitative features – the total number of iterations S_{Main} and the number of tiers in the classification tree structure L_{Kol} .

Table 1

Initial parameters of classification problems

Type of classification problem	The dimension of the feature space N	The power of data array of the primary IS – M	The total number of classes by data splitting IS – l	Relation of objects of different classes IS – $(H_1 / H_2 / \dots / H_i)$
The problem of geological data classification (Z1)	22	1250	2	756/494
The problem of chemical analysis of the quality of hydrocarbon fuels (Z2)	14	4863	6	823/648/1412/918/583/764
The problem of classification of flood situations in the Tisza river basin of Transcarpathian region (Z3)	18	6118	3	76/108/5934
The problem of classification of flood situations in the Uzh river basin of Transcarpathian region (observation post №1) (Z4)	18	4252	3	73/102/4107
The problem of classification of flood situations in the Uzh river basin of Transcarpathian region (observation post №2) (Z5)	18	4139	3	68/97/3974

Integrated feature of the classification tree quality Q_{Main} displays the basic parameters (characteristics) of classification trees and can be used as an optimality criterion in the evaluation procedure of an arbitrary tree-like recognition scheme [20]. Let us note that the main idea of classification tree methods based on autonomous algorithms in their structure lies in a step-by-step approximation by the selected set of algorithms the data set of primary IS [22].

Table 2

Comparative table of structure classification schemes LCT/ACT

№ of problem	Method of synthesis of classification tree structure	Integral feature of model quality Q_{Main}	Convergence of the classification tree structure (number of iterations) S_{Main}	Number of tiers in structure LCT/ACT L_{Kol}
Z1	Method of complete LCT based on the selection of elementary features	0,004789	79	22
Z1	LCT model with a one-time assessment of the importance of the features	0,002263	102	16
Z2	Limited LCT construction method	0,003244	91	17
Z2	Algorithmic tree method (type I)	0,005119	46	9
Z3	Algorithmic tree method (type II)	0,002941	72	15
Z3	Method of extensive selection of features (step-by-step assessment)	0,003612	84	13
Z4	Algorithm tree (type I)	0,005054	43	10
Z5	Algorithm tree (type II)	0,002813	75	16

The obtained structures of classification trees (ACT/LCT models) are from one side are characterized by high versatility in terms of practical problems and relatively compact structure of the

model itself, but from the other side it requires significant hardware costs for storing generalized features and initial assessment of the fixed classification algorithms quality according to data of IS comparing to the neural network concept [31-35]. Therefore in comparison with the ACT concept LCT method has a high speed of classification schemes, relatively insignificant hardware costs for storage and operation of the tree structure itself, and high quality of discrete objects classification.

Let us note that all ACT / LCT models and ACT / LCT structures were built in the software application "Orion III". Based on the classification tree method and the principle of modularity, Uzhhorod National University has developed a software application "Orion III" to generate autonomous recognition systems. The algorithmic library of the system has 15 recognition algorithms, including algorithmic implementations of both the LCT structure and the ACT structure [23].

From Table 2 we can see that the methods of algorithm trees (of two types) have shown a high rate of convergence for constructing the classification tree structure on represented IS compared to the LCT schemes. Consider that the first type of the ACT structure shows a good result in terms of structural complexity (number of tiers, vertices, generalized features) of constructed classification model in comparison with logical classification trees and the tree of algorithms of the second type. In general, we can conclude about the rapid convergence of ACT structures compared to LCT models and advantage due to this the structural complexity of the constructed classification tree and the informational capacity of generalized feature sets.

6. Conclusion

Therefore, taking into consideration all the above in the research, we can assume the following points:

On the condition of weak class separation in the LCT case, if on every n -th step the selected elementary feature φ_n weakly separates the set (subset) of objects of primary IS, then, in this case, the process of classification tree construction coincides relating to the primary IS and terminates no more than in $m-1$ steps, where m – the quantity of all training pairs of primary IS.

Classification tree (of LCT structure) on condition of strong classes separation of primary IS sets of objects, which has m full tiers, levels (that is the case, when on i -th tier vertices are located), has y vertices – so array recognition of primary IS on condition that ($y > 1$) by means of full LCT takes place no more than in steps, where m is calculated by means of expression $m = R(\frac{\log_2 n_1}{1 + \log_2 y})$.

A general number of all final vertices of logical structure (sheets of recognition tree) of constructed classification scheme will unambiguously determine the final power of classification tree method scheme (models LCT/ACT).

Power of some GF (the set of constructed GF) for the fixed step of the ACT method scheme is considered the general quantity of training pairs $(x_i, f_R(x_i))$ of primary IS (a subset of primary IS) that look like (1), which are approximated (correctly classified) by given generalized feature (sequence of generalized features).

In case of weak class separation of primary IS for the ACT scheme the process of classification tree construction coincides relating to IS data set and terminates in no more than M steps, where M – is the quantity of all primary IS training pairs.

In the case of strong classes, separation of primary IS for the ACT scheme, when the power of constructed GF (or set of GF) is limited only by the practical possibility of the classification algorithm itself a_i and initial parameters of IS, the ACT scheme (model) will be constructed in t steps, where the value t is defined by the ratio (9).

7. Acknowledgments

The paper was carried out within the framework of the research "Modeling and forecasting of emergency situations in the Carpathian region and countries of Central and Eastern Europe", the state registration number of the work is 0106V00285, the category of work is fundamental research (ID-2201020), 01 – Fundamental research on the most important problems of natural, social and humanitarian sciences.

8. References

- [1] T. Hastie, R. Tibshirani, J. Friedman, *The Elements of Statistical Learning*, Berlin, Springer, 2008.
- [2] J.R. Quinlan, *Induction of Decision Trees*, *Machine Learning* 1 (1986) 81-10.
- [3] L.L. Breiman, J.H. Friedman, R.A. Olshen, C.J. Stone, *Classification and regression trees*, Boca Raton, Chapman and Hall/CRC, 1984.
- [4] M. Lupei, A. Mitsa, V. Repariuk, V. Sharkan, Identification of authorship of Ukrainian-language texts of journalistic style using neural networks, *Eastern-European Journal of Enterprise Technologies* 1 (2 (103)) (2020) 30-36. doi: 10.15587/1729-4061.2020.195041.
- [5] M. Jafari, H. Xu, Intelligent Control for Unmanned Aerial Systems with System Uncertainties and Disturbances Using Artificial Neural Network, *Drones* 3 (2018) 24-36.
- [6] M. Miyakawa, Criteria for selecting a variable in the construction of efficient decision trees, *IEEE Transactions on Computers* 38(1) (1989) 130-141.
- [7] H. Koskimaki, I. Juutilainen, P. Laurinen, J. Rönig, Two-level clustering approach to training data instance selection: a case study for the steel industry, in: *Proceedings of the International Joint Conference on Neural Networks, IJCNN 2008*, part of the IEEE World Congress on Computational Intelligence, IEEE, Los Alamitos. (2008) 3044-3049. doi: 10.1109/ijcnn.2008.4634228.
- [8] S.F. Jaman, M. Anshari, Facebook as marketing tools for organizations: Knowledge management analysis, in: *Dynamic perspectives on globalization and sustainable business in Asia*, IGI Global, Hershey. (2019) 92-105. doi: 10.4018/978-1-5225-7095-0.ch007.
- [9] W. Sirichotedumrong, T. Maekawa, Y. Kinoshita, H. Kiya. Privacy-preserving deep neural networks with pixel-based image encryption considering data augmentation in the encrypted domain, in: *Proceedings of the 26th IEEE International Conference on Image Processing, ICIP 2019*, IEEE, Los Alamitos. (2019) 674-678. doi: 10.48550/arXiv.1905.01827.
- [10] R.L. De Mántaras, A distance-based attribute selection measure for decision tree induction, *Machine learning* 6(1) (1991) 81-92.
- [11] K. Karimi, H. Hamilton, Generation and Interpretation of Temporal Decision Rules, *International Journal of Computer Information Systems and Industrial Management Applications* 3 (2011) 314-323.
- [12] B. Kamiński, M. Jakubczyk, P. Szufel, A framework for sensitivity analysis of decision trees, *Central European Journal of Operations Research* 26 (1) (2017) 135-159.
- [13] H. Deng, G. Runger, E. Tuv, Bias of importance measures for multi-valued attributes and solutions, in: *Proceedings of the 21st International Conference on Artificial Neural Networks*, volume 2 of ICANN 2011, Springer-Verlag, Berlin. (2011) 293-300. doi: 10.1007/978-3-642-21738-8_38.
- [14] S.A. Subbotin, Construction of decision trees for the case of low-information features, *Radio Electronics, Computer Science, Control* 1 (2019) 121-130.
- [15] M. Jafari, H. Xu. Intelligent Control for Unmanned Aerial Systems with System Uncertainties and Disturbances Using Artificial Neural Network. *Drones*. 3 (2018) 24–36.
- [16] A. Painsky, S. Rosset, Cross-validated variable selection in tree-based methods improves predictive performance, *IEEE Transactions on Pattern Analysis and Machine Intelligence* 39(11) (2017) 2142-2153. doi:10.1109/tpami.2016.2636831.
- [17] D. Imamovic, E. Babovic, N. Bijedic, Prediction of mortality in patients with cardiovascular disease using data mining methods, in: *Proceedings of the 19th International Symposium INFOTEH-JAHORINA, INFOTEH 2020*, IEEE, Los Alamitos. (2020) 1-4.
- [18] S.B. Kotsiantis, *Supervised Machine Learning: A Review of Classification Techniques*,

Informatica 31 (2007) 249-268.

[19] Y.I. Zhuravlev, V.V. Nikiforov, Recognition algorithms based on the calculation of estimates, *Cybernetics* 3 (1971) 1-11.

[20] Y.A. Vasilenko, E.Y. Vasilenko, I.F. Povkhan, Branched feature selection method in mathematical modeling of multi-level image recognition systems, *Artificial Intelligence* 7 (2003) 246-249.

[21] I. Povkhan, A constrained method of constructing the logic classification trees on the basis of elementary attribute selection, *CEUR Workshop Proceedings, Proceedings of the Second International Workshop on Computer Modeling and Intelligent Systems (CMIS-2020)*, 2608 (2020) 843-857.

[22] Y.A. Vasilenko, E.Y. Vasilenko, I.F. Povkhan, Conceptual basis of image recognition systems based on the branched feature selection method, *European Journal of Enterprise Technologies* 7(1) (2004) 13-15.

[23] I. Povkhan, M. Lupei, The algorithmic classification trees, in: *Proceedings of the IEEE Third International Conference on Data Stream Mining & Processing, DSMP 2020*, IEEE, Los Alamitos, (2020) 37-44.

[24] I. Povkhan, M. Lupei, M. Kliap, V. Laver, The issue of efficient generation of generalized features in algorithmic classification tree methods, , in: *Proceedings of the International Conference on Data Stream Mining and Processing, DSMP 2020*, Springer, Cham. (2020) 98-113.

[25] I. Povkhan, Classification models of flood-related events based on algorithmic trees, *Eastern-European Journal of Enterprise Technologies* 6(4) (2020) 58-68. doi: <https://doi.org/10.15587/1729-4061.2020.219525>.

[26] J. Rabcan, V. Levashenko, E. Zaitseva, M. Kvassay, S. Subbotin, Application of Fuzzy Decision Tree for Signal Classification, *IEEE Transactions on Industrial Informatics* 15(10) (2019) 5425-5434. doi: 10.1109/tii.2019.2904845.

[27] P.E. Utgoff, Incremental induction of decision trees, *Machine learning* 4(2) (1989) 161-186. doi:10.1023/A:1022699900025.

[28] L. Hyafil, R.L. Rivest, Constructing optimal binary decision trees is NP-complete, *Information Processing Letters* 5(1) (1976) 15-17.

[29] H. Wang, M. Hong, Online ad effectiveness evaluation with a two-stage method using a Gaussian filter and decision tree approach, *Electronic Commerce Research and Applications* 1(35) (2019) Article 100852. doi: 10.1016/j.elerap.2019.100852.

[30] I.L. Kaftannikov, A.V. Parasich, Decision Tree's Features of Application in Classification Problems, *Bulletin of the South Ural State University, Ser. Computer Technologies, Automatic Control, Radio Electronics* 15(3) (2015) 26-32. doi: 10.14529/ctcr150304.

[31] I.F. Povhan, Logical recognition tree construction on the basis a step-to-step elementary attribute selection, *Radio Electronics, Computer Science, Control* 2 (2020) 95-106.

[32] S.L. Lin, H.W. Huang, Improving Deep Learning for Forecasting Accuracy in Financial Data, *Discrete Dynamics in Nature and Society* Volume 2020 (2020) Article ID 5803407. doi: 10.1155/2020/5803407.

[33] R. Srikant, R. Agrawal, Mining generalized association rules, *Future Generation Computer Systems*, 13(2) (1997) 161-180.

[34] Ameisen E, *Building Machine Learning Powered Applications: Going from Idea to Product*, O'Reilly Media, California. (2020).

[35] M. Sewak, *Deep Reinforcement Learning*, *Frontiers of Artificial Intelligence*, Springer, Berlin, 2020. doi: 10.1007/978-981-13-8285-7.

[36] J.B. Harley, D. Sparkman, Machine learning and NDE: Past, present, and future, *AIP Conference Proceedings*, 2102(1) (2019). doi:10.1063/1.5099819.

Neural Network Forecasting Method for Inventory Management in the Supply Chain

Oleg Grygor¹, Eugene Fedorov¹, Olga Nechyporenko¹ and Mykola Grygorian²

¹*Cherkasy State Technological University, Shevchenko Blvd., 460, Cherkasy, 18006, Ukraine*

²*Cherkasy Institute of Fire Safety named after Chornobyl Heroes of National University of Civil Defence of Ukraine, Onoprienko str., 8, Cherkasy, 18034, Ukraine*

Abstract

Determining the optimal level of inventory comes down to the timeliness of the procurement and replenishment procedures, which ensure the minimum total costs associated with procurement and storage. The problem of insufficient prediction accuracy for inventory management arising in supply chains is considered. A neural network prediction model based on a Time-Delay Restricted Boltzmann Machine with unit delays cascades in the output and input layers is proposed. During the structural identification of this model, the neurons count in the hidden layer was calculated, and the parametric identification was performed based on the CUDA parallel processing technology. This improves the prediction efficiency by increasing the prediction accuracy and decreasing the computational complexity. Software has been developed by the Matlab package that realizes the offered method. The created software is used to implement the prediction in the supply chain management problem.

Keywords

prediction accuracy, supply chain management problem, neural network prediction model, Restricted Boltzmann Machine, stochastic learning

1. Introduction

Despite a large number of studies devoted to the problem of improving the efficiency of supply chains and reducing logistics costs associated with inventory management, some questions remain open. Currently, the problem in the supply chain management (SCM) is not enough high prediction accuracy [1-4]. Thus SCM can be effectless, for example for the concept “Material Requirements Planning” (MRP). Therefore, the development of supply chain prediction methods is an urgent problem.

As of today, many approaches are known as prediction tools, among which are:

- methods based on a modified liquid state machine and a modified gated recurrent unit [5, 6];
- autoregressive and regressive prediction methods [7];
- structural prediction methods Markov chains [8];
- prediction methods exponential smoothing [9];
- logical prediction methods regression and classification trees [10];
- artificial neural network prediction methods [11, 12].

The use of artificial neural networks in prediction provides tangible advantages:

- the relations between features are investigated on ready-made models;
- no guess-work about the features distribution are required;

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: chdtu-cherkassy@ukr.net (O. Grygor); fedorovee75@ukr.net (E. Fedorov); olne@ukr.net (O. Nechyporenko); hryhorian_mykola@chipb.org.in (M. Grygorian)
ORCID: 0000-0002-5233-290X (O. Grygor); 0000-0003-3841-7373 (E. Fedorov); 0000-0002-3954-3796 (O. Nechyporenko); 0000-0003-0359-5735 (M. Grygorian)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

- initial data about the features may be missing;
- source data can be noisy, incomplete or highly correlated;
- systems analysis with a big nonlinearity degree is possible;
- rapid model create;
- high adaptability;
- systems analysis with a big number of features is possible;
- full list of all possible models is not demanded;
- systems analysis with heterogeneous features is possible.

Given the advantages above, the paper will use a artificial neural network prediction method.

The goal of the work is to create a artificial neural network prediction method in the supply chain to improve prediction accuracy. To achieve the aim, the next tasks were solved:

- analyze existing prediction methods;
- propose a artificial neural network prediction model;
- choose a criterion for evaluating the artificial neural network prediction model quality;
- propose a method for calculating parameters of the artificial neural network prediction model;
- conduct a numerical study.

2. Formal problem statement

Let the learning set $S = \{(x_\mu, d_\mu)\}$, $\mu \in \overline{1, P}$ be given for the prediction.

The problem of enhancement the prediction accuracy for the Time-Delay Restricted Boltzmann Machine (TDRBM) model $g(x, W)$, where x – input vector, W – parameters vector, is showed as the problem of searching for this model such a parameters vector W^* that meets the criterion

$$F = \frac{1}{P} \sum_{\mu=1}^P (g(x_\mu, W^*) - d_\mu)^2 \rightarrow \min.$$

3. Literature review

The often used prediction artificial neural networks are:

- Jordon's neural network (JNN) [13, 14]. It is a recurrent two-layer network and is founded on multilayer perceptron (MLP). JNN is recommended for prediction stationary short-term signals;
- Elman's neural network (ENN) [15, 16]. It is a recurrent two-layer network and is founded on MLP. SRN is recommended for prediction stationary short-term signals;
- recurrent multilayer perceptron (RMLP) [17, 18]. It is a recurrent multilayer network and is founded on MLP. RMLP is recommended for prediction stationary short-term signals;
- nonlinear autoregressive model (NAR) [19, 20]. It is a non-recurrent two-layer network and is founded on MLP. NAR is recommended for prediction stationary short-term signals;
- nonlinear autoregressive moving average model (NARMA) [21, 22]. It is a recurrent two-layer network and is founded of MLP. NARMA is recommended for prediction stationary short-term signals;
- bidirectional recurrent neural network (BRNN) [23, 24]. It is a recurrent two-layer network and is founded on two Elman networks. BRNN is recommended for prediction non-stationary long-term signals.

Table 1 shows the comparative features of prediction artificial neural networks.

As follows from Table 1, most neural networks have one or more of the following disadvantages:

- have a high probability of hitting a local extremum;
- do not have the ability to train in batch mode;
- have no feedback.

Due to this, it is relevant to create a neural network that will eliminate the indicated disadvantages.

Table 1

Comparative features of prediction artificial neural networks

Criterion \ Network	RMLP	JNN	ENN (SRN)	NAR	NARMA	BRNN
The presence of feedback	+	+	-	-	+	+
Unit delay cascade	-	-	-	+	+	-
Low probability of local extremum	-	-	-	-	-	-
High training speed	+	+	+	+	+	+
Possibility of batch learning	-	-	-	+	-	-

4. Block diagram of the neural network prediction model

Figures 1 and 2 show the block diagrams of a prediction model based on a Time-Delay Restricted Boltzmann Machine (TDRBM) of two types, which are stochastic recurrent neural networks with one hidden layer.

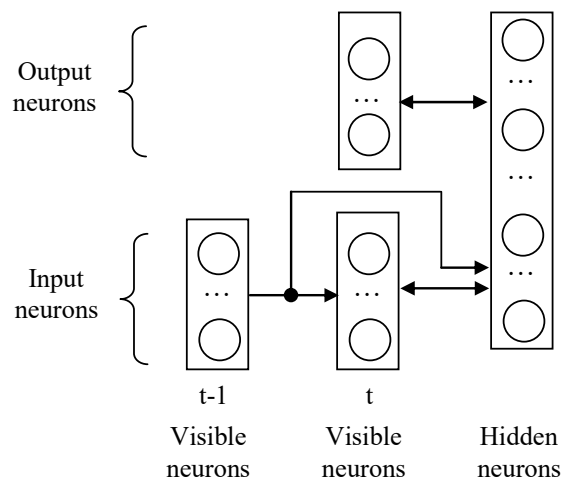


Figure 1: Block diagram of a prediction model TDRBM for visible input neurons (TDRBM type first)

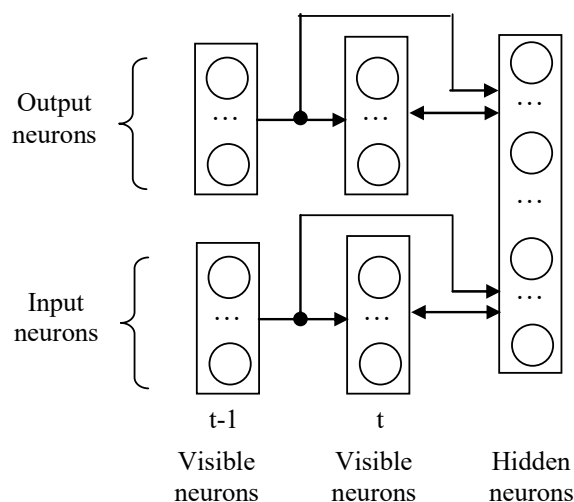


Figure 2: Block diagram of a prediction model TDRBM for the visible input neurons and visible output neurons (TDRBM type second)

Unlike the traditional Restricted Boltzmann Machine (RBM) [25, 26], unit delay cascades are used for neurons in the visible layer. TDRBM type first has unit delay cascade in the visible input layer. TDRBM type second has unit delay cascade in the visible input layers and visible output layers.

5. Neural network prediction models

5.1. Components of the prediction model TDRBM

The components of the TDRBM model are stochastic neurons, the state of which is described based on the Bernoulli distribution in the form

$$x_j = \begin{cases} 1, & \text{with probability } P_j \\ 0, & \text{with probability } 1 - P_j \end{cases}.$$

The probability of transition of the j^{th} stochastic neuron to state 1 is defined as

$$P_j = \frac{1}{1 + \exp(-\Delta E_j)},$$

where ΔE_j – is the TDRBM energy increment when the state of the j^{th} stochastic neuron changes from 0 to 1.

5.2. Prediction model TDRBM type first

Positive phase (steps 1-4)

1. Initial value assignment time moment

$$\mu = 1.$$

Initial value assignment of the time delay input neurons state

$$\mathbf{x}^{in}(\mu - t) = 0, \quad t \in \overline{1, M^{in}}.$$

2. Initial value assignment of the input and output neurons state

$$\mathbf{x}^{in}(\mu) = \mathbf{x}_{\mu}^{in}, \quad \mathbf{x}^{out}(\mu) = \mathbf{0}.$$

3. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{ij}^{in-h} x_i^{in}(\mu - t) - \sum_{i=1}^{N^{out}} w_{ij}^{out-h} x_i^{out}(\mu)\right)}, \quad j \in \overline{1, N^h}$$

where w_{ij}^{in-h} – weights between the input layer and the hidden layer,

w_{ij}^{out-h} – weights between the output layer and the hidden layer,

b_j^h – bias of the hidden layer neurons,

M^{in} – the unit delays count for the input layer,

N^{in} – the neurons count in the input layer,

N^h – the neurons count in the hidden layer,

N^{out} – the neurons count in the output layer.

4. Calculation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, \quad j \in \overline{1, N^h}$$

where $U(0,1)$ – standard uniform distribution on a line $[0,1]$.

Negative phase (step 5)

5. Calculation of the output neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{out} - \sum_{i=1}^{N^h} w_{ij}^{out-h} x_i^h(\mu)\right)}, \quad j \in \overline{1, N^{out}}$$

where b_j^{out} – bias of the output layer neurons.

6. Calculation of the output neurons state

$$x_j^{out}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, \quad j \in \overline{1, N^{out}}.$$

The outcome is $(x_1^{out}(\mu), \dots, x_{N^{out}}^{out}(\mu))$.

5.3. Prediction model based on TDRBM type second

Positive phase (steps 1-4)

1. Initial value assignment time moment

$$\mu = 1.$$

Initial value assignment time moment of the time delay input and output neurons state

$$\mathbf{x}^{in}(\mu - t) = 0, \quad t \in \overline{1, M^{in}},$$

$$\mathbf{x}^{out}(\mu - t) = 0, \quad t \in \overline{1, M^{out}}.$$

2. Initial value assignment of the input and output neurons state

$$\mathbf{x}^{in}(\mu) = \mathbf{x}_\mu^{in}, \quad \mathbf{x}^{out}(\mu) = \mathbf{0}.$$

3. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-h} x_i^{in}(\mu - t) - \sum_{t=0}^{M^{out}} \sum_{i=1}^{N^{out}} w_{tij}^{out-h} x_i^{out}(\mu - t)\right)}, \quad j \in \overline{1, N^h},$$

where w_{tij}^{in-h} – weights between the input layer and the hidden layer,

w_{tij}^{out-h} – weights between the output layer and the hidden layer,

b_j^h – bias of the hidden layer neurons,

M^{in} – the unit delays count for the input layer,

M^{out} – the unit delays count for the output layer,

N^{in} – the neurons count in the input layer,

N^h – the neurons count in the hidden layer,

N^{out} – the neurons count in the output layer.

4. Computation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, \quad j \in \overline{1, N^h}$$

where $U(0,1)$ – standard uniform distribution on a line $[0,1]$.

Negative phase (step 5)

5. Calculation of the output neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{out} - \sum_{t=1}^{M^{out}} \sum_{i=1}^{N^{out}} w_{tij}^{out-out} x_i^{out}(\mu - t) - \sum_{i=1}^{N^h} w_{0ij}^{out-h} x_i^h(\mu)\right)}, \quad j \in \overline{1, N^{out}}$$

6. Calculation of the output neurons state

$$x_j^{out}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^{out}}.$$

The outcome is $(x_1^{out}(\mu), \dots, x_{N^{out}}^{out}(\mu))$.

6. Criterion for the artificial neural network prediction model quality

For learning the TDRBM, a model adequacy criterion was selection, which means the selection of such parameters $W = \{w_{tij}^{in-h}, w_{tij}^{out-h}, w_{tij}^{in-in}, w_{tij}^{out-out}\}$, which get a mean squared error (MSE) minimum:

$$F = \frac{1}{P} \sum_{\mu=1}^P \sum_{i=1}^{N^{out}} (x_{\mu i}^{out} - d_{\mu i})^2 \rightarrow \min_W. \quad (1)$$

The learning of the TDRBM model is taking into account the criterion (1).

7. Methods for calculating the neural network prediction model parameters

7.1. Principles of calculating the neural network prediction model parameters

The determination of the TDRBM parameters is perform on the found of the CD-1 method, which speeds up stochastic learning with a teacher, since instead of stabilizing the state of neurons, it performs only one step of adjusting their state. TDRBM operates in two phases: positive and negative.

For *TDRBM type first*, in the positive phase, visible input neurons are fixed at times from $t - M^{in}$ to t and visible output neurons at time t , and the TDRBM performs one step of tuning hidden neurons. In the negative phase, the visible input neurons are first fixed at times from $t - M^{in}$ to $t - 1$, the hidden neurons are trained in the positive phase, and the TDRBM performs one step of tuning the visible input and output neurons at time t . After that, the visible input neurons are fixed at times from $t - 1$ to $t - M^{in}$, and the visible input and output neurons trained in the negative phase at time t , and the TDRBM performs one step of tuning the hidden neurons.

For *TDRBM type second*, in the positive phase visible input neurons are fixed at times $t - M^{in}$ to t , visible output neurons at times from $t - M^{out}$ to t , and the TDRBM performs one step of tuning hidden neurons. In the negative phase, the visible input neurons are first fixed at times from $t - M^{in}$ to $t - 1$, the visible output neurons at times from $t - M^{out}$ to $t - 1$, the hidden neurons are trained in the positive phase, and the TDRBM performs one step of tuning the visible input and output neurons at time t . After that, visible input neurons are recorded at times from $t - 1$ to $t - M^{in}$, visible output neurons at times from $t - 1$ to $t - M^{out}$, and visible input and output neurons trained in the negative phase at time t , and TDRBM performs one step of tuning hidden neurons.

7.2. Method for calculating the prediction model parameters found on TDRBM type first

1. Learning iteration number $n = 1$, initial value assignment by uniform distribution means on the interval $(0,1)$ or $[-0.5, 0.5]$ bias $b_i^{in}(n)$, $i \in \overline{1, N^{in}}$, $b_i^{out}(n)$, $i \in \overline{1, N^{out}}$, $b_j^h(n)$, $j \in \overline{1, N^h}$, and weights $w_{tij}^{in-h}(n)$, $t \in \overline{0, M^{in}}$, $i \in \overline{1, N^{in}}$, $j \in \overline{1, N^h}$, $w_{ij}^{out-h}(n)$, $i \in \overline{1, N^{out}}$, $j \in \overline{1, N^h}$, $w_{tij}^{in-in}(n)$,

$t \in \overline{1, M^{in}}$, $i, j \in \overline{1, N^{in}}$, $w_{tii}^{in-h}(n) = 0$, $w_{ii}^{out-h}(n) = 0$, $w_{tii}^{in-in}(n) = 0$, $w_{0ij}^{in-h}(n) = w_{0ji}^{in-h}(n)$, $w_{ij}^{out-h}(n) = w_{ji}^{out-h}(n)$, where M^{in} is the unit delays count for input neurons.

2. A learning set $\{(\mathbf{x}_\mu^{in}, \mathbf{x}_\mu^{out}) | \mathbf{x}_\mu^{in} \in \{0,1\}^{N^{in}}, \mathbf{x}_\mu^{out} \in \{0,1\}^{N^{out}}\}$, $\mu \in \overline{1, P}$ is defined, where $\mathbf{x}_\mu^{in}$ – μ^{th} learning input vector, $\mathbf{x}_\mu^{out}$ – μ^{th} learning output vector, P is the learning set capacity.

Positive phase (steps 3-7)

3. Initial value assignment time moment

$$\mu = 1.$$

Initial value assignment of the time delay input neurons state

$$\mathbf{x}^{in}(\mu - t) = 0, \mathbf{x}^{out}(\mu - t) = 0, t \in \overline{1, M^{in}}.$$

4. Initial value assignment of the input and output neurons state

$$\mathbf{x}^{in}(\mu) = \mathbf{x}_\mu^{in}, \mathbf{x}^{out}(\mu) = \mathbf{x}_\mu^{out}.$$

5. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h(n) - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-h}(n) x_i^{in}(\mu - t) - \sum_{i=1}^{N^{out}} w_{ij}^{out-h}(n) x_i^{out}(\mu)\right)}, j \in \overline{1, N^h}$$

6. Calculation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^h}$$

7. Saving the neurons state in the positive phase, i.e. $\mathbf{x}^{in}(\mu) = \mathbf{x}^{in}(\mu)$, $\mathbf{x}^{out}(\mu) = \mathbf{x}^{out}(\mu)$, $\mathbf{x}^h(\mu) = \mathbf{x}^h(\mu)$. If $\mu < P$, then $\mu = \mu + 1$, go to 4.

Negative phase (steps 8-16)

8. Initial value assignment time moment

$$\mu = 1.$$

Initial value assignment of the time delay input neurons state

$$\mathbf{x}^{in}(\mu - t) = 0, \mathbf{x}^{out}(\mu - t) = 0, t \in \overline{1, M^{in}}.$$

9. Initial value assignment of the of hidden, input and output neurons state

$$\mathbf{x}^{in}(\mu) = \mathbf{x}^{in}(\mu), \mathbf{x}^{out}(\mu) = \mathbf{x}^{out}(\mu), \mathbf{x}^h(\mu) = \mathbf{x}^h(\mu).$$

10. Calculation of the output neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{out}(n) - \sum_{i=1}^{N^h} w_{ij}^{out-h}(n) x_i^h(\mu)\right)}, j \in \overline{1, N^{out}}.$$

11. Calculation of the of output neurons state

$$x_j^{out}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^{out}}.$$

12. Calculation of the input neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{in}(n) - \sum_{t=1}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-in}(n) x_i^{in}(\mu - t) - \sum_{i=1}^{N^h} w_{0ij}^{in-h}(n) x_i^h(\mu)\right)}, j \in \overline{1, N^h}.$$

13. Calculation of the input neurons state

$$x_j^{in}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^h}.$$

14. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h(n) - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{ij}^{in-h}(n) x_i^{in}(\mu-t) - \sum_{i=1}^{N^{out}} w_{ij}^{out-h}(n) x_i^{out}(\mu)\right)}, j \in \overline{1, N^h}.$$

15. Calculation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^h}.$$

16. Saving the neurons state in the negative phase, i.e. $\mathbf{x}2^{in}(\mu) = \mathbf{x}^{in}(\mu)$, $\mathbf{x}2^{out}(\mu) = \mathbf{x}^{out}(\mu)$, $\mathbf{x}2^h(\mu) = \mathbf{x}^h(\mu)$. If $\mu < P$, then $\mu = \mu + 1$, go to 9.

17. Tuning of bias input layer founded on Boltzmann's rule

$$b_i^{in}(n) = b_i^{in}(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_i^{in}(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_i^{in}(\mu) \right), i \in \overline{1, N^{in}}.$$

18. Tuning of bias output layer founded on Boltzmann's rule

$$b_i^{out}(n) = b_i^{out}(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_i^{out}(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_i^{out}(\mu) \right), i \in \overline{1, N^{out}}.$$

19. Tuning of bias hidden layer neurons founded on Boltzmann's rule

$$b_j^h(n) = b_j^h(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_j^h(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_j^h(\mu) \right), j \in \overline{1, N^h}.$$

20. Tuning of weights between the input layer and the hidden layer founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{in-h}(n) &= w_{ij}^{in-h}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{0, M^{in}}, i \in \overline{1, N^{in}}, j \in \overline{1, N^h}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{in}(\mu-t), x1_j^h(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{in}(\mu-t), x2_j^h(\mu)), \\ \varphi(x_i, x_j) &= \begin{cases} 1, & x_i = x_j \\ 0, & x_i \neq x_j \end{cases}. \end{aligned}$$

21. Tuning of weights between the output layer and the hidden layer founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{out-h}(n) &= w_{ij}^{out-h}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), i \in \overline{1, N^{out}}, j \in \overline{1, N^h}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{out}(\mu), x1_j^h(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{out}(\mu), x2_j^h(\mu)). \end{aligned}$$

22. Tuning of weights between input layers founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{in-in}(n) &= w_{ij}^{in-in}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{1, M^{in}}, i, j \in \overline{1, N^{in}}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{in}(\mu-t), x1_j^{in}(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{in}(\mu-t), x2_j^{in}(\mu)), \end{aligned}$$

23. If the neurons states in the output layer of the positive and negative phases differ slightly, i.e.

$$\frac{1}{P} \sum_{\mu=1}^P \left(\sum_{i=1}^{N^{out}} |x1_i^{out}(\mu) - x2_i^{out}(\mu)| \right) > \varepsilon, \text{ then } n = n + 1, \text{ go to 2.}$$

The method for calculating the prediction model parameter based on TDRBM type first is shown in Figure 3.

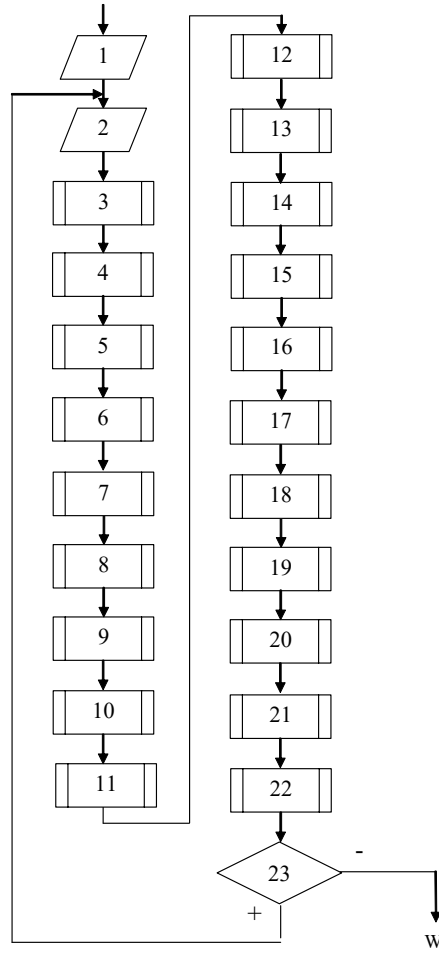


Figure 3: Flowchart of the method for calculating the prediction model parameters found on TDRBM type first

7.3. Method for calculating the prediction model parameters found on TDRBM type second

1. Learning iteration number $n = 1$, initial value assignment by uniform distribution means on the interval $(0,1)$ or $[-0.5, 0.5]$ bias $b_i^{in}(n)$, $i \in \overline{1, N^{in}}$, $b_i^{out}(n)$, $i \in \overline{1, N^{out}}$, $b_j^h(n)$, $j \in \overline{1, N^h}$, and weights $w_{tij}^{in-h}(n)$, $t \in \overline{0, M^{in}}$, $i \in \overline{1, N^{in}}$, $j \in \overline{1, N^h}$, $w_{tij}^{out-h}(n)$, $t \in \overline{0, M^{out}}$, $i \in \overline{1, N^{out}}$, $j \in \overline{1, N^h}$, $w_{tij}^{in-in}(n)$, $t \in \overline{1, M^{in}}$, $i, j \in \overline{1, N^{in}}$, $w_{tij}^{out-out}(n)$, $t \in \overline{1, M^{out}}$, $i, j \in \overline{1, N^{out}}$, $w_{iii}^{in-h}(n) = 0$, $w_{ii}^{out-h}(n) = 0$, $w_{iii}^{in-in}(n) = 0$, $w_{0ij}^{in-h}(n) = w_{0ji}^{in-h}(n)$, $w_{ij}^{out-h}(n) = w_{ji}^{out-h}(n)$, where M^{in} is the unit delays count for input neurons, M^{out} is the unit delays count for output neurons.
2. A learning set $\{(\mathbf{x}_\mu^{in}, \mathbf{x}_\mu^{out}) | \mathbf{x}_\mu^{in} \in \{0,1\}^{N^{in}}, \mathbf{x}_\mu^{out} \in \{0,1\}^{N^{out}}\}$ is defined, $\mu \in \overline{1, P}$, where $\mathbf{x}_\mu^{in} - \mu^{\text{th}}$ learning input vector, $\mathbf{x}_\mu^{out} - \mu^{\text{th}}$ learning output vector, P is the learning set capacity.

Positive phase (steps 3-7)

3. Initial value assignment time moment

$$\mu = 1.$$

Initial value assignment of the time delay input and output neurons state

$$\mathbf{x}^{in}(\mu - t) = 0, \mathbf{x}^{out}(\mu - t) = 0, t \in \overline{1, M^{in}},$$

$$\mathbf{x}^{out}(\mu-t)=0, \mathbf{x}1^{out}(\mu-t)=0, t \in \overline{1, M^{out}}.$$

4. Initial value assignment of the input and output neurons state

$$\mathbf{x}^{in}(\mu)=\mathbf{x}_{\mu}^{in}, \mathbf{x}^{out}(\mu)=\mathbf{x}_{\mu}^{out}.$$

5. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h(n) - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-h}(n) x_i^{in}(\mu-t) - \sum_{t=0}^{M^{out}} \sum_{i=1}^{N^{out}} w_{tij}^{out-h}(n) x_i^{out}(\mu-t)\right)}, j \in \overline{1, N^h}.$$

6. Calculation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^h}.$$

7. Saving the neurons state in the positive phase, i.e. $\mathbf{x}1^{in}(\mu)=\mathbf{x}^{in}(\mu)$, $\mathbf{x}1^{out}(\mu)=\mathbf{x}^{out}(\mu)$, $\mathbf{x}1^h(\mu)=\mathbf{x}^h(\mu)$. If $\mu < P$, then $\mu = \mu + 1$, go to 4.

Negative phase (steps 8-16)

8. Initial value assignment time moment

$$\mu=1.$$

Initial value assignment of the time delay input and output neurons state

$$\mathbf{x}^{in}(\mu-t)=0, \mathbf{x}2^{in}(\mu-t)=0, t \in \overline{1, M^{in}},$$

$$\mathbf{x}^{out}(\mu-t)=0, \mathbf{x}2^{out}(\mu-t)=0, t \in \overline{1, M^{out}}.$$

9. Initial value assignment of the of hidden, input and output neurons state

$$\mathbf{x}^{in}(\mu)=\mathbf{x}1^{in}(\mu), \mathbf{x}^{out}(\mu)=\mathbf{x}1^{out}(\mu), \mathbf{x}^h(\mu)=\mathbf{x}1^h(\mu).$$

10. Calculation of the output neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{out}(n) - \sum_{t=1}^{M^{out}} \sum_{i=1}^{N^{out}} w_{tij}^{out-out}(n) x_i^{out}(\mu-t) - \sum_{i=1}^{N^h} w_{0ij}^{out-h}(n) x_i^h(\mu)\right)}, j \in \overline{1, N^{out}}.$$

11. Calculation of the of output neurons state

$$x_j^{out}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^{out}}.$$

12. Calculation of the input neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^{in}(n) - \sum_{t=1}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-in}(n) x_i^{in}(\mu-t) - \sum_{i=1}^{N^h} w_{0ij}^{in-h}(n) x_i^h(\mu)\right)}, j \in \overline{1, N^{in}}.$$

13. Calculation of the input neurons state

$$x_j^{in}(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^{in}}.$$

14. Calculation of the hidden neurons transition probability to state 1

$$P_j = \frac{1}{1 + \exp\left(-b_j^h(n) - \sum_{t=0}^{M^{in}} \sum_{i=1}^{N^{in}} w_{tij}^{in-h}(n) x_i^{in}(\mu-t) - \sum_{t=0}^{M^{out}} \sum_{i=1}^{N^{out}} w_{tij}^{out-h}(n) x_i^{out}(\mu-t)\right)}, j \in \overline{1, N^h}.$$

15. Calculation of the hidden neurons state

$$x_j^h(\mu) = \begin{cases} 1, & P_j \geq U(0,1) \\ 0, & P_j < U(0,1) \end{cases}, j \in \overline{1, N^h}.$$

16. Saving the neurons state in the negative phase, i.e. $\mathbf{x}2^{in}(\mu) = \mathbf{x}^{in}(\mu)$, $\mathbf{x}2^{out}(\mu) = \mathbf{x}^{out}(\mu)$, $\mathbf{x}2^h(\mu) = \mathbf{x}^h(\mu)$. If $\mu < P$, then $\mu = \mu + 1$, go to 9.

17. Tuning of bias input layer founded on Boltzmann's rule

$$b_i^{in}(n) = b_i^{in}(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_i^{in}(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_i^{in}(\mu) \right), i \in \overline{1, N^{in}}.$$

18. Tuning of bias output layer founded on Boltzmann's rule

$$b_i^{out}(n) = b_i^{out}(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_i^{out}(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_i^{out}(\mu) \right), i \in \overline{1, N^{out}}.$$

19. Tuning of bias hidden layer neurons founded on Boltzmann's rule

$$b_j^h(n) = b_j^h(n) + \eta \left(\frac{1}{P} \sum_{\mu=1}^P x1_j^h(\mu) - \frac{1}{P} \sum_{\mu=1}^P x2_j^h(\mu) \right), j \in \overline{1, N^h}.$$

20. Tuning of weights between the input layer and the hidden layer founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{in-h}(n) &= w_{ij}^{in-h}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{0, M^{in}}, i \in \overline{1, N^{in}}, j \in \overline{1, N^h}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{in}(\mu - t), x1_j^h(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{in}(\mu - t), x2_j^h(\mu)), \\ \varphi(x_i, x_j) &= \begin{cases} 1, & x_i = x_j \\ 0, & x_i \neq x_j \end{cases}. \end{aligned}$$

21. Tuning of weights between the output layer and the hidden layer founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{out-h}(n) &= w_{ij}^{out-h}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{0, M^{out}}, i \in \overline{1, N^{out}}, j \in \overline{1, N^h}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{out}(\mu), x1_j^h(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{out}(\mu), x2_j^h(\mu)). \end{aligned}$$

22. Tuning of weights between input layers founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{in-in}(n) &= w_{ij}^{in-in}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{1, M^{in}}, i, j \in \overline{1, N^{in}}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{in}(\mu - t), x1_j^{in}(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{in}(\mu - t), x2_j^{in}(\mu)). \end{aligned}$$

23. Tuning of weights between output layers founded on Boltzmann's rule

$$\begin{aligned} w_{ij}^{out-out}(n) &= w_{ij}^{out-out}(n) + \eta(\rho_{ij}^+ - \rho_{ij}^-), t \in \overline{1, M^{out}}, i, j \in \overline{1, N^{out}}, \\ \rho_{ij}^+ &= \frac{1}{P} \sum_{\mu=1}^P \varphi(x1_i^{out}(\mu - t), x1_j^{out}(\mu)), \rho_{ij}^- = \frac{1}{P} \sum_{\mu=1}^P \varphi(x2_i^{out}(\mu - t), x2_j^{out}(\mu)). \end{aligned}$$

24. If the neurons states in the output layer of the positive and negative phases differ slightly, i.e.

$$\frac{1}{P} \sum_{\mu=1}^P \left(\sum_{i=1}^{N^{out}} |x1_i^{out}(\mu) - x2_i^{out}(\mu)| \right) > \varepsilon, \text{ then } n = n + 1, \text{ go to 2.}$$

The method for calculating the prediction model parameters found on TDRBM type second is shown in Figure 4.

8. Experiments and results

A numerical study of the proposed method for determining the parameter values was carried out using the CUDA technology of parallel processing of information in the Matlab package, wherein the number of threads in the block was 1024, the summation was performed based on the parallel reduction algorithm.

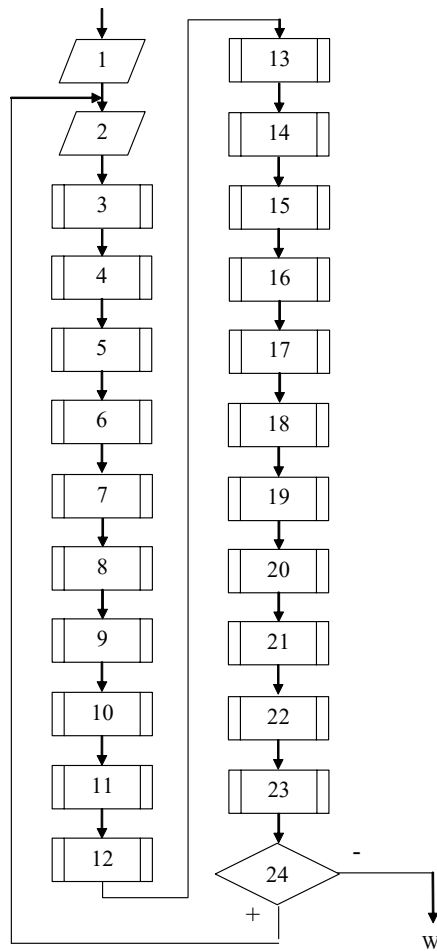


Figure 4: Flowchart of the method for calculating the prediction model parameters found on TDRBM type second

To determine the prediction model structure TDRBM with 16 input neurons (the bits count of the integer predicted metric), i.e. defining the number of hidden neurons, a experiments count were performed, the results of which are showed in Figure 5.

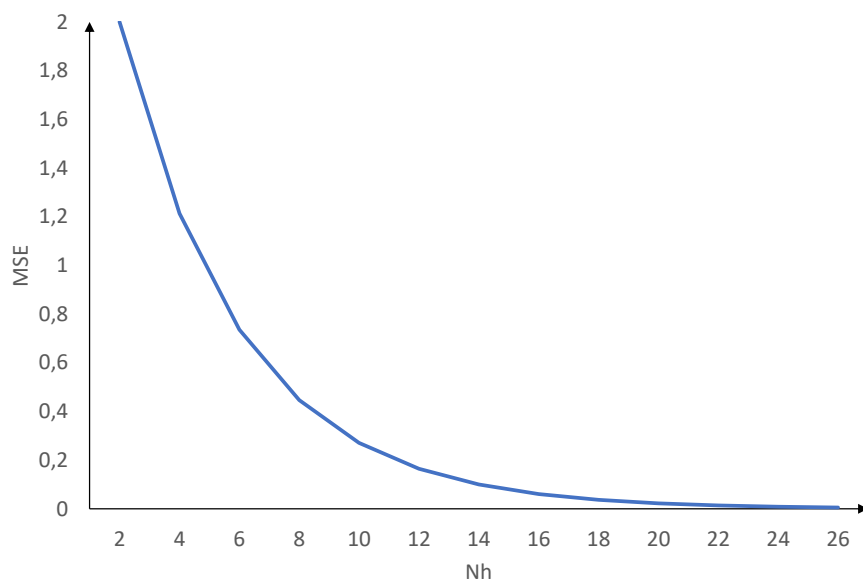


Figure 5: The mean squared error dependence on the hidden neurons number

As source data to calculate the artificial neural network prediction model parameters, a values sample of the logistics company "Ekol Ukraine" economic activities was used. The dataset capacity for the "goods sold" indicator was 1800 (in five years). The dataset was spilt into three parts - valid data (20%), training data (60%), test data (20%). The learning was consistent and occurred over 1000 epochs.

The criterion for selection the artificial neural network model structure of the was the prediction minimum mean squared error (MSE) (Table 2). According to Figure 5, with an increase in the hidden neurons count the error value decreases. For the prediction, it is enough to choose 16 hidden neurons, because with a further increase in their count the change in the MSE value is small.

Table 2

Comparative features of prediction artificial neural networks

Network Criterion	RMLP	JNN	ENN (SRN)	NAR	NARMA	BRNN	TDRBM type first / second
Minimum MSE of the prediction	0.17	0.12	0.10	0.15	0.07	0.05	0.04 / 0.02

According to Table 2, recurrent neural networks with deterministic neurons (JNN, ENN(SRN), NARMA, BRNN) are more efficient than non-recurrent neural networks with deterministic neurons (RMLP, NAR); neural networks with stochastic neurons (TDRBM) are more efficient than neural networks with deterministic neurons (RMLP, JNN, ENN(SRN), NAR, NARMA, BRNN), and TDRBM type second has the highest prediction accuracy.

9. Conclusion

1. To solve the problem of enhancement the prediction accuracy in the supply chain, the modern prediction methods were researched. These researches have shown that today the use of neural networks is the most effective.
2. To improve the efficiency of the prediction, a artificial neural network RBM was modified by the unit delays cascades in the input and output layers. In the experiments, its model structure was determined.
3. A method was offered for calculating the parameters of the offered artificial neural network prediction model based on the CUDA technology of parallel processing of information. This made it possible to ensure accuracy of the prediction and high speed.
4. The offered approach can be used in different intelligent prediction systems. For example, in supply chain inventory management systems such as MRP and Lean Production, where prediction plays an important role.

10. References

- [1] J. F. Cox, J.G. Schleher, Theory of Constraints Handbook, New York, NY, McGraw-Hill, 2010.
- [2] S. Smerichevska et al, Cluster Policy of Innovative Development of the National Economy: Integration and Infrastructure Aspects: monograph, S. Smerichevska (Eds.), Wydawnictwo naukowe WSPIA, 2020.
- [3] E. M. Goldratt, My saga to improve production, Selected Readings in Constraints Management, Falls Church, VA: APICS (1996) 43-48.
- [4] E. M. Goldratt, Production: The TOC Way (Revised Edition) including CD-ROM Simulator and Workbook, Revised edition, Great Barrington, MA: North River Press, 2003.
- [5] T. Neskorođieva, E. Fedorov, I. Izonin, Forecast method for audit data analysis by modified liquid state machine, in: CEUR Workshop Proceedings, 2020, volume 2631, pp. 145-158.

- [6] E. Fedorov, T. Utkina, O. Nechyporenko, Forecast method for natural language constructions based on a modified gated recursive block, in: *CEUR Workshop Proceedings*, vol. 2604, 2020, pp. 199-214.
- [7] R. T. Baillie, G. Kapetanios, F. Papailias, Modified information criteria and selection of long memory time series models, in: *Computational Statistics and Data Analysis*, volume 76, 2014, pp. 116–131. doi: 10.1016/j.csda.2013.04.012.
- [8] A. Wilinski, Time series modelling and forecasting based on a Markov chain with changing transition matrices, in: *Expert Systems with Applications*, volume 133, 2019, pp. 163–172. doi: 10.1016/j.eswa.2019.04.067.
- [9] P. Bidyuk, T. Prosyankina-Zharova, O. Terentiev, Modelling nonlinear nonstationary processes in macroeconomy and finances, in: *Advances in Computer Science for Engineering and Education*, volume 754, Springer, 2019, pp. 735–745. doi: 10.1007/978-3-319-91008-6_72.
- [10] B. Choubin, G. Zehtabian, A. Azareh, E. Rafiei-Sardooi, F. Sajedi-Hosseini, Ö. Kişi, Precipitation forecasting using classification and regression trees (CART) model: a comparative study of different approaches, *Environ Earth Sciences* 77, 314 (2018). doi: 10.1007/s12665-018-7498-z.
- [11] L. Lyubchik, E. Bodyansky, A. Rivtis, Adaptive harmonic components detection and forecasting in wave non-periodic time series using neural networks, in: *Proceedings of the ISCDMCI'2002*, Evpatoria, 2002, pp. 433-435.
- [12] S. N. Sivanandam, S. Sumathi, S. N. Deepa, *Introduction to Neural Networks using Matlab 6.0*, The McGraw-Hill Comp., Inc., New Delhi, 2006.
- [13] E. C. Carcano, P. Bartolini, M. Muselli and L. Piroddi, Jordan recurrent neural network versus ihacres in modelling daily streamflows, *Journal of Hydrology* 362 (2008) 291-307.
- [14] H. Hikawa and Y. Araga, Study on gesture recognition system using posture classifier and Jordan recurrent neural network, in: *Proceedings of the International Joint Conference on Neural Networks*, 2011, pp. 405-412. doi: 10.1109/IJCNN.2011.6033250.
- [15] Z. Zhang, Z. Tang, C. Vairappan, A novel learning method for Elman neural network using local search, in: *Neural Information Processing – Letters and Reviews*, vol. 11, 2007, pp. 181–188.
- [16] A. Wysocki and M. Ławryńczuk, Predictive control of a multivariable neutralisation process using Elman neural networks, in: *Advances in Intelligent Systems and Computing*, Springer: Heidelberg, 2015, pp. 335-344. doi: 10.1007/978-3-319-15796-234.
- [17] S. Haykin, *Neural networks and Learning Machines*, Upper Saddle River, New Jersey: Pearson Education, Inc., 2009.
- [18] K.-L. Du, K. M. S. Swamy, *Neural Networks and Statistical Learning*, Springer-Verlag, London, 2014.
- [19] J. A. Pucheta, C. M. Rodríguez Rivero, M. R. Herrera, C. A. Salas, H. D. Patiño and B. R. Kuchen, A feed-forward neural networks-based nonlinear autoregressive model for forecasting time series, *Computación y Sistemas* 14 (4) (2011) 423-435.
- [20] G. P. Zhang, Time series forecasting using a hybrid ARIMA and neural network model, *Neurocomputing* 50 (2003) 159–175. doi: 10.1016/S0925-2312(01)00702-0.
- [21] W.K. Wong, M. Xia, W. C. Chu, Adaptive neural network model for time-series forecasting, *European Journal of Operational Research* 207(2) (2010) 807-816. doi: 10.1016/j.ejor.2010.05.022.
- [22] H. Jaeger, H. Haass, Harnessing nonlinearity: Predicting chaotic systems and saving energy in wireless communication, *Science* 304 (2004) 78–80. doi: 10.1126/science.1091277.
- [23] M. Sundermeyer, T. Alkhoul, J. Wuebker, H. Ney, Translation modeling with bidirectional recurrent neural networks, in: *Proceedings of the Conference on Empirical Methods on Natural Language Processing*, 2014, pp. 14-25.
- [24] M. Berglund, T. Raiko, M. Honkala, L. Kärkkäinen, A. Vetek, J. Karhunen, Bidirectional Recurrent Neural Networks as Generative Models, in: *Proceedings of the 28th International Conference on Neural Information Processing Systems*, 2015, pp. 856–864.
- [25] G. E. Hinton, *A Practical Guide to Training Restricted Boltzmann Machines*, Technical Report UTML TR 2010–003, University of Toronto, 2010.
- [26] A. Fischer, C. Igel, *Training Restricted Boltzmann Machines: An Introduction*, *Pattern Recognition* 47 (2014) 25-39. doi: 10.1016/j.patcog.2013.05.025.

Helicopters Aircraft Engines Self-Organizing Neural Network Automatic Control System

Serhii Vladov^a, Yurii Shmelov^a and Ruslan Yakovliev^a

^a *Kremenchuk Flight College of Kharkiv National University of Internal Affairs, vul. Peremohy, 17/6, Kremenchuk, Poltavaska Oblast, Ukraine, 39605*

Abstract

The purpose of this work is to dedicate to the improvement of the automatic control system for helicopters aircraft engines by dividing the control object into actuating mechanism and gas turbine engine. This allows you to take into account the dynamics of the executive part of the system and the engine, it becomes possible to use the mismatch between parts of the structural diagram of the automatic control system, thereby increasing the reliability and stability of the system in various modes. As a hardware-software implementation of the automatic control system for helicopters aircraft engines, it is proposed to use a self-tuning neural network control system for multiply connected dynamic objects, the adaptation of which as a modified neural network controller of helicopters aircraft engines by introducing an integrator into the system structure made it possible to bring the graph of the real transient process in the engine closer to the ideal one, thereby increasing the reliability and stability of the system in various modes.

Keywords

Helicopter aircraft engine, automatic control system, neural network, transfer function, transient processes.

1. Introduction

The process of ensuring the stability of the operation parameters of helicopters aircraft turboshaft engines (TE) by maintaining the required (stable) compressor rotor speed and the dosage of fuel supply to the combustion chamber has always been a difficult task. Of particular difficulty are the launch modes and transient modes of engine operation, taking into account external factors (the impact of atmospheric conditions and aircraft flight modes). In view of this, automatic control systems (ACS) are used to adjust the engine [1, 2]. The values of engine thermogasdynamic parameters required for helicopter flight and the reliable and stable operation of the power plant over the entire range of operating conditions are ensured with appropriate engine adjustment carried out by the ACS. It establishes and maintains some relationships between engine parameters. This regulation is formed taking into account the requirements for specific fuel consumption and other thermogasdynamic parameters, strength limitations, the required accuracy of maintaining parameters, and other factors [3, 4].

At present, the ACS of TE that implement the specified control laws are subject to rather stringent requirements both for permissible deviations of parameters in steady-state operating modes and for dynamic errors during transients [5]. The ACS for helicopters TE is no exception. As a rule, the following requirements are imposed of TE ACS: high accuracy of maintaining the specified parameters; minimal complexity of technical execution; the possibility of switching from one mode to

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL: ser26101968@gmail.com (S. Vladov); nviddil.klk@gmail.com (Yu. Shmelov); ateu.nv.klk@gmail.com (R. Yakovliev)
ORCID: 0000-0001-8009-5254 (S. Vladov); 0000-0002-3942-2003 (Yu. Shmelov); 0000-0002-3788-2583 (R. Yakovliev)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

another (when performing a maneuver) without reducing the quality of control. To fulfill all the above requirements, it is necessary to create a new approach to the choice of the ACS structure, the synthesis of control methods and their technical implementation. This statement is based on the analysis of the results of full-scale tests and previous theoretical studies.

2. Literature review

Modern methods and experience in development automatic control systems for gas turbine engines originate in the works of scientists, for example, A. Shevyakov, S. Sirotin, O. Gurevich, F. Golberg, O. Selivanov, G. Dobryansky, V. Dedesh, V. Rutkovsky, S. Zemlyakov, B. Petrov, B. Cherkasov, V. Avgustinovich, Yu. Gusev, F. Shaimardanov, B. Ilyasov, V. Vasiliev, G. Kulikov, Yu. Kabalnov, V. Krymsky, V. Efanov and others. Problems of scientists from foreign universities, research organizations and firms involved in the creation of dynamic objects, engines and airborne equipment.

In many practical cases, it becomes necessary to automate processes occurring in complex dynamic systems, which include several subsystems that are interconnected and interact with each other [6]. The characteristic properties of such systems are non-linearity, multidimensionality, multi-connectivity and multifunctionality, i.e., during normal operation, both the layout of the system and the dynamic properties of the separate subsystems themselves change [7]. Examples of such modern systems are multi-connected automatic control systems (MCAS) for complex dynamic objects, such as aircraft gas turbine engine, power complexes, synchronous generators, and so on. According to the American company Honeywell, which analyzed the operation of more than 100000 control loops in 350 production processes, about 49 % of the control loops are configured incorrectly or erroneously [8]. The main difficulty in this case lies in ensuring the stability and the desired quality of functioning of both the MCAS as a whole and its separate subsystems in various operating modes [9]. Therefore, achieving the desired quality of functioning of a multiply connected system is an urgent practical and theoretical task. In the article, to solve this problem, it is proposed to use logical controllers in separate channels.

In modern MCAS, to improve the dynamic properties, nonlinear elements and connections are often used, implemented in the form of nonlinear controllers [10]. The use of nonlinear algorithms significantly expands the possibilities of purposefully changing the quality of control processes, and also improves the dynamic and static properties of the system [11].

Among this class of regulators, regulators with logical switching of transmission coefficients either in a direct circuit or in a feedback circuit are widely used [12]. Switching in such systems occurs at certain ratios of the coordinates of the system, which are determined by the logical control law. There are many different logical control laws [13–15], which have in common that switching occurs depending on the value of the error coordinate $\varepsilon(t)$ and its derivative $\varepsilon'(t)$. However, these logic control laws are developed for systems with one input and output, and do not take into account the mutual influence of separate channels, which is typical for MCAS. And also, for them it is necessary to calculate the values of the coefficients each time when the parameters of the control object change to ensure high quality control. Under these conditions, the use of the apparatus of neural networks is appropriate.

3. Synthesis of a multidimensional neural network controller for helicopters aircraft turboshaft engines

It is assumed that the dynamic properties of helicopters aircraft TE as multidimensional control objects are described by the following differential equations "input-output":

$$\varphi = \left(\mathbf{Y}^{(n)}, \mathbf{Y}^{(n-1)}, \dots, \mathbf{Y}; \mathbf{U}^{(n)}, \mathbf{U}^{(n-1)}, \dots, \mathbf{U} \right); \quad (1)$$

where $\mathbf{U} = (u_1(t), u_2(t), \dots, u_N(t))^T$, $\mathbf{Y} = (y_1(t), y_2(t), \dots, y_N(t))^T$ – respectively, vectors of inputs (control actions) and outputs (controlled variables); m and n – maximum orders of derivatives $u_k^{(i)}$,

$y_e^{(j)}$ for input and output variables $u_k(t)$ and $y_e(t)$, ($m \leq n$); N – number of engine control channels, that is, the dimension of the ACS; $\varphi(\cdot)$ – nonlinear vector function.

It is also considered that for helicopters aircraft TE, the condition of observation and control is made [16]. In [17–19], a block diagram of a multidimensional ACS for helicopters aircraft TE (using the TV3-117 aircraft engine as an example) is proposed with the inclusion of N integrators (I) in the system – one in each of the N channels of the control system (fig. 1), implemented in the form of a multi-mode neural network controller using a dynamic recurrent neural network based on a perceptron (fig. 2), which provides control of the object (1), subject to the following requirements for the synthesized ACS:

- astatism (zero static error);
- physical implementation of the neural network controller;
- stability and given quality of control processes on a fixed set of modes $M = \{M_1, \dots, M_R\}$ of engine operation;
- minimum complexity of the multidimensional neural network controller.

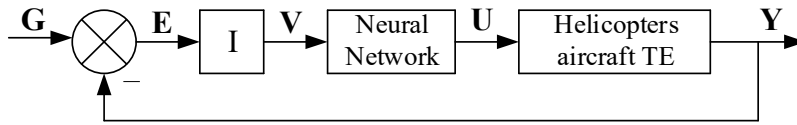


Figure 1: Structural diagram of the multidimensional ACS of helicopters aircraft TE

The requirement of astaticism of a multidimensional ACS is reduced to the inclusion of N integrators (I) in the system – one in each of the N channels of the control system. Then for the circuit in fig. 1 can be accepted:

$$V_i(z) = \frac{T_0}{1 - z^{-1}} E_i(z); \quad (2)$$

where $i = 1, 2, \dots, N$; z^{-1} – time shift operator, one cycle delay T_0 ; $\frac{T_0}{1 - z^{-1}}$ – integrator discrete transfer function.

The requirement of the physical capability of the neural network controller to be implemented is based on the assumption that a dynamic recurrent neural network based on a perceptron is taken as a neural network (fig. 2).

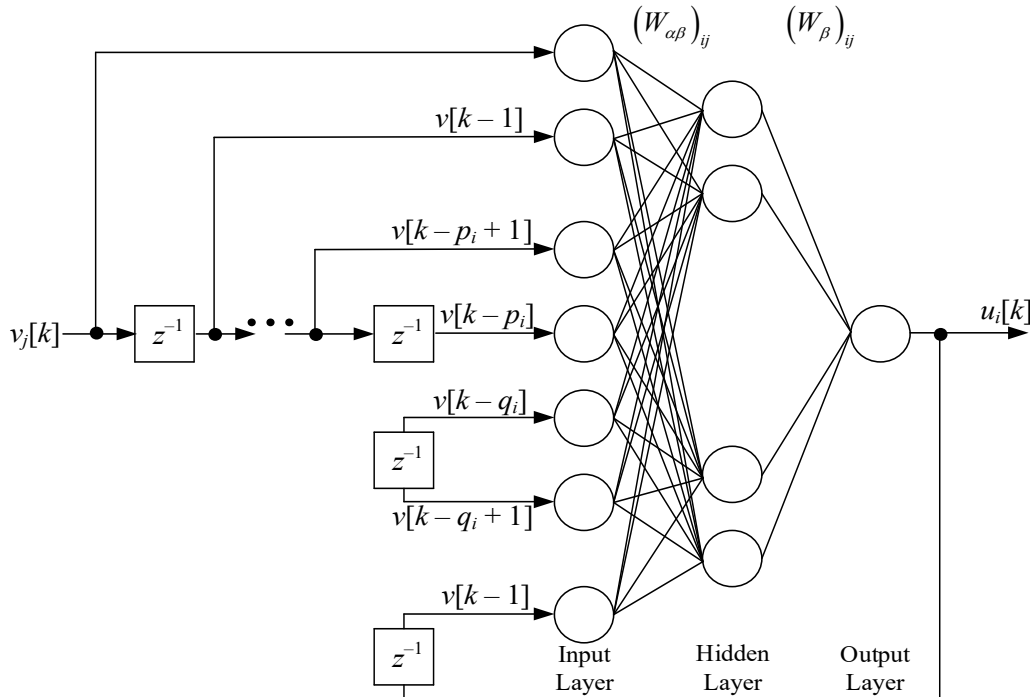


Figure 2: Multimode neural network controller based on perceptron

The neural network controller has N inputs and N outputs, including $N + \sum_{i=1}^N (p_i + q_i)$ neurons in the input layer, σ neurons in the common hidden layer, and N neurons in the output layer, the connections between which are carried out using adjustable (training) weights $W_{\alpha\beta}$, W_{β} ($\alpha = 1, 2, \dots, N + \sum_{i=1}^N (p_i + q_i)$; $\beta = 1, 2, \dots, N_{\sigma}$). It is obvious that the number of neurons in the hidden layer σ must satisfy the constraints $\sigma > N$ (otherwise, it is impossible to provide independent formation of control actions $u_1, u_2, \dots, u_N$ when changing the neural network inputs $v_1, v_2, \dots, v_N$) [17–19]. The total number of unknown parameters of the neural network controller, that is, the number of adjustable neural network weights, is:

$$K_p = \sigma \left(2N + \sum_{i=1}^N (p_i + q_i) \right). \quad (3)$$

To ensure stability on a given set of $M_1, \dots, M_R$ modes of operation of helicopter aircraft TE ACS, it is necessary to fulfill the following relationship between the values of p_i, q_i, σ, n, R and N :

$$\sigma \left(2N + \sum_{i=1}^N p_i \right) + (\sigma - R) \sum_{i=1}^N q_i \geq R(N + n); \quad (4)$$

from which it is possible to determine the unknown integer values p_i, q_i and σ , which determine the structure of the neural network controller.

To fulfill the requirement of the minimum complexity criterion, it is assumed that the complexity of the neural network controller is determined by the number of adjustable neural network parameters (K_p), the desired solution to the structural synthesis problem based on the minimum complexity criterion should be considered a neural network controller described by a set of numbers $\langle p_1, \dots, p_N; q_1, \dots, q_N; \sigma \rangle$ minimizing the value of the objective function (3) when executing the constraint (4).

For helicopters aircraft TE, the vector of inputs (control actions) has the form:

$$\mathbf{U} = (G_T)^T; \quad (5)$$

where G_T – fuel consumption, and the state vector and the vector of outputs (controlled variables) of the engine are written, respectively, as $\mathbf{X} = (x_1, x_2)^T$ and $\mathbf{Y} = (n_{TC}, T_G^*)$, where n_{TC} – gas generator r.p.m.; T_G^* – gas temperature in front of the compressor turbine.

The block diagram of the ACS is shown in fig. 3, where D_n and D_T – sensors gas generator r.p.m. and gas temperature in front of the compressor turbine; Act_{G_T} – actuator that provides the formation of actions along the coordinate $\overline{G_T}$; $\mathbf{G} = (\overline{n_{TC}}, \overline{T_G^*})^T$ – vector of settings (given influences); $\overline{n_{TC}}^0$ and $\overline{T_G^*}^0$ – required (tasks) values gas generator r.p.m.; T_G^* – gas temperature in front of the compressor turbine.

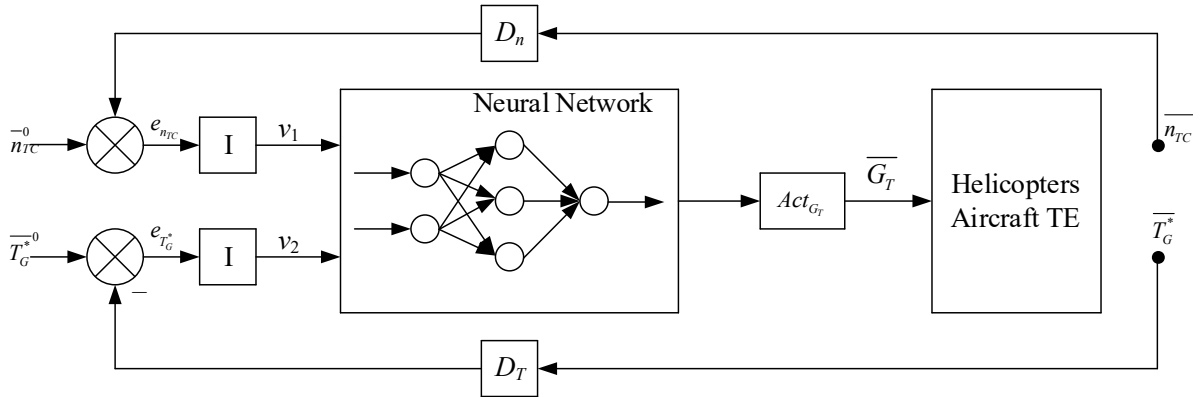


Figure 3: Structural diagram of synthesized ACS for helicopters aircraft TE

4. Modification of automatic control system for helicopters turboshaft engines

According to the concept developed in [20], the actuating mechanism (AM) and TE were considered as a single whole: an invariable part of the system. This approach has proven itself in the synthesis of TE control algorithms for helicopters of civil aviation or transport aviation. For such control objects, the dynamic processes in the fuel system proceed much faster than in the engine; therefore, their influence on TE was simply neglected. But in TE, transient processes in the fuel and engine assembly occur almost simultaneously. This statement has been repeatedly confirmed by the results of full-scale tests [21]. On the basis of the foregoing, we single out the TE and AM directly into separate links – the fuel metering unit (FMU) and modify the structural diagram of the ACS shown in fig. 1 (fig. 4).

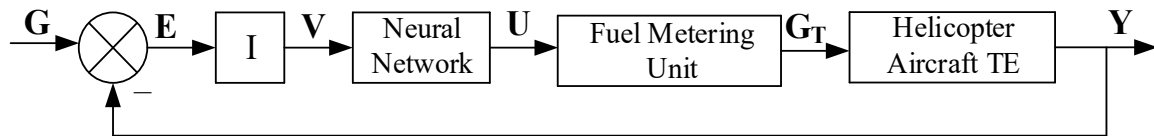


Figure 4: Structural diagram of a modified multidimensional ACS of helicopters aircraft TE with a divided control object on FMU and TE

When conducting a simple study of the operation of the ACS TE, which consists in various combinations of parameters for transfer functions for TE and FMU, it was found that the quality of control (accuracy, overshoot, stability margins) changes dramatically when switching from mode to mode. Thus, the task of analyzing the quality of control and synthesizing control algorithms for objects of this class becomes very relevant.

In this paper, the ACS of TE is studied and the quality of control is analyzed taking into account the dynamics of the FMU and TE. Consider the automatic control system of the gas turbine engine shown in fig. 4. The system consists of a comparison element (CE), a regulator, FMU and TE. The initial value of r.p.m. and gas temperature in front of the compressor turbine and the obtained values of the number of these parameters are received at the input of the CE, the inconsistency of the incoming parameters is formed at the output and the system error is formed – ξ [20, 21]. The error is fed to the input of the controller, the signal u is generated at the output, which is fed to the input of the FMU, the signal of fuel consumption G_T is generated at the output, which is fed to the input of TE and, accordingly, the signal Y is generated, which is fed to the input of the CE. Taking into account that in the proposed ACS scheme of TE the control object was divided, it is advisable to introduce nonlinear models separately for the TE and FMU and simulate the operation of the system, taking into account the dynamics of its elements. In order to investigate the above-described ACS TE, it is also proposed to introduce mathematical models of the FMU and TE into the structure of the system in order to improve the quality of control of the entire system as a whole [20, 21].

On fig. 5 shows the ACS TE scheme developed in this work. In the logical block (LB) the input signals are analyzed as follows: a knowledge base is built on the basis of experimental data and conclusions. In relation to it, membership functions are formed for the input parameters of the LB, as well as output signals [22]. Having formed the necessary change, the LB sends response signals to the input of the comparison element, forming a control signal that is fed to the input of the FMU and its model. The LB receives two signals: the inconsistency of the FMU and TE models with the FMU and TE models – model error (ξ_{mod}) and the inconsistency of the FMU with the FMU model – FMU error (ξ_{FMU}). As practice shows, the TE error is small and is not taken into account in the course of the study.

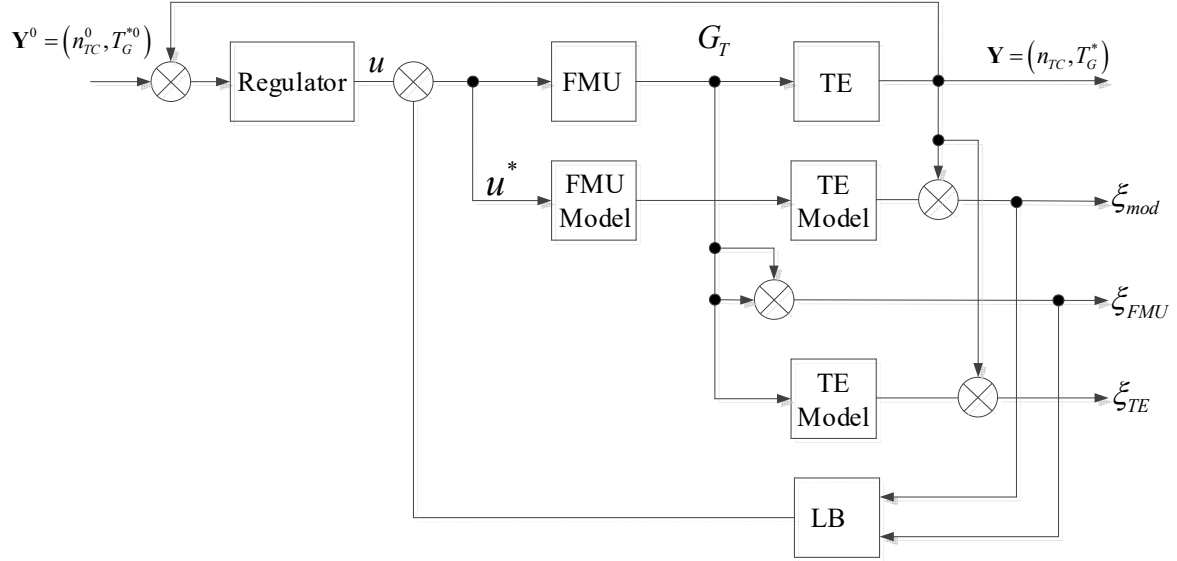


Figure 5: Proposed ACS TE, including the regulator, FMU, TE, FMU model, TE model and LB

5. Mathematical description of helicopters aircraft turboshaft engines automatic control system

Synthesis of one channel of the automatic control system.

Professor Valery Petunin made a significant contribution to the synthesis of logical-dynamic systems of automatic control of gas turbine engines based on the coordination and adaptation of control channels [23–27]. According to [23], the block diagram of a separate channel of ACS TE has the form of fig. 6.

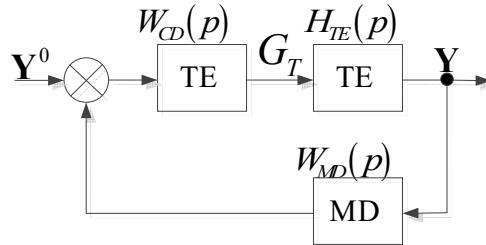


Figure 6: Structural diagram of a separate ACS channel

The transfer function of a closed ACS is determined by the expression:

$$\Phi(p) = \frac{W_{CD}(p) \cdot H_{TE}(p)}{1 + W_{CD}(p) \cdot H_{TE}(p) \cdot W_{MD}(p)}. \quad (6)$$

The transfer function of the measuring device must be equal to unity, i.e., $W_{MD}(p) = 1$. Let us equate the transfer function of the closed system to the transfer function of the desired system $\Phi^*(p)$,

i.e., $\Phi(p) = \Phi^*(p)$, then $\Phi^*(p) = \frac{W_{CD}(p) \cdot H_{TE}(p)}{1 + W_{CD}(p) \cdot H_{TE}(p) \cdot W_{MD}(p)}$. After transformations, we obtain the

following expression for the transfer function of the control device [24]:

$$W_{CD} = \frac{1}{H_{TE}(p)} \cdot \frac{\Phi^*(p)}{1 - \Phi^*(p)}; \quad (7)$$

where $\frac{1}{H_{TE}(p)}$ – inverse transfer function of the control object; $\frac{\Phi^*(p)}{1 - \Phi^*(p)}$ – desired open-loop transfer function.

For one channel, the GTE transfer function $H_{TE}(p) = k_{TE} \cdot \frac{A(p)}{B(p)}$, then $\frac{1}{H_{TE}(p)} = \frac{1}{k_{TE}} \cdot \frac{B(p)}{A(p)}$. If an open-loop transfer function $W^*(p)$ is required:

$$W^*(p) = \frac{\Phi^*(p)}{1 - \Phi^*(p)} = \frac{k}{p \cdot C(p)}; \quad (8)$$

where $\frac{k}{p}$ – determines system astatism, and $\frac{1}{C(p)}$ – its inertia, then the transfer function of the control device will take the form:

$$W_{CD} = \frac{1}{k_{TE}} \cdot \frac{B(p)}{A(p)} \cdot \frac{k}{p \cdot C(p)}. \quad (9)$$

In the ACS of complex technical objects, professor Valery Petunin proposed the implementation of a common isodromic controller [23] with a transfer function:

$$W_{IC}(p) = \frac{k_{IC} \cdot (T_{IC} \cdot p + 1)}{p \cdot C(p)}; \quad (10)$$

where the boosting link $(T_{IC} \cdot p + 1)$ corrects the inertia of TE and regulators of individual channels:

$$W_{MD} = \frac{1}{k_{TE_i}} \cdot \frac{B(p)}{A_i(p)} \cdot \frac{k}{k_{IC} \cdot (T_{IC} \cdot p + 1)}. \quad (11)$$

Mathematical description of the desired automatic control system.

The block diagram of the desired ACS of one channel is shown in fig. 7, where $W^*(p)$ – transfer function of the desired open-loop system.

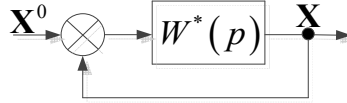


Figure 7: Structural diagram of the desired ACS

Then the transfer function of the closed desired system can be represented as:

$$\Phi^*(p) = \frac{W^*(p)}{1 + W^*(p)}. \quad (12)$$

According to (8), $W^*(p) = \frac{\Phi^*(p)}{1 - \Phi^*(p)}$, and if $W^*(p) = \frac{k}{p \cdot C(p)}$, then

$$\Phi^*(p) = \frac{k}{p \cdot C(p) + k}. \quad (13)$$

The inertia of the actuators must be taken into account in $C(p)$, if necessary, it can be adjusted. The transfer function of a closed system $\Phi^*(p)$ should be close to the standard transfer function, taking into account the requirements for the quality of transient processes.

Synthesis of controlling channel for gas generator r.p.m.

For the fuel dosing circuit into the main combustion chamber by the speed channel, the transfer function of helicopter aircraft TE can be represented as:

$$H_{nG_T}(p) = k_n \cdot \frac{A_n(p)}{B(p)}; \quad (14)$$

where the order of the polynomial $A_n(p)$ is one less than the order of $B(p)$ (table 1).

Table 1

Correspondence of polynomial order on the number of shafts of helicopters aircraft TE

Number of helicopters aircraft TE shafts	Polynomial order	
1	0	1
2	1	2

Transfer function of a one-shaft TE (for example, GTE-350) in terms of gas generator r.p.m. can be obtained according to [23] in the form:

$$H_{n_{G_T}}(p) = \frac{0,4}{0,5 \cdot p + 1}. \quad (15)$$

Then the transfer function of the control device for gas generator r.p.m. channel is represented as:

$$W_{CDn}(p) = \frac{1}{k_n} \cdot \frac{B(p)}{A_n(p)} \cdot \frac{k}{p \cdot C(p)}. \quad (16)$$

If the transfer function of the desired system $W^*(p) = \frac{3}{p \cdot (0,02 \cdot p + 1)}$ and the transfer function of the general isodromic controller $W_{IC}(p) = \frac{3 \cdot (0,56 \cdot p + 1)}{p \cdot (0,02 \cdot p + 1)}$ [23, 25], then the transfer function of the controller over the channel of gas generator r.p.m.:

$$W_n(p) = W_1(p) = \frac{1}{0,4} = 2,5. \quad (17)$$

For emergency operation, the transfer function of two-shaft TE (for example, TV3-117) in terms of gas generator r.p.m. can be obtained in accordance with [23, 25]:

$$H_{n_{TC}G_T}(p) = \frac{0,186 \cdot p + 0,875}{0,133 \cdot p^2 + 0,94 \cdot p + 1}. \quad (18)$$

Then the transfer function of the control device for gas generator r.p.m. channel is:

$$W_{CDn_{TC}}(p) = \frac{1}{k_{n_{TC}}} \cdot \frac{B(p)}{A_{n_{TC}}(p)} \cdot \frac{k}{p \cdot C(p)}. \quad (19)$$

If the transfer function of the desired system $W^*(p) = \frac{3}{p \cdot (0,02 \cdot p + 1)}$ and the transfer function of the general isodromic controller $W_{IC}(p) = \frac{3 \cdot (0,766 \cdot p + 1)}{p \cdot (0,02 \cdot p + 1)}$ [23, 25], then the transfer function of the controller over the channel of gas generator r.p.m.:

$$W_{n_{TC}}(p) = W_1(p) = \frac{1}{0,766} \cdot \frac{0,175 \cdot p + 1}{0,210 \cdot p + 1} = \frac{0,229 \cdot p + 1,306}{0,210 \cdot p + 1}. \quad (20)$$

Synthesis of the gas temperature control channel before the compressor turbine.

For the fuel dosing circuit into the combustion chamber by the gas temperature channel before the compressor turbine, the transfer function of helicopter aircraft TE can be represented as:

$$H_{T_G^*G_T}(p) = k_{T_G^*} \cdot \frac{A_{T_G^*G_T}(p)}{B(p)}; \quad (21)$$

where $A_{T_G^*G_T}(p)$ and $B(p)$ are polynomials of the same order (table 2).

Table 2

Correspondence of polynomial order on the number of shafts of helicopters aircraft TE

Number of helicopters aircraft TE shafts	Polynomial order	
1	1	1
2	2	2

For emergency operation, the transfer function of a one-shaft TE in terms of gas temperature in front of the compressor turbine can be obtained according to [23, 25]:

$$H_{T_G^*G_T}(p) = 0,35 \cdot \frac{0,829 \cdot p}{0,56 \cdot p + 1} = \frac{0,29 \cdot p}{0,56 \cdot p + 1}. \quad (22)$$

Then the transfer function of the control device for the gas temperature channel before the compressor turbine has the form:

$$H_{IC_{T_G}^*}(p) = \frac{1}{k_{T_G}^*} \cdot \frac{B(p)}{A_{T_G^* G_T}(p)} \cdot \frac{k}{p \cdot C(p)}. \quad (23)$$

If the transfer function of the desired system $W^*(p) = \frac{3}{p \cdot (0.02 \cdot p + 1)}$ and the transfer function of the general isodromic controller $W_{IC}(p) = \frac{3 \cdot (0.56 \cdot p + 1)}{p \cdot (0.02 \cdot p + 1)}$ [23, 25], then transfer function of the controller through the gas temperature channel before the compressor turbine will look like:

$$W_{T_G}^*(p) = W_2(p) = \frac{1}{0.35} \cdot \frac{1}{0.829 \cdot p + 1} = \frac{2.857}{0.829 \cdot p + 1}. \quad (24)$$

For emergency operation, the transfer function of a two-shaft TE in terms of gas temperature in front of the compressor turbine can be obtained in accordance with [13, 15]:

$$H_{T_G^* G_T}(p) = 0.333 \cdot \frac{0.064 \cdot p^2 + 0.667 \cdot p + 1}{0.133 \cdot p^2 + 0.94 \cdot p + 1} = \frac{0.021 \cdot p^2 + 0.222 \cdot p + 0.333}{(0.766 \cdot p + 1) \cdot (0.174 \cdot p + 1)}. \quad (25)$$

Then the transfer function of the control device for the gas temperature channel before the compressor turbine looks like this:

$$H_{IC_{T_G}^*}(p) = \frac{1}{k_{T_G}^*} \cdot \frac{B(p)}{A_{T_G^* G_T}(p)} \cdot \frac{k}{p \cdot C(p)}. \quad (26)$$

If the transfer function of the desired system $W^*(p) = \frac{3}{p \cdot (0.02 \cdot p + 1)}$ and the transfer function of the general isodromic controller $W_{IC}(p) = \frac{3 \cdot (0.56 \cdot p + 1)}{p \cdot (0.02 \cdot p + 1)}$ [23, 25], then transfer function of the controller through the gas temperature channel before the compressor turbine will look like:

$$W_{T_G}^*(p) = W_2(p) = \frac{1}{0.333} \cdot \frac{0.174 \cdot p + 1}{0.064 \cdot p^2 + 0.667 \cdot p + 1} = \frac{0.522 \cdot p + 3}{0.064 \cdot p^2 + 0.667 \cdot p + 1}. \quad (27)$$

The principles of construction of fast-response GTE gas temperature meters were studied in [17].
Synthesis of cross-links of control channels.

According to [23], it is assumed that $W_{K1}(p)$ and $W_{K2}(p)$ – transfer functions of corrective links, the results of the synthesis of which are given in [16], which have the form:

$$W_{K1}(p) = \frac{W_1(p) - 1}{W_2(p)}; \quad W_{K2}(p) = \frac{1 - W_2(p)}{W_2(p)}. \quad (28)$$

Then the expressions for the transfer functions for these links look like this:

1. For a one-shaft TE:

$$W_{K1}(p) = 0.428 \cdot (0.829 \cdot p + 1) = 0.355 \cdot p + 0.428; \quad W_{K2}(p) = 0.35 \cdot (0.829 \cdot p - 1.857) = 0.29 \cdot p - 0.65.$$

2. For a two-shaft TE:

$$W_{K1}(p) = \frac{-0.0018 \cdot p + 0.0517}{0.206 \cdot p + 1} \cdot \frac{0.064 \cdot p^2 + 0.667 \cdot p + 1}{0.174 \cdot p + 1}; \quad W_{K2}(p) = 0.333 \cdot \frac{0.064 \cdot p^2 + 0.145 \cdot p - 2}{0.174 \cdot p + 1}.$$

According to [23], the corrective element $W_1(p)$ of the ACS control channel is permanently switched on, and the corrective element $W_2(p)$ is implemented when the limitation channel is turned on by parallel connection to $W_1(p)$ of a differential dynamic link with a transfer function according to the output logical signal of the selector L :

$$W_{\Delta}(p) = W_1(p) - W_2(p). \quad (29)$$

$$1. \text{ For a one-shaft TE: } W_{\Delta}(p) = \frac{1.842 \cdot p - 0.653}{0.829 \cdot p + 1}.$$

$$2. \text{ For a two-shaft TE: } W_{\Delta}(p) = \frac{0.174 \cdot p + 1}{0.206 \cdot p + 1} \cdot \frac{0.0739 \cdot p^2 + 0.152 \cdot p - 1.845}{0.064 \cdot p^2 + 0.667 \cdot p + 1}.$$

6. Development of a neural network for implementation of multidimensional automatic control system for helicopters turboshaft engines

Fig. 8 shows a block diagram of the proposed self-adjusting neural network control system for helicopters aircraft TE with interconnected coordinates.

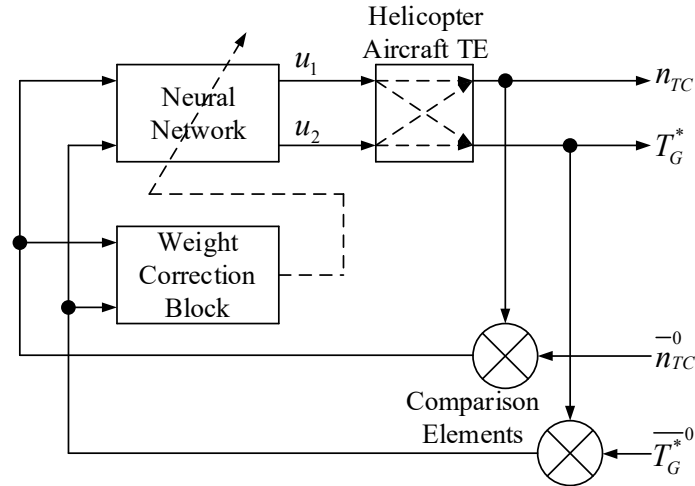


Figure 8: Proposed neural network control system for helicopters aircraft TE with interconnected coordinates

The control error vector $\mathbf{e} = [e_1, e_2, \dots, e_n]^T$ after the comparison elements is fed to the input of the neural network and the weight correction block, in which, depending on the signal $\mathbf{e}(t)$, at each discrete time t , the weight coefficients of the neural network are adjusted. networks. The vector of signals $\mathbf{u} = [u_1, u_2, \dots, u_n]^T$ from neural network output is the control one and is fed to the input of the control object (helicopter aircraft TE). The neural network for controlling aviation gas turbine engines of helicopters as multiply connected objects is a multilayer neural network with one intermediate layer containing N_0 neurons in the input layer and N_2 neurons in the output layer, with $N_2 = N_0 = n$. The network is characterized by the number of neurons N_1 in the inner layer (layer 1) [28]. The structure of a multilayer neural network is shown in fig. 9.

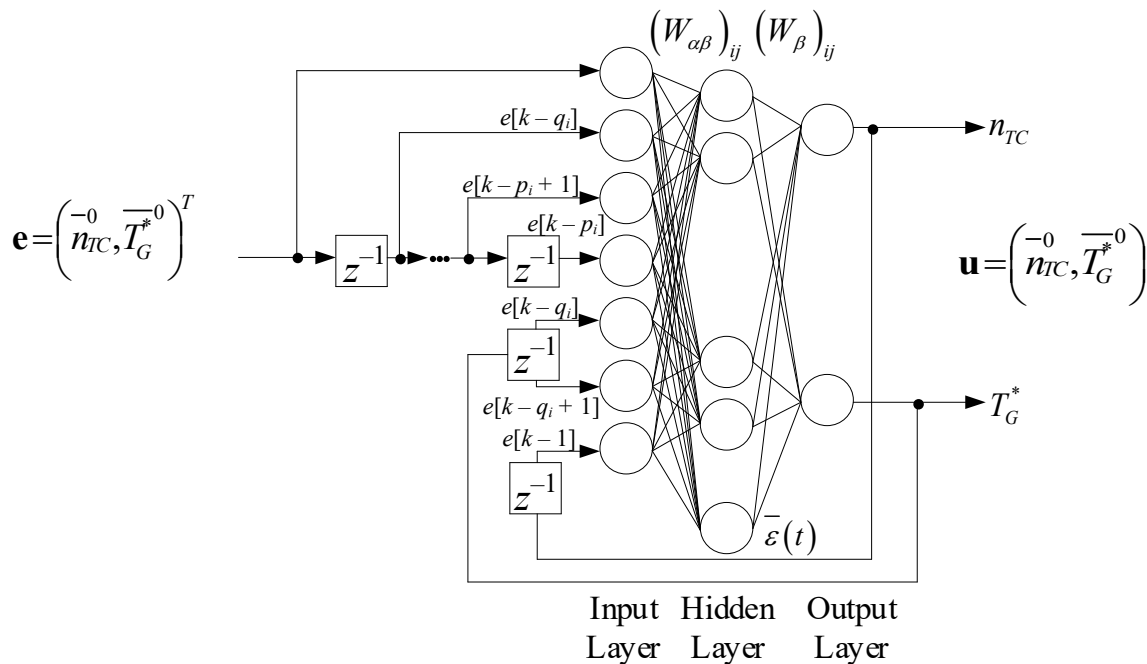


Figure 9: Structure of the neural network of a self-tuning system regulation

The input layer (layer 0) consists of nodes – signal receivers e_i ($i = \overline{1, n}$), the output layer – of neurons – signal sources u_i ($i = \overline{1, n}$). Signals e_i are fed to the input of the neural network of a discrete moment of time t ($t = 0, 1, 2, \dots$), which are converted by the network into control signals u_i . The discrete time t is related to the continuous time Θ as follows: $\Theta = \delta \cdot t$, where δ – quantization step. Each i neuron of the l -th layer ($l = \overline{1, 2}$) transforms the input vector into the original scalar value. At the first stage, the superposition of the input signals of the neuron is calculated:

$$z_i^l = \sum_{j=1}^{N_{l-1}} w_{ij}^l \cdot o_j^{l-1} - g_i^l; \quad (30)$$

where w_{ij}^l – weight coefficient, which is an adjustable parameter and characterizes the connection of the j -th neuron ($l-1$) of the layer with the i -th neuron of the l -th layer; g_i^l – shift amount.

Assuming that $w_{i0}^l = -g_i^l$ and $o_0^{l-1} = 1$, expression (30) is rewritten as:

$$z_i^l = \sum_{j=1}^{N_{l-1}} w_{ij}^l \cdot o_j^{l-1}. \quad (31)$$

Next, the value of z is converted to the initial value of the neuron:

$$o_i^l = f(z_i^l). \quad (32)$$

The nonlinear transformation (32) is defined by an activation function, often defined as a sigmoid function:

$$f(z) = \frac{1}{1 + e^{-z}}. \quad (33)$$

An important property of this function is the simplicity of determining the derivative of this function, i.e.

$$f'(z) = f(z) \cdot (1 - f(z)). \quad (34)$$

With the accepted notation, the mathematical description of the neural network is written using the system of equations:

$$\begin{cases} u_j(t) = o_j^2; j = \overline{1, n}; \\ o_i^l = f(z_i^l); i = \overline{1, N_l}; l = \overline{1, 2}; \\ z_i^l = \sum_{j=1}^{N_{l-1}} w_{ij}^l \cdot o_j^{l-1}; i = \overline{1, N_l}; l = \overline{1, 2}; \\ o_j^0 = e_j(t); j = \overline{1, n}; \\ o_0^0 = o_0^1 = 1. \end{cases} \quad (35)$$

During the operation of a self-adjusting neural network control system, the system settings – the weight coefficients of the neural network change in such a way that the value $E = \|\mathbf{e}\| \rightarrow 0$, while the value of $E = \frac{1}{2} \cdot \sum_{i=1}^n \alpha_i \cdot e_i^2$, where α_i – coefficients that determine the weight of each control channel of the total error E , is taken as the norm of the vector E .

Correction of neural network weight coefficients w_{ij}^l (training of the neural network) is carried out in the block for correcting the weight coefficients by backpropagation error method [29, 30]. The main calculated ratios in this case have the form:

$$w_{ij}^l(t) = w_{ij}^l(t-1) - \gamma \cdot \delta \cdot \frac{\partial E}{\partial w_{ij}^l}. \quad (36)$$

For the weight coefficients of the neuron of the original layer (layer 2), the values $\frac{\partial E}{\partial w_{ij}^{(2)}}$ are determined according to the expression:

$$\frac{\partial E}{\partial w_{ij}^{(2)}} = \sum_{k=1}^n \left(\alpha_k \cdot e_k \cdot \frac{\partial y_k}{\partial u_i} \right) \cdot f'(z_i^{(2)}) \cdot o_j^{(1)}; \quad k = \overline{1, n}; \quad i = \overline{1, n}; \quad j = \overline{1, N_1}; \quad (37)$$

where e_k – control error for the k -th initial variable; α_k – weight coefficient for the k -th output variable; $\frac{\partial y_k}{\partial u_i}$ – derivative of the k -th output variable of the object with respect to the i -th input action;

$f'(z_i^{(2)})$ – derivative of the activation function for the i -th neuron of the second (initial) layer; $o_j^{(1)}$ – initial value of the j -th neuron of the first layer.

For the inner layer (layer 1), the values $\frac{\partial E}{\partial w_{ij}^{(1)}}$ are determined by the scalar product:

$$\frac{\partial E}{\partial w_{ij}^{(1)}} = \left(\bar{\mathbf{a}} \times \mathbf{e}, \mathbf{Y}_u \times \mathbf{U}_{w_{ij}^{(1)}} \right); \quad i = \overline{1, N_1}; \quad j = \overline{0, n}; \quad (38)$$

where $\bar{\mathbf{a}}$ – matrix of weight coefficients of variable adjustments; $\mathbf{e}$ – adjustment error vector; $\mathbf{Y}_u$ – matrix of derivative variables of regulation with respect to incoming control actions; $\mathbf{U}_{w_{ij}^{(1)}}$ – vector of derivatives of the output signals of the neural network by the weight $w_{ij}^{(1)}$ of the layer 1 neuron; $\bar{\mathbf{a}} \times \mathbf{e}$ – vector product of matrix $\bar{\mathbf{a}}$ and vector $\mathbf{e}$; $\mathbf{Y}_u \times \mathbf{U}_{w_{ij}^{(1)}}$ – vector product of matrix $\mathbf{Y}_u$ and vector $\mathbf{U}_{w_{ij}^{(1)}}$.

The presented vectors and matrices look like this:

$$\bar{\mathbf{a}} = \begin{pmatrix} \alpha_1 & 0 & \dots & 0 \\ 0 & \alpha_2 & \dots & 0 \\ M & M & \dots & M \\ 0 & 0 & \dots & \alpha_n \end{pmatrix}; \quad \mathbf{e} = \begin{pmatrix} e_1 \\ e_2 \\ M \\ e_n \end{pmatrix}; \quad \mathbf{Y}_u = \begin{pmatrix} \frac{\partial y_1}{\partial u_1} & \frac{\partial y_1}{\partial u_2} & \dots & \frac{\partial y_1}{\partial u_n} \\ \frac{\partial y_2}{\partial u_1} & \frac{\partial y_2}{\partial u_2} & \dots & \frac{\partial y_2}{\partial u_n} \\ M & M & \dots & M \\ \frac{\partial y_n}{\partial u_1} & \frac{\partial y_n}{\partial u_2} & \dots & \frac{\partial y_n}{\partial u_n} \end{pmatrix}; \quad \mathbf{U}_{w_{ij}^{(1)}} = \begin{pmatrix} \frac{\partial u_1}{\partial w_{ij}^{(2)}} \\ \frac{\partial u_2}{\partial w_{ij}^{(2)}} \\ M \\ \frac{\partial u_n}{\partial w_{ij}^{(2)}} \end{pmatrix}.$$

With the machine implementation of the presented learning algorithm, the matrix $\mathbf{Y}_u$ will be a matrix consisting of a set of zeros and ones, and the ones will be equal to the element corresponding to the main (direct) control channels. For an object with n inputs and outputs, the $\mathbf{Y}_u$ matrix might look like this:

$$\mathbf{Y}_u = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ M & M & \dots & M \\ 0 & 0 & \dots & 1 \end{pmatrix}. \quad (39)$$

The vector $\mathbf{U}_{w_{ij}^{(1)}}$ elements $\frac{\partial u_k}{\partial w_{ij}^{(1)}}$ are determined according to the expression:

$$\frac{\partial u_k}{\partial w_{ij}^{(1)}} = f'(z_k^{(2)}) \cdot w_{ki}^{(2)} \cdot f'(z_i^{(1)}) \cdot o_j^{(0)}; \quad k = \overline{1, n}; \quad i = \overline{1, N_1}; \quad j = \overline{0, n}. \quad (40)$$

Before the start of the self-adjusting control system, the system parameters are set: in the weight correction block – weight coefficients setting parameter γ , in the neural network block – number of neurons in the inner layer N_1 and the weight coefficients $w_{ij}^{(l)}$ of the neurons of layers 1 and 2. Coefficients $w_{ij}^{(l)}$ are selected by a random sensor $[-1, 1]$ according to the uniform distribution law.

7. Simulation studies of the neural network control system of helicopters aircraft turboshaft engines

TV3-117 aircraft TE, which is part of the power plant of Mi-8MTV helicopter with two inputs (n_{TC} , T_G^*), one output ($\bar{G}_T$) and the presence of significant cross-links (fig. 10) [17–19], was used as research object.

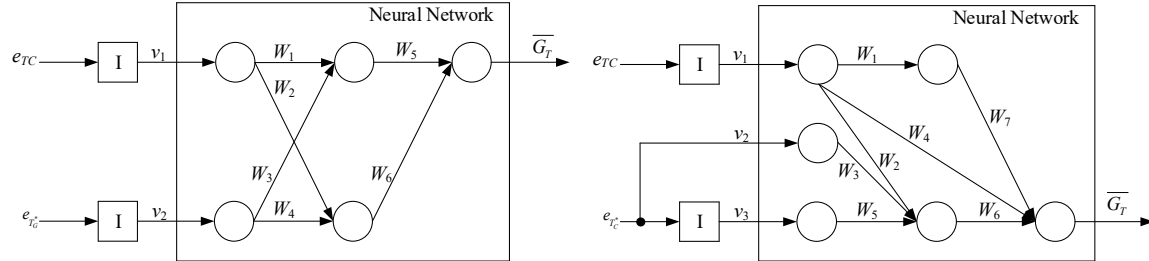


Figure 10: Structural diagram of control object (TV3-117 aircraft TE)

On fig. 11, the diagrams of transient processes in the system during testing set the impact without preliminary self-tuning according to the model (curves 1, 2) and with preliminary self-tuning (curves 3, 4). Quite often, at the initial stage of self-adjustment (when the neuroregulator is first switched on in the control loop), large outliers of controlled values from the set values are observed (fig. 11, curve 1). This is due to the fact that when the neuroregulator is turned on, the weight coefficients of the neural network are initialized randomly. In this case, there is a possibility of organizing positive feedbacks, which are the source of too large deviations of controlled values at the initial stage of self-tuning. The occurrence of such a situation can lead to emergency situations on the helicopter.

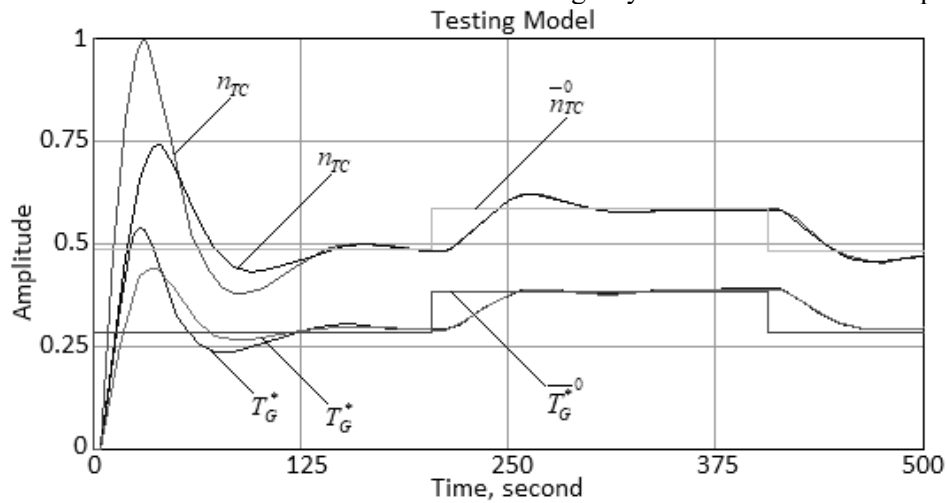


Figure 11: Influence of preliminary self-tuning according to the model on the quality of the control process

The availability of a priori information about helicopter aircraft TE, its dynamic and static characteristics will significantly reduce the deviation of controlled values in the process of self-tuning. This is ensured by the fact that at the previous stage the self-tuning of the system is carried out according to the approximate model of the gas turbine engine. Then there is a switch to work with the control object (helicopter aircraft TE), the channels of which can be described by aperiodic links of the first order with a delay. The parameters of TE model are determined approximately. An error in estimating the parameters within $\pm 100\%$ has little effect on subsequent self-tuning results, with the highest sensitivity observed behind the gain. The parameters of the control system built using a neural network are the neurons number in the hidden layer N_1 and the correction parameter of neural network weights coefficient γ . On fig. 12 show the influence of the tuning parameters N_1 and γ on the quality of the regulation process. With an increase in the values of the parameters N_1 and γ , the speed of the system increases, but the values of the overshoot and dynamic error also increase, and the value of the degree of extinction of transient processes decreases. With a significant increase in the values of the tuning parameters, undamped oscillations and instability of the control process may occur.

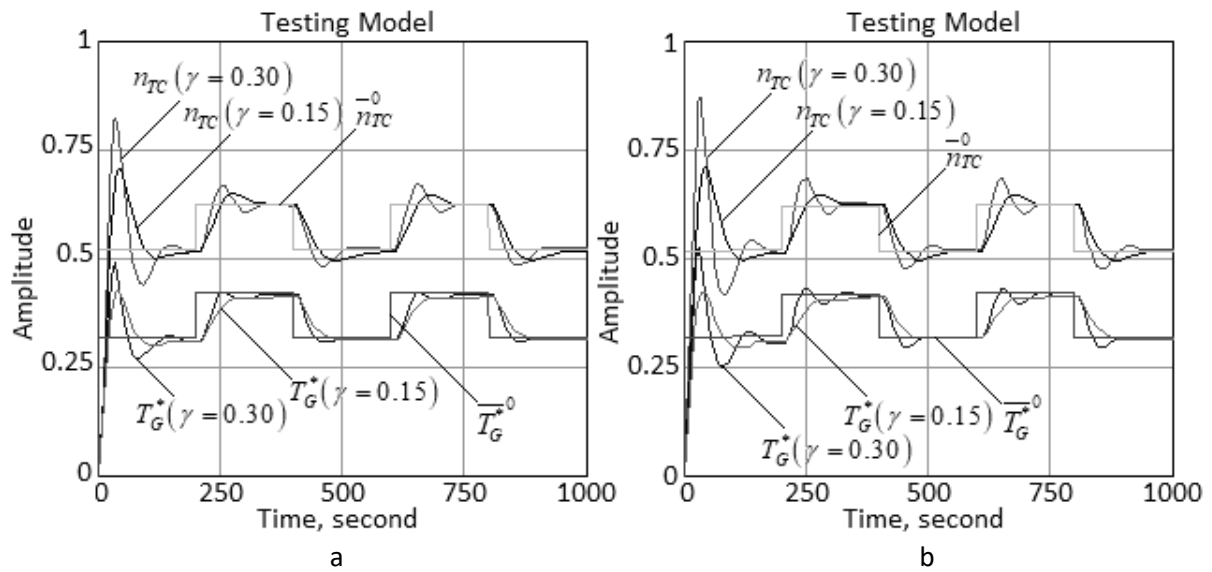


Figure 12: Influence diagrams: a – number of neurons N_1 in the intermediate layer on the quality of the regulation process; b – weight correction factor for the quality of the control process

Neural network control system for helicopters aircraft TE works out quite well both external disturbing influences (fig. 13, a), and setting influences (fig. 13, b), as well as antiphase setting actions (fig. 13, c).

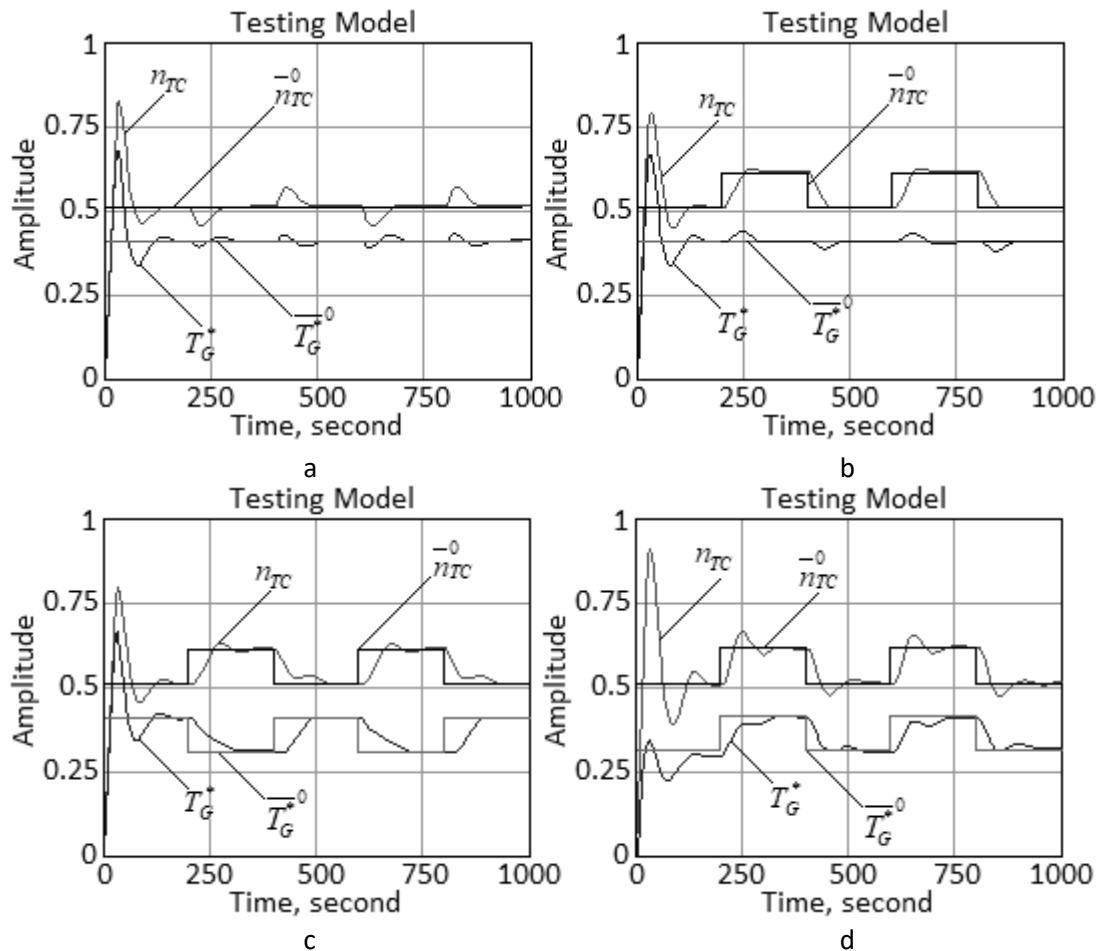


Figure 13: Process diagrams: a – square wave perturbations; b – one master influence; c – two setting influences going out of phase; d – setting influences when changing the gain coefficients of TV3-117 aircraft TE channels

In a neural network control system, as in classical systems, cross channels have a significant impact on the quality of regulation. With a decrease in the influence of cross-links (a decrease in the coupling coefficient $\frac{K_{12}K_{21}}{K_{11}K_{22}}$), the quality of regulation improves significantly. An important property

of a neural network control system is the possibility of its effective adaptation to changes in the properties of the control object. In the course of simulation studies, the influence of changing the properties of the control object on the quality of the regulation process in the system was considered. On fig. 13, d shows the transient process in the system during the development of two setting actions and in the form of square waves. At the same time, during the process, the gain factors along the channels of the control object changed: the coefficients K_{11} and K_{22} linearly decreased from 2.0 to 1.0 and from 1.5 to 1.0, respectively, and the coefficients K_{12} and K_{21} linearly increased from 0.3 to 0.5 and from 0.4 to 0.6, respectively. As can be seen from fig. 13, d, the high quality of regulation is maintained with a sufficiently large variation (30...100 %) of the gains of the object's channels, that is, the neural network control system adapts to changes in the gains of the object's channels.

The proposed neural network automatic control system of helicopters aircraft TE adapts well to changes in time constants and delays – changing them even by two or three times does not have a significant effect on changing the quality of regulation in the system.

8. Neural network training results

The input data for training the neural network is massive of n_{TC} and T_G^* parameters recorded on board the helicopter. *Accuracy*, *Precision*, *Recall*, *F*-measure metrics are used to assess the quality of neural network training. *Precision* can be interpreted as the proportion of parameters that the neural network called positive and at the same time are really positive, and *Recall* shows what proportion of the parameters of the positive class of all objects of the positive class was found by the algorithm [31]. The results of training the neural network according to the *Accuracy* and *Loss* indicators are shown in fig. 14. As can be seen from fig. 14, the *Accuracy* indicator approaches one, and *Loss* indicator – tends to zero, which indicates the high accuracy of the model and its minimal error.

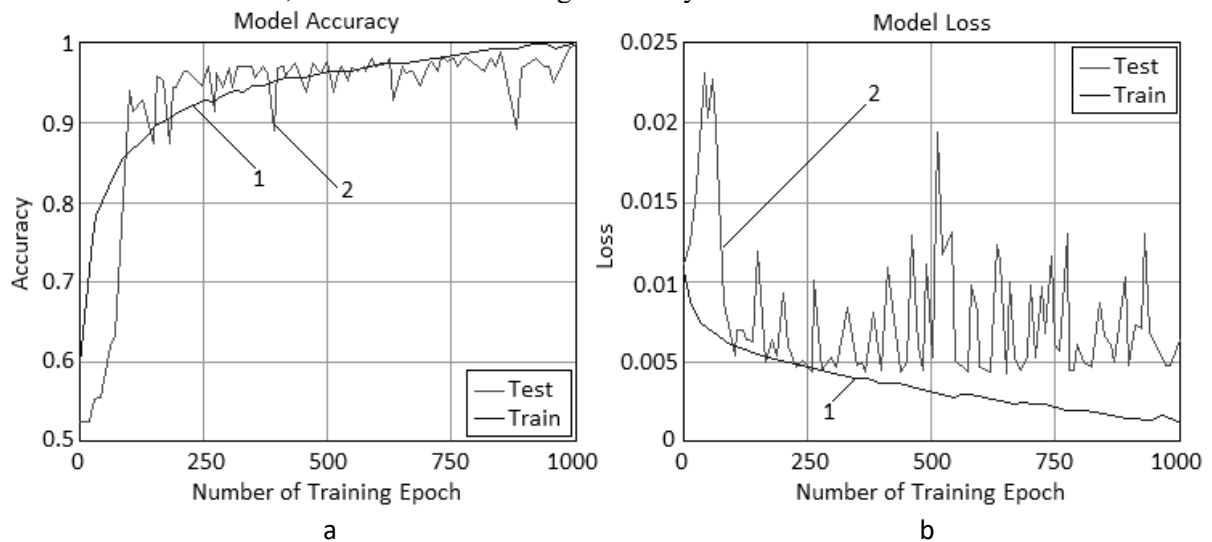


Figure 14: Neural network classifier training results (1 – test; 2 – train): a – *Accuracy* indicator; b – *Loss* indicator

9. Results and discussion

To assess the quality of ACS helicopters aircraft TE control, we introduce the following requirements:

- amplitude stability margin: not less than 20 dB;

- phase stability margin: from 35 to 80°;
- overshoot: no more than 5%;
- static error: no more than $\pm 5\%$ (± 0.05);
- regulation time: no more than 5 s.

When modeling the system (fig. 15, b), it was found that only with the values of the time constant (T) for the transfer functions of the FMU and TE: $T = 0.7$ s, $T = 0.5$ s, $T = 1$ s and transfer coefficient $k = 1$ the system works optimally, meeting the requirements of control quality and system stability. This indicates that the system changes parameters when operating in other modes, the quality of control of which may not meet the requirements. Therefore, we will take for ACS helicopters aircraft TE value of the time constant $T = 0.7$ s and the gain $k = 1$, and we will consider the system ideal, taken as a standard in the forthcoming study. Using the experimental data obtained during various passages of the routes, the points associated with the change in altitude and flight speed were selected: for a time of 50, 200, 500 s. According to [32], using the experimental data at the selected points, the values of the time constant and gain for the FMU and TE were obtained. When modeling in the ACS scheme of helicopters aircraft TE, the models of FMU and TE changed alternately with the obtained experimental parameters of the wind turbine and gas turbine engine, which made it possible to analyze the system according to the requirements described above. In the future, we will use the simulation time of 50 s, since it will be enough for the study.

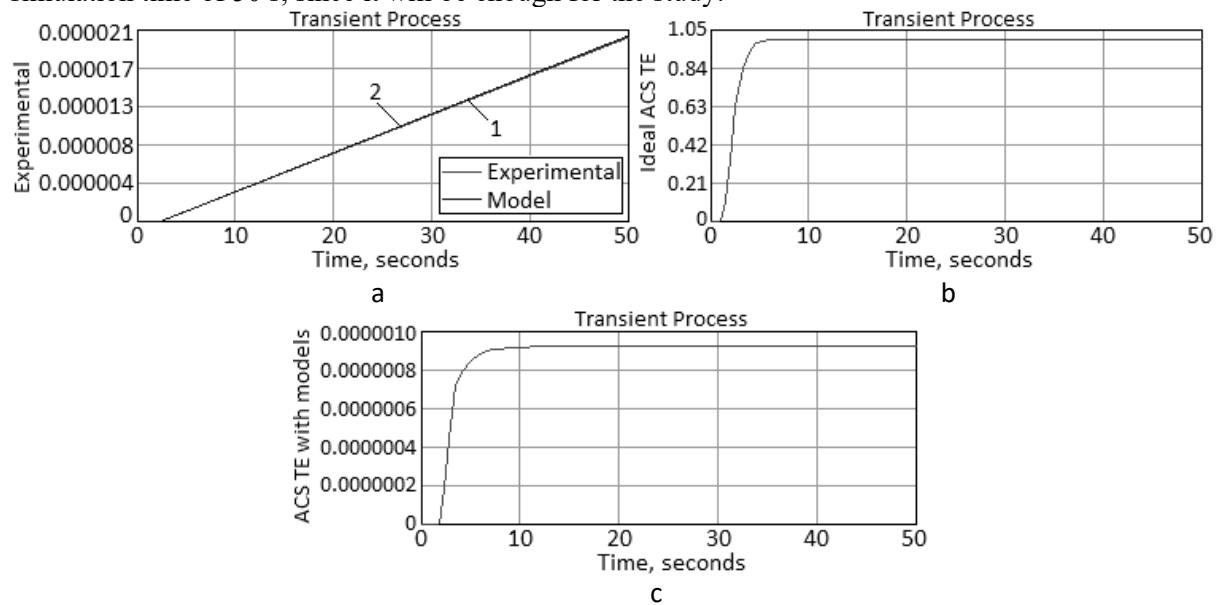


Figure 15: Results of the simulation of helicopters aircraft turboshaft engines automatic control system during the simulation time of 50 s: a – transient process of helicopters aircraft turboshaft engines automatic control system with experimental data (1), helicopters aircraft turboshaft engines automatic control system with models of FMU and TE (2); b – ideal ACS TE; c – ACS TE with models

The results of simulation of ACS TE for 50 s are shown in fig. 15. Modeling of the system was carried out in three stages: for an ideal scheme, with the parameters used in the design of helicopters aircraft turboshaft engines automatic control system, as well as for the system with experimental data and the system using the above approach with mathematical models of FMU and TE to adjust the operation of the entire system. As can be seen from the fig. 15, the transient process with ideal transfer function parameters for FMU and TE is established during the regulation time, which is 5 s; the system with experimental values is quite inertial and does not meet the requirements of control quality and stability, to adjust the helicopters aircraft turboshaft engines automatic control system, mathematical models of FMU and TE were introduced, which reduced the control time and began to meet the requirements. As can be seen from fig. 15, c, the transient process of the proposed ACS TE is inferior in quality: the value does not reach unity. Thus, in order to increase the accuracy of the transient process, it is proposed to introduce an LB based on fuzzy logic, the knowledge base and membership functions of which for input and output parameters will correspond to the graph of the dependence of errors on the control signal (fig. 16).

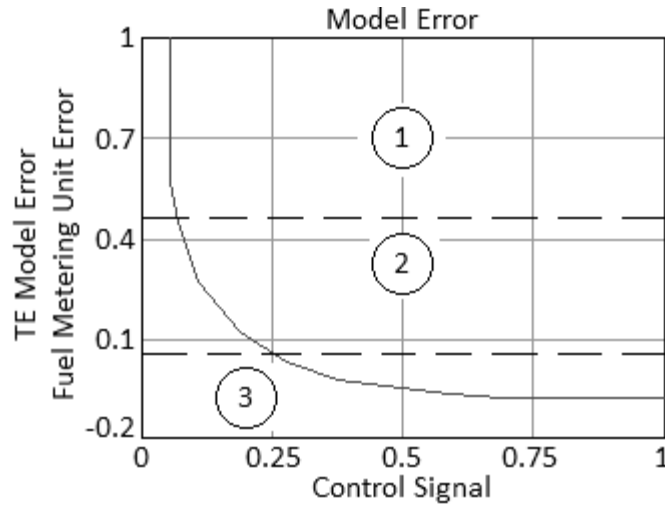


Figure 16: Dependence of model and ADT errors (ξ_{mod} , ξ_{FMU}) on control signal divided into zones: 1 – minimum, 2 – average, 3 – maximum

To ensure an acceptable nature of the transition process of the proposed ACS TE, it is proposed to introduce one more regulator: an integrating link. Experimental modeling showed that for the integrator the value of the gain (k) equal to 150 became sufficient to increase the quality of the output parameters. On fig. 17 shows such a transient process, and several points are plotted on the graph, characterizing the ideal process. Such a parametric and structural change made it possible to qualitatively change the output parameters of the system with experimental data and approach the ideal parameters chosen in the article.

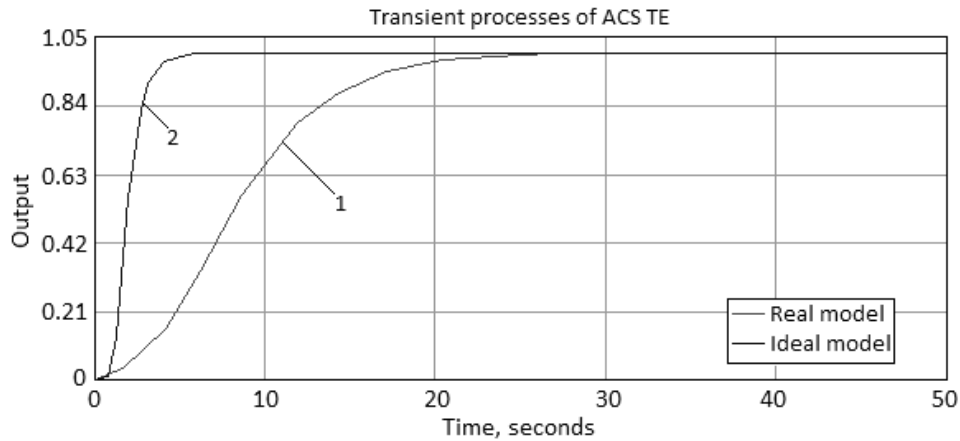


Figure 17: Transient processes of ACS TE with models and introduction of an integrator into the structure (1), ideal TE (2)

In [17, 18], the modeling of transient processes in steady-state operating modes (for nominal and emergency modes) was carried out using a multi-mode neural network controller based on a perceptron. In this paper, a similar study was carried out using the developed ACS TE based on a self-tuning neural network control system. Input data similar to [7, 8] are given in table 3.

Table 3

Input data of TV3-117 aircraft engine operating modes

TV3-117 aircraft engine operating modes	n_{TC}	T_G^*	G_T
M_1 (Nominal mode)	0.709	0.769	0.096
M_2 (Emergency mode)	1.092	1.087	0.334

Fig. 18 shows the results of modeling the automatic control system for TV3-117 gas turbine engine based on a self-adjusting neural network control system, from which it follows that the use of

the developed automatic control system increases the accuracy of modeling transient processes in the gas turbine engine control system, the graphs of which are close to the standard.

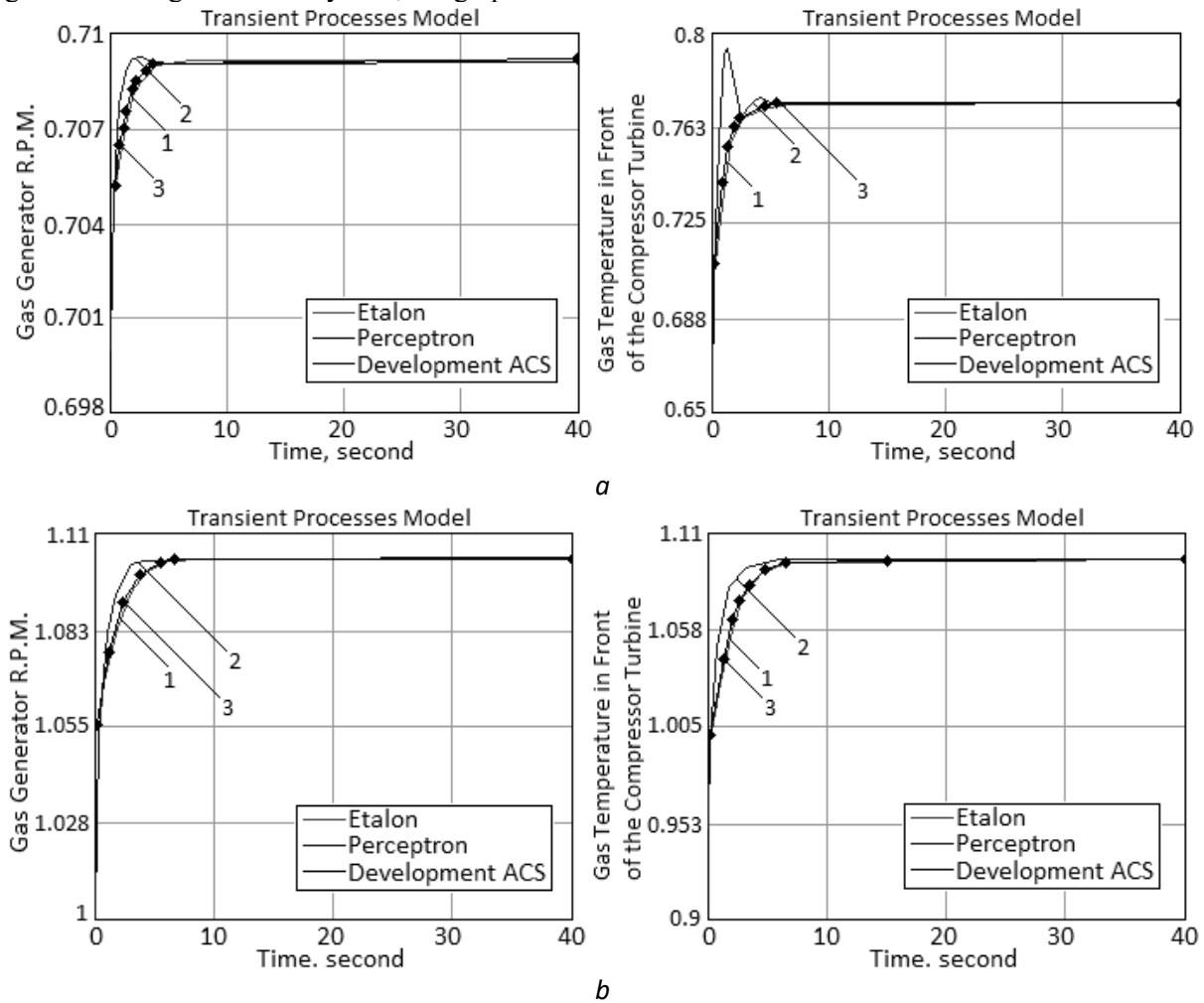


Figure 18: Results of modeling the automatic control system of TV3-117 aircraft engine: a – mode M_1 ; b – mode M_2 ; 1 – reference model; 2 – model with a multi-mode neural network controller based on perceptron; 3 – model with self-tuning neural network control system

10. Conclusions

The automatic control system of helicopters aircraft engines has been improved, in which the division of the control object into actuating mechanism and gas turbine engines makes it possible to take into account the dynamics of the executive part of the system and the engine, it becomes possible to use the mismatch between parts of the structural diagram of the automatic control system, thereby increasing the reliability and stability of the system in various modes.

The method for constructing a mathematical model of the automatic control system for helicopters aircraft engines, based on the selector of control channels according to the engine's thermogasdynamic parameters, was further developed by modifying the transfer functions, which made it possible to adapt the developed automatic control system for helicopters aircraft engines to a change in time constants and delays, namely, changing them even by two or three times does not have a significant effect on changing the quality of regulation in the system.

The self-adjusting neural network control system for multiply connected dynamic objects was further developed, the adaptation of which as a modified neural network controller of helicopters aircraft engines (by introducing an integrator into the system structure) made it possible to bring the graph of the real transient process in the engine closer to the ideal one, thereby increasing the reliability and stability of the system at various modes. The intelligent approach made it possible to

form a logical block, which qualitatively improved the output parameters of the system and made it possible to approach the ideal ones with a sufficient degree of accuracy.

11. References

- [1] S. N. Vassilyev, A. Yu. Kelina, Y. I. Kudinov, F. F. Pashchenko, Intelligent Control Systems, *Procedia Computer Science*, vol. 103 (2017) 623–628. doi: 10.1016/j.procs.2017.01.088
- [2] S. Pang, Q. Li, B. Ni, Improved nonlinear MPC for aircraft gas turbine engine based on semi-alternative optimization strategy, *Aerospace Science and Technology*, vol. 118 (2021), 106983. doi: 10.1016/j.ast.2021.106983
- [3] I. A. Krivosheev, K. E. Rozhkov, N. B. Simonov, Complex Diagnostic Index for Technical Condition Assessment for GTE, *Procedia Engineering*, vol. 206 (2017) 176–181. doi: 10.1016/j.proeng.2017.10.456
- [4] E. L. Ntantis, Diagnostic methods for an aircraft engine performance, *Journal of Engineering Science and Technology*, vol. 8, no. 4 (2015) 64–72. doi: 10.25103/jestr.084.10
- [5] S. Hosseinimaab, A. Tousi, A new approach to off-design performance analysis of gas turbine engines and its application, *Energy Conversion and Management*, vol. 243 (2021), 114411. doi: 10.1016/j.enconman.2021.114411
- [6] G. Dong, Y. Li, S. Sui, Fault detection and fuzzy tolerant control for complex stochastic multivariable nonlinear systems, *Neurocomputing*, vol. 275 (2018) 2392–2400. doi: 10.1007/s00500-019-04618-8
- [7] C. Rodriguez, J. Viola, Y. Q. Chen, Data-driven modelling for a high order multivariable thermal system and control, *IFAC-PapersOnLine*, vol. 54, issue 20 (2021) 753–758. doi: 10.1016/j.ifacol.2021.11.262
- [8] G.-Q. Zeng, X.-Q. Xie, M.-R. Chen, J. Weng, Adaptive population extremal optimization-based PID neural network for multivariable nonlinear control systems, *Swarm and Evolutionary Computation*, vol. 44 (2019) 320–334. doi: 10.1016/j.swevo.2018.04.008
- [9] R. Ortega, D. N. Gerasimov, N. E. Barabanov, V. O. Nikiforov, Adaptive control of linear multivariable systems using dynamic regressor extension and mixing estimators: Removing the high-frequency gain assumptions, *Automatica*, vol. 110 (2019) 108589. doi: 10.1016/j.automatica.2019.108589
- [10] S. H. Nagarsheth, S. N. Sharma, Relative Normalized Gain Array-based interaction indicator for non-square multivariable control systems: properties and application, *IFAC-PapersOnLine*, vol. 53, issue 2 (2020) 4032–4037.
- [11] T. Ma, Recursive filtering adaptive neural fault-tolerant control for uncertain multivariable nonlinear systems, *European Journal of Control*, vol. 59 (2021) 274–289. doi: 10.1016/j.ejcon.2020.10.003
- [12] S. Labiod, T. M. Guerra, Adaptive Fuzzy Control for Multivariable Nonlinear Systems with Indefinite Control Gain Matrix and Unknown Control Direction, *IFAC-PapersOnLine*, vol. 53, issue 20 (2020) 8019–8024. doi: 10.1016/j.ifacol.2020.12.2232
- [13] A. M. Shaheen, R. A. El-Sehiemy, M. M. Alharthi, S. S. M. Ghoneim, A. R. Ginidi, Multi-objective jellyfish search optimizer for efficient power system operation based on multi-dimensional OPF framework, *Energy*, vol. 237 (2021) 121478. doi: 10.1016/j.energy.2021.121478
- [14] Q. Zhang, J. Ding, W. Kong, Y. Liu, Q. Wang, T. Jiang, Epilepsy prediction through optimized multidimensional sample entropy and Bi-LSTM, *Biomedical Signal Processing and Control*, vol. 64 (2021) 102293. doi: 10.1016/j.bspc.2020.102293
- [15] A. Mohammadi, F. Sheikholeslam, S. Mirjalili S., Inclined planes system optimization: Theory, literature review, and state-of-the-art versions for IIR system identification, *Expert Systems with Applications*, vol. 200 (2022) 117127. doi: 10.1016/j.eswa.2022.117127
- [16] Vasiliev S., Badamshin R., Valeev S., Vasiliev V., Gvozdev V., Guzairov M., Zhernakov S., Ilyasov B., Kusimov S., Munasypov R., Raspopov E., Frid A. and Chernyahovskaya L. (2008) Intelligent systems for control and monitoring of gas turbine engines, Ufa, Ufa State Aviation Technical University, 550 p.

- [17] S. Vladov, N. Nazarenko, N. Tutova, V. Moskalyk, A. Ponomarenko, Multi-dimensional TV3-117 aircraft engine automatic control system based of the neural network regulator, Transactions of Kremenchuk Mykhailo Ostrohradskyi National University, issue 2/2020 (121) (2020) 79–84. doi: 10.30929/1995-0519.2020.2.79-84
- [18] T. Shmelova, Yu. Shmelov, S. Vladov, Concept of building intelligent control systems for aircraft, unmanned aerial vehicles and aircraft engines, 2020 IEEE 6th International Conference on Methods and Systems of Navigation and Motion Control (MSNMC), Kiev, Ukraine, October 2020 14–19. doi: 10.1109/MSNMC50359.2020.9255509
- [19] S. Vladov, A. Matusiev, I. Yakovenko, Synthesis of channels for control of helicopters aircraft engines gas generator r.p.m, Physical processes and fields of technical and biological objects: XX International scientific and technical conference, Kremenchuk, Ukraine, November, 2021, 11–13.
- [20] B. Ilyasov, V. Vasilyev, S. Valeev, Design of multi-level intelligent control systems for complex technical objects on the basis of theoretical-information approach, Acta Polytechnica Hungarica, vol. 17, no 8 (2020) 137–150. doi: 10.12700/APH.17.8.2020.8.10
- [21] Y. Shen, K. Khorasani, Hybrid multi-mode machine learning-based fault diagnosis strategies with application to aircraft gas turbine engines, Neural Networks, vol. 130 (2020) 126–142. doi: 10.1016/j.neunet.2020.07.001
- [22] E. Denisova, M. Chernikova, Automatic control system for gas turbine engines with the insertion of mathematical models to the control, Fundamental research, no. 9-2 (2016) 243–248.
- [23] V. Petunin, Synthesis of automatic control systems for gas turbine engines with channel selector, Bulletin of USATU, vol. 11, no 1 (28) (2008) 3–10.
- [24] V. Petunin, A. Frid, Method for constructing adaptive logical-dynamic automatic control systems with selectors, Journal of Instrument Engineering, vol. 54, no 5 (2011) 49–56.
- [25] V. Petunin, Mathematical models of multiconnected automatic control systems with channel selectors, Bulletin of USATU, vol. 15, no 2 (42) (2011) 52–58.
- [26] V. Petunin, Synthesis of automatic control systems for aircraft with automata limiting parameters, Journal of Instrument Engineering, vol. 53, no 10 (2010) 18–24.
- [27] V. Petunin, Synthesis of an adaptive gas-turbine engine gas temperature controller, Reshetnev Readings: Proceedings of the XIII International Scientific Conference, Krasnoyarsk, Russia, November, 2009 158–159.
- [28] I. Elizarov, M. Soludanov, Self-adjustable neural network control system of multilinked dynamic object, Information processes and management, no 1 (2006) 30–44.
- [29] I. Siddikov, Z. Iskandarov, Algorithm for the synthesis of a self-tuning neural network control system for multicomponent dynamic objects, 11th World Conference “Intelligent System for Industrial Automation” (WCIS-2020), Tashkent, Uzbekistan, October, 2021 435–441. doi: 10.1007/978-3-030-68004-6_57
- [30] Y. Xue, Y. Wang, J. Liang, A self-adaptive gradient descent search algorithm for fully-connected neural networks, Neurocomputing, vol. 478 (2022) 70–80. doi: 10.1016/j.neucom.2022.01.001
- [31] L. Demidova, D. Marchev, Application of recurrent networks in the classification problem of failures in the complex technical systems within the framework of proactive maintenance, Vestnik of Ryazan State Radio Engineering University, no 69 (2019) 135–148. doi: 10.21667/1995-4565-2019-69-135-148
- [32] B. Ilyasov, G. Saitova, I. Sabitov, Control of multi-connected systems based on logic controllers, Bulletin of USATU, vol. 18, no 2 (63) (2014) 98–102.

Feature Engineering and Missing Data Imputation Method of Medical Data Analysis

Nataliya Shakhovska¹, Nataliia Melnykova¹

¹ Lviv Polytechnic National University, S.Bandera str,12, Lviv, 79013, Ukraine

Abstract

This work provides an alternative way to preprocessing procedure for consolidated data. Two methods are proposed. The first one is used for feature selection based on ensemble of machine learning algorithms. And the second one organizes missing data imputation based on combination of functional dependencies and associative rules. Ensemble methods for processing multimodal data based on a hierarchical classifier, a set of weak classifiers and a number of methods for selecting important characteristics with a much higher value of accuracy on unbalanced data sets compared to existing machine learning methods are developed. The methods are validated on medical dataset. The percentage of recovery data is on 1.2% comparing with associative rules. The proposed missing data imputation method creates additional data values operating a based domain and functional dependencies and includes these values to available training data. The correctness of the filled-in values is proved on the predictor built on the original dataset. The proposed PPD method conducts 12% better than RF and EM models for 30% missing data.

Keywords

feature selection, missing data, machine learning, ensemble, data preprocessing.

1. Introduction

Methods for detecting dependencies in data have existed for more than a century. The first of them emerged as methods of mathematical and statistical analysis (correlation and factor analysis). The computerization of data analysis processes has dramatically increased the amount of data analyzed, the quality of analysis and the scope of methods of knowledge extraction.

The information boom has led to a thousand-fold growth in data collected in different domains. Such areas include computer technology, economics, sociology, medicine, astronomy, etc. The increase in the amount of information collected continues to grow exponentially. For example, based on a Digital Universe Study commissioned by EMC, the total global data in 2005 was 130 exabytes, up to 2227 EB by 2015, and increased again last year to 7 ZB (zettabytes). The forecast made by the same study shows that by 2025 the amount of digital data will increase to 10.9 ZB. The size of separate databases is growing just as fast and has overcome the petabyte barrier. Most of the collected data is not currently analyzed or is only a simple analysis. This is the most important for medical data, primarily when hospitals do not support medical standards for data exchanging (HL7, for example).

The main problems that arise in data processing are

- the absence of analytical methods appropriate for use in several subject areas,
- the demand for significant human resources to support the data investigation process,
- the high computational complexity of existing investigation algorithms and the fast growth of composed data.

To the constant increase in research time, even with frequent updates of server hardware and the need to work with distributed databases, the capacities of which most existing data analysis methods

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAILnataliya.b.shakhovska@lpnu.ua (N. Shakhovska); melnykovanatalia@gmail.com (N. Melnykova)

ORCID: 0000-0002-6875-8534 (N. Shakhovska); 0000-0002-2114-3436 (N. Melnykova)



© 2020 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

are not used inefficiently. Thus, there is a problem of developing an effective strategy of data analysis that can be applied to dispersed databases of various domains.

Data analysis methods are effectively used on clear and previously processed data. That is why the paper aimed to analyses and develop methods for data preprocessing. Particularly, we have taken into account feature selection (feature engineering) and missing data imputation.

The main contribution of the paper is as follows:

1. A new hybrid ensemble feature selection model for a machine learning-based post-COVID-19 prediction system is proposed as an automatic feature cut-off rank identifier.
2. A method based on probabilistic dependence is developed and tested. The percentage of recovery data is on 1.2% comparing with associative rules.
3. Development of a method for finding dependencies in large data sets with gaps and uncertainties based on an ensemble of clustering methods and auto-associative dependencies. Boruta, Decision Tree and Random Forest are used to select an object. The importance of variables is different for different methods (logistic regression, Support vector machine, Naive Bayes, XGBoost, Random Forest, neural network, decision tree). This means that the relationship between the parameters is supported only for part of the data set. That is why we propose to find the dependency for separate clusters and use this dependence for classification.

2. Literature Review

Selecting the appropriated features can be a more significant task than lessening computation time or enhancing classification or predictive accuracy. For example, in medicine [1], finding the optimal set of optimal features for the classification or predictive problem can help develop a diagnostic test.

Selecting important features (for example, determining genes appropriate for a specific type of cancer) can help decipher the mechanisms underlying the welfare problem for research. A complete enumeration of features can implement the selection method, that is, having checked all possible sets, selecting those signs for which the error is minimal. This method is simple to implement, but it is entirely ineffective on big data. Therefore, in this case, other algorithms are most often used.

The three primary classes of feature extract algorithms – filters, wrappers, and built-in algorithms are used [2].

Filters are based on some metrics that are independent of the classification method. For example, the correlation of features with the target vector and information content criteria. They are applied before classification. The most significant benefit of filtering is that it can be used as preprocessing to reduce space dimensionality and overcome overfitting. Filtering methods are generally fast. Filters are used to choose features in clustering or to build an initial approximation [3]. Unfortunately, such methods are not designed to detect complex relationships between elements and, as a rule, are not sensitive enough to specify all dependences in the data.

Embedded algorithms organize feature extraction during the classifier training approach, and it is they explicitly optimize the set of features used to achieve better accuracy [4]. The benefits of built-in algorithms are that, as a rule, they discover solutions quickly, decrease the possibility of retraining data, eliminate the need to separate data into training and test subsamples. Nevertheless, these algorithms are not ubiquitous.

Wrappers rely on feature significance information from several classifications or regression models and can thus find deeper patterns in dataset than filters. Wrappers can be built on any classifier that defines the degree of significance of the features [5].

Imputation is a procedure for assessing unknown or missing values based on available data, which allows you to form a complete set of data with some plausible estimates.

Methods of imputation are divided into non-model and model-based approaches. There are approaches based on single and multiple imputation methods in terms of the quantity of values received from imputation methods [6]. One-time filling algorithms provide a single complete set of data, where a deal replaces each space. The benefits of this method are the use of methods of comprehensive data analysis in the subsequent stages of processing. Replenishment algorithms form several complete data sets analyzed separately and later combined according to specific rules. This minimizes standard errors in the following steps by processing comprehensive datasets. Nevertheless,

mentioned methods require multiple resources to create more data sets, spend more time performing the analysis, and have more memory to store the results [7].

The method based on mean substitution solves the problem of incomplete data by replacing each missing variable with an average value. There are the following types of substitutions: median value, mean value for subgroup [8], value with the highest frequency, and replacement with minimum / maximum value. This method can lead to undesirable results [8], such as variance change, negative correlation shift, and misrepresentation of the population.

Hot-deck imputation approach replaces each gap with a random value taken from an existing dataset [9]. Its significant disadvantage is the warping of correlations and covariates.

Cold-deck imputation (CD) approach implements the replacement of each gap by some constant value from an external source [9]. The specific case of CD is zero replacement. It has the same disadvantages as hot-deck imputation.

The regression model applies to replace data gaps with expected values emanated from a regression equation constructed from a complete dataset. Disadvantages of regression completion include the need to accurately define regression models, exaggerate correlation and covariance, the likelihood of moving to predicted values outside the logical sequence, and the need for extensive quantities of data to obtain consistent estimates.

The association rules (AR) mining method uses the constructed associative rules for the data imputation [10, 11]. However, the temporal complexity of this method needs to be improved [12].

Diabetes increases the likelihood of severe COVID19. New clinical data and investigations show that this might work in the opposite direction: scientists are recording new cases where COVID-19 has sharply provoked type 1 diabetes in humans. The World Health Organization views diabetes as one of the existing diseases, on a par with respectable age, making someone more vulnerable to severe COVID 19 infection. Cellular immunity is essential for protecting against viral diseases, and its effectiveness decreases with age. In particular, a decrease in T-cell receptors explains the significant increase in mortality of COVID19 with age.

Therefore, another important factor is the search for possible relationships between biomarkers of aging and COVID resistance.

Existing research and the data sets collected for this use standard machine learning methods, while demonstrating not very high accuracy of prediction. Thus, in the paper [13], there are used feature selection, XGBoost and decision tree to determine COVID biomarkers, F1-score does not rise above 0.7. In the paper [14], markers of CH4 and CH8 immunodeficiency and their association with coronavirus infections were analyzed using statistical models, including the Cox model. Accordingly, it is impossible to prevent negative situations. In paper [15], there is used the empirical mode decomposition (EEMD) and the artificial neural network (ANN) to predict the COVID-19 epidemic. Thus, in order to prolong the active period of life, it is necessary to track the dynamics of changes in molecular and biochemical markers, anthropometric indicators, behavioral factors, environmental parameters and habitat, and so on. As a result, it is necessary to use a big data-based approach to collect information from disparate datasets, process them, and further analyze them. At the same time, it is necessary to analyze small data samples, which will include time multimodal series of changes in human parameters.

The analysis of literature sources showed the lack of a comprehensive approach to solving the problem of prolonging the active period of life and preventing exacerbation of chronic diseases. At the same time, as we see in the case of COVID, chronic diseases (diabetes and obesity) can not only reduce the ability to work, but also increase the likelihood of severe course of other diseases.

3. Materials and Methods

Biomarkers of aging are used to predict possible changes in the body that lead to disability due to functional age-related changes. Biomarkers of aging are markers that can predict the functional capacity of an organism at a certain age better than chronological age. Since the immune system is a change the language here for English version

leading factor in aging, the main impact of which is realized through increased inflammation and reduced effectiveness of cellular immunity, it becomes clear the need to involve relevant markers to develop interventions to increase the duration of healthy longevity.

Thus, control of chronic age-related diseases (diabetes and obesity) and biological markers can be used to predict functional changes in the body, and analysis of other personal indicators will determine how to reduce the negative effects of such changes by extending the period of active longevity.

Sociological research also shows that people in certain regions remain active for a long time and there are far fewer people who are obese. Therefore, it is also advisable to analyze the parameters of the environment and habitat and its impact on the parameters of the organism.

The baseline of the hybrid ensemble feature selection model looks like the following:

- Several selectors using,
- Aggregation of the results.

Several wrapper algorithms will be used in the preprocessing stage for the first stage.

The correlation matrix shows the numerical value of the correlation coefficient for all possible combinations of variables. It is mainly used to find out the relationship between more than two variables. The decision tree returns the feature weight as the criterion for evaluating features. It allows building a ranked list of selected features using different measures. Classification and Regression Trees (CART) was used for feature selection with Gini-index as a measure in our case.

Random Forest is an ensemble of numerous training-sensitive algorithms (decision trees). Mentioned approach has a slight compensation. The bias of the training method is the deviation of the average response of the trained algorithm from the reaction of the ideal algorithm. Each of these classifiers is built on a random subset of entities and a random subset of characters.

Boruta is a heuristic algorithm for choosing important features based on Random Forest approach [16]. Features with the Z-measure less than the maximum Z-measure among the added features are removed at each iteration. To calculate the Z-measure, we need to calculate its importance, obtained using the built-in algorithm in Random Forest, and divide it by the standard deviation of the feature importance. Added features are obtained as follows: the characteristics available in the selection are copied, and then every new attribute is filled by mixing its values. This procedure is repeated several times to get statistically significant results, and variables are generated independently at each iteration.

The Jaccard index [17] measures the similarity of the feature subsets chosen by separated feature selectors (each selector is organized as a separated iteration):

$$(S_1, \dots, S_{1n}) = \frac{|S_1 \cap \dots \cap S_n|}{|S_1 \cup \dots \cup S_n|}, \quad (1)$$

where S_i is the subset of features at the i -th iteration, for $i=1, \dots, n$. The value of the Jaccard index varies from 0 to 1, where 1 implies the absolute similarity of subsets.

The schema of the hybrid ensemble feature selection model is given in Fig. 1.

Next, the method of missing data imputation is developed. This method is based on classical functional dependencies in relational databases and association rules from non-relational databases. It consists of two parts:

- Probabilistic Production Dependencies mining;
- the Probabilistic Production Dependencies usage for missing data imputation.

Investigating extensive data requires specifying attribute clusters that form functional dependencies (FD) [18]. However, datasets are extensively standard in the real world, with essential dependencies defined only on a subset of crucial attribute group values. Moreover, the reliance can appear not only for tuples in relational data sources but also between subsets of values in different tuples. We will name them Probabilistic Production Dependency (PPD).

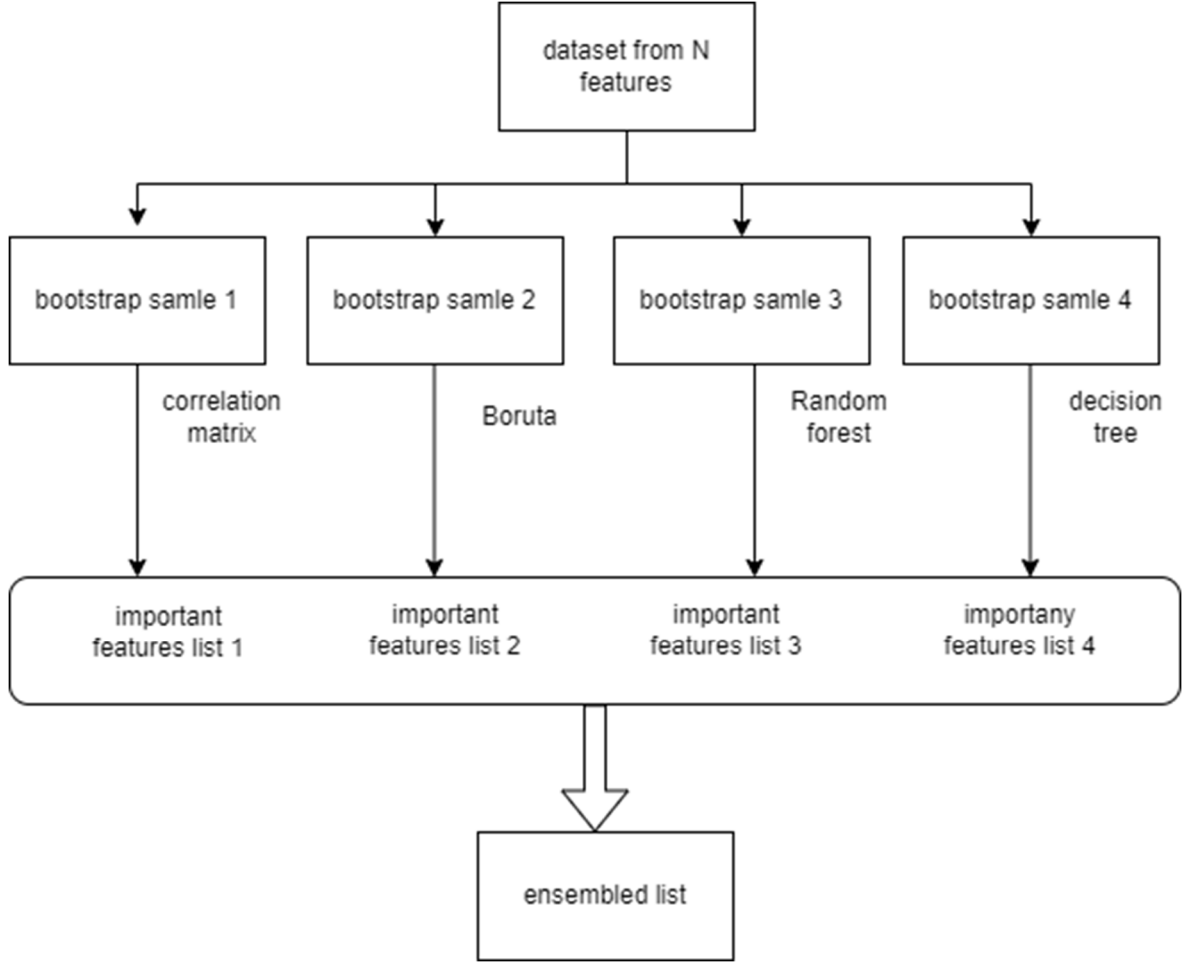


Figure 1: The hybrid ensemble feature selection model.

Probabilistic Production Dependency is a dependence similar to associative rule in the primary ratio selection that is proper for many entities. The significance threshold should be specified expertly or based on estimation of the probability of erroneous selection of this dependence. The main distinction between associative rules and PPD is that PPD will be generated from existing functional dependencies (FD) in the dataset [9].

$$F_l : K = \{a_i\}, a_i \in A, D = \{a_j\}, \quad (2)$$

$$a_j \in A, : P(k \in K \rightarrow d \in D) = p'$$

where k and d are the tuples of groups of attributes K and D , respectively.

The gaps and missing data presented among the values of the attribute Y of the relation r are classified using PPD. The following algorithm for PPD mining is proposed.

Algorithm 1. PPD mining algorithm

1. Entities with the same X-values will be grouped;
2. To choose attributes from FD with same X and add them to Y;
3. To calculate the *Support* and *Confidence*, *Imputation* of the tuple selected in step 2);
4. To identify the tuples with the highest value of *Confidence*;
5. To add $X \rightarrow Y$ to *PPDset*.

To fill in missing data, the PPD should be built.

Next, the novel algorithm for missing data imputation is developed.

Algorithm 2. Data imputation algorithm

```

Completeness=0
While Completeness/100< Imputation
Arrange all attributes from PPDset by Confidence level
For each group:
    If percentage of non-empty Y-value is higher or equal to Support
        fill in empty values using PPDset
        Completeness++
    Else
        Merge PPD using Armstrong rules

```

Next, a hierarchical classifier as a two-stage data prediction algorithm is developed. The first stage is clustering; the next step is to build a classification model for each separate cluster. K-means [19] together with the random forest do not dominate other cluster models. The hierarchical classifier is constructed as follows:

- the appropriate number of clusters was found using gap statistics;
- the density of distribution is calculated;
- XGboost and Random Forest are used for each cluster separately;
- hard voting for the results. Based on it, the class with the highest number of votes will be selected. If the voices are the same, the result of the classifier with the minimum value of depth will be selected.

4. Experiments and results

To validate the proposed methods, dataset with medical data is used.

Dataset consists of 35 features and 122 instances collected from Lviv regional rehabilitation center for post-COVID patients with short- and long-term (more than 20 days) treatment and rehabilitation.

The personal data were removed from the dataset and replaced with unique random identifiers. The next feature, sex, is processed using one-hot encoding technics and in the final dataset is given in two components – female and male. Features like age, weight, height, BMI, CAT, pulse, the function of external respiration are taken as physiological parameters measured before inpatient treatment. The rest of the features were immune-based biomarkers as described below.

Zero cells (0-lymphocytes) do not carry T- and B-cells markers. Zero cells make 10–20 % of the total lymphocytes in human peripheral blood. Some researchers consider them immature or overripe T- or B-lymphocytes because they have a small number of antigens common to B- and T-cells. Zero cells include K-cells and NK-cells.

CD3+ is a surface marker specific to all T-lymphocyte subpopulation cells. By function, it belongs to the family of proteins that form a complex of membrane signaling associated with the T-cell receptor. Mature T-lymphocytes are "responsible" for cellular immune reactions and conduct immunological monitoring of antigenic homeostasis in the body.

CD4+ is a characteristic of helper T-cells; also represented on monocytes, macrophages, dendritic cells. It binds to class II MHC molecules expressed on antigen-presenting cells, facilitating the recognition of peptide antigens. Helper T-lymphocytes (CD4+) are helpers (inducers) of the immune response, cells that regulate the strength of the body's immune response to a foreign antigen, as well as control the stability of the body's internal environment (antigenic homeostasis) and cause increased antibody synthesis [20].

CD8+ is a characteristic of suppressor and cytotoxic T-cells, NK-cells, mostly thymocytes. It is a T-cell activation receptor that facilitates the recognition of cell-bound class I MHC antigens.

CD16+ natural killers are part of innate immunity; they are involved in early response against viral infections and intracellular bacteria. Compared with cells of specific immunity (T- and B-lymphocytes), they have the advantage that they do not require long-term activation. Besides, NK-

cells complement the action of T-cytotoxic cells can also regulate the immune response by producing various cytokines, including interferon- γ . They are the primary cells of antitumor protection. Their role is vital in manifesting cellular immunity in viral, protozoan, fungal and bacterial diseases caused by intracellular parasites. Their action is enhanced by interferon. The functions performed by natural killers can be divided into two main types: the production of cytokines that regulate the work of other cells of the immune system and the direct destruction of damaged cells.

Mature B lymphocytes express CD22+ markers. B-lymphocytes are responsible for the humoral adaptive immune response, primarily at removing extracellular infectious agents. After binding to a specific antigen, B-lymphocytes, in cooperation with T-lymphocytes and T-helpers proliferate, differentiate into plasma cells that secrete antibodies/immunoglobulins and memory cells. Defects of humoral immunity associated with the B-cells are sporadic, so a common hypoinmunoglobulinemia is mainly caused by other reasons.

CD4/CD8 immunoregulatory index reflects the ratio of CD4+ cells (T-helpers) to CD8+ cells (T-cytotoxic cells). It is a relative indicator that has an indicative value. Its small increase or decrease has no independent diagnostic value. Changes in the index significance the clinician to focus on the reasons for the deviation of this index. The immunoregulatory index is assessed relative to the phase of the immune response. In the period of exacerbation and remission of clinical manifestations, the immunoregulatory index reaches high values due to the high percentage of T-helpers (CD4+ T-cells). During the recovery period, the indicator's value decreases due to the increase in the level of CD8+ T-cells (killers). Violation of this pattern indicates the inadequacy of the immune response and the possibility of chronic infection due to incomplete removal of the pathogen [21].

We implemented our approach in Rstudio. The essential packages we used were caret, rpart, Metrics, Boruta, Random forest, rules and ggplot2 for visualization. To generate PPD, the minimal support threshold equal to 0.0001 in the a priori algorithm is chosen. In addition, all rules with confidence below 0.001 are filtered out.

First, the predictive accuracy for whole dataset and selected features was analyzed (Table 1).

Table 1

The comparison of ML models prediction results using the whole set of features and selected features subset

Model	For the whole dataset			For selected features		
	MSE	MAE	R2	MSE	MAE	R2
k-NN	7.029	2.155	-0.162	5.781	1.825	0.044
SVM	5.753	1.828	0.049	5.521	1.812	0.038
SGD	16.007	3.280	-1.647	12.328	2.603	-1.039
Linear	17.826	3.571	-1.948	14.836	2.760	-1.453
Regression						
MLP NN	5.717	1.777	0.055	5.469	1.034	0.034

The much higher accuracy is obtained with SVM (Support vector machine, polynomial kernel) and artificial neural network (one hidden layer with 12 neurons in it, sigmoid activation function) for selected features [22]. The used measures are the following [23]:

- mean squared error MSE,
- mean absolute error MAE,
- R2.

Next, missing data imputation method is used. The developed method was compared with the existing ones: associative rules (AR), random forest (RF), support vector machine (SVM), multilayered perceptron (MLP), expectation-maximization (EM) and k-nearest neighbor (KNN) (Fig. 2). The recovery error is presented using normalized root-mean-square error (NRMSE).

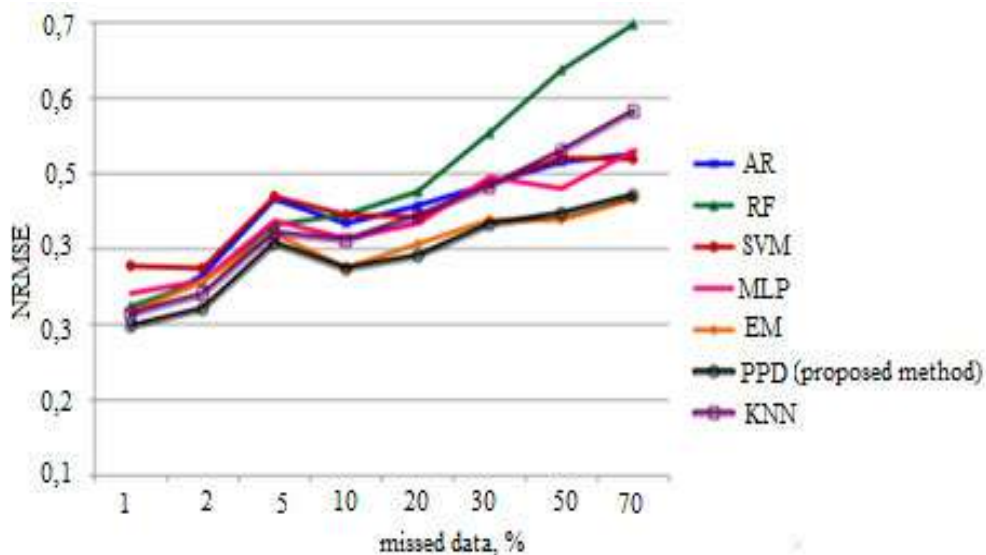


Figure 2: NRMSE recovery error.

The proposed missing data imputation method creates additional data values operating a based domain and functional dependencies and includes these values to available training data. The correctness of the filled-in values is proved on the predictor built on the original dataset. The proposed PPD method conducts 12% better than RF and EM models for 30% missing data.

5. Conclusion

Current trends in the development of information technology and databases are analyzed. As a result, unsolved problems in the field of dependency search in large databases of several subject areas, in particular, in medicine, were revealed. The analysis of existing methods and means of detection of dependences in data is carried out. This made it possible to identify a new subclass of dependencies - Probabilistic Production dependencies. A method for deriving PPD in relational databases has been developed.

A new hybrid ensemble feature selection model for a machine learning-based post-COVID prediction system is proposed as an automatic feature cut-off rank identifier. Ensemble methods for processing of multimodal data based on a hierarchical classifier, a set of weak classifiers and a number of methods for selecting important characteristics with a much higher value of accuracy on unbalanced data sets compared to existing machine learning methods are developed.

The associative rules are found together with weak predictors usage to improve the classification quality.

The proposed missing data imputation method creates additional data values operating a based domain and functional dependencies and includes these values to available training data.

The correctness of the filled-in values is proved on the predictor built on the original dataset. The proposed PPD method conducts 12% better than RF and EM models for 30% missing data. The EM method looks the best for more additional missing data (the range of about 40%-50% missing data), and the PPD has equivalent results with the SVM (support vector machine).

Comparison of approaches for investigating and modeling the statistical processes by qualitative criteria verified the proposed method has the subsequent benefits:

- retaining the characteristics of resistance to errors in the data;
- allowing the parallel implementation in distributed databases;
- automating and performing the analysis of various data types.

6. References

- [1] M.B. Kursu, W.R. Rudnicki, The all relevant feature selection using random forest, arXiv preprint arXiv:1106.5112, 2011.
- [2] G. Chandrashekar, F. Sahin, A survey on feature selection methods. *Computers Electrical Engineering*, 40(1), 16-28 (2014).
- [3] A. Bommert, X. Sun, B. Bischl, J. Rahnenführer, M. Lang, Benchmark for filter methods for feature selection in high-dimensional classification data. *Computational Statistics Data Analysis*, 143, 106839 (2020).
- [4] B. Venkatesh, J. Anuradha, A review of feature selection and its methods. *Cybernetics and Information Technologies*, 19(1), 3-26 (2019).
- [5] L. N. Sanchez-Pinto, L. R. Venable, J. Fahrenbach, M. M. Churpek, Comparison of variable selection methods for clinical predictive modeling. *International journal of medical informatics*, 116, 10-17 (2018).
- [6] P. Hayati Rezvan, K. J. Lee, J. A. Simpson, The Rise of Multiple Imputation: A Review of the Reporting and Implementation of the Method in Medical Research. *BMC Med Res Methodol*, , 15, 30 (2015). <https://doi.org/10.1186/s12874-015-0022-1>
- [7] N. Ahmat Zainuri, A. A. Jemain, N. A Muda, Comparison of Various Imputation Methods for Missing Values in Air Quality Data. *JSM*, 44, 449–456 (2015). <https://doi.org/10.17576/jsm-2015-4403-17>
- [8] C. A. Leke, T. Marwala, Introduction to Missing Data Estimation. *Deep Learning and Missing Data in Engineering Systems*, 1-20 (2019). <https://doi.org/10.1117/12.2053057>
- [9] C. Wang, N. Shakhovska, A. Sachenko, M. Komar, A New Approach for Missing Data Imputation in Big Data Interface. *Information Technology and Control*, 49(4), 541-555 (2020).
- [10] M. Azmi, G. C. Runger, A. Berrado, Interpretable regularized class association rules algorithm for classification in a categorical data space. *Information Sciences*, 483, 313-331 (2019).
- [11] F. Thabtah, P. Cowling, Y. Peng, MCAR: multi-class classification based on association rule. In *The 3rd ACS/IEEE International Conference on Computer Systems and Applications*, (2005)/
- [12] K. Mittal, G. Aggarwal, P. Mahajan, A comparative study of association rule mining techniques and predictive mining approaches for association classification. *International Journal of Advanced Research in Computer Science*, 8(9) (2017).
- [13] L. Yan, H. T. Zhang, J. Goncalves, Y. Xiao, M. Wang, Y. Guo, Y. Yuan, An interpretable mortality prediction model for COVID-19 patients, *Nature machine intelligence*, 2 (5), 283-288 (2020).
- [14] A. Trickey, M. T. May, P. Schommers, J. Tate, S. M. Ingle, J. L. Guest, J. A. Sterne, CD4: CD8 ratio and CD8 count as prognostic markers for mortality in human immunodeficiency virus – infected patients on antiretroviral therapy: the Antiretroviral Therapy Cohort Collaboration (ART-CC), *Clinical Infectious Diseases*, 65 (6), 959-966 (2017).
- [15] N. Hasan, A Methodological Approach for Predicting COVID-19 Epidemic Using EEMD-ANN Hybrid Model. *Internet of Things* 11, 100228 (2020). <https://doi.org/10.1016/j.iot.2020.100228>
- [16] L. N. Sanchez-Pinto, L. R. Venable, J. Fahrenbach, M. M. Churpek, Comparison of variable selection methods for clinical predictive modeling. *International journal of medical informatics*, 116, 10-17 (2018).
- [17] A., Dariush, A. Bastanfard, An objective method to evaluate exemplar-based inpainted images quality using Jaccard index, *Multimedia Tools and Applications* 80.17, 26199-26212 (2021).
- [18] A., Federico, et al., Association rules extraction for the identification of functional dependencies in complex technical infrastructures, *Reliability Engineering System Safety* 209 (2021).
- [19] I.-D. Borlea, et al., A unified form of fuzzy C-means and K-means algorithms and its partitional implementation, *Knowledge-Based Systems* 214 (2021): 106731.
- [20] A.M. Miggelbrink, et al., CD4 T-cell exhaustion: Does it exist and what are its roles in cancer?, *Clinical Cancer Research* 27.21, 5742-5752 (2021).
- [21] L. I. U. Huan-xia, et al., Analysis of CD4/CD8 ratio in HIV-infected patients who accepted initial antiretroviral therapy for 48 weeks, *China Tropical Medicine* 21.3 (2021).

- [22] D.A. Otchere, et al., Application of supervised machine learning paradigms in the prediction of petroleum reservoir properties: Comparative analysis of ANN and SVM models, *Journal of Petroleum Science and Engineering* 200 (2021): 108182.
- [23] D. Chicco, M. J. Warrens, G. Jurman, The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation. *PeerJ Computer Science*, 7, e623 (2021).

Model for Finding Frequent Sets in FP-growth for Multimodal Data

Nataliya Boyko, Oleksandr Tkachyk

Lviv Polytechnic National University, Profesorska Street 1, Lviv, 79013, Ukraine

Abstract

The article considers the process of generating candidates from a structured database and discusses the basic concepts and metrics (support, confidence, lift and conviction). The principle of operation of the FP-Growth algorithm is described in detail, both its mathematical part and the software part on the example of the Python language. Two primary levels of scientific knowledge are considered for research: empirical and theoretical. The advantages and disadvantages of the Apriori algorithm are presented. These shortcomings can be overcome with the FP-Growth algorithm. This algorithm is an improved method of Apriori. Frequent patterns are formed without the need to generate candidates. The study considers the main stages of the alternative algorithm FP-Growth search for associative rules. The process of transaction arrears in the database is considered in detail. The method of transforming multimodal data into a tree structure is thoroughly discussed. The example assumes a detailed iteration process for each transaction. The procedure for generating candidates is analyzed in such a way as to reduce the number of passes for a given data set. The paper examines the process of compressing transactions into a simple structure. The efficient and complete output of partial data sets is provided.

Keywords ¹

Machine Learning, Artificial Intelligence, Frequent Pattern-Growth Algorithm, Frequent Pattern-Growth Algorithm Tree, Apriori, Associative Rule, Associations Rules Learning, Netflix, Amazon.

1. Introduction

There is no doubt that in the last few years, machine learning (ML) / artificial intelligence (AI) has been gaining in popularity. Today, the use of machine learning in processing multimodal data is a multi-stage complex process that allows you to perform the necessary tasks for forecasting. Some of the most common ML algorithms are Netflix's algorithms for creating movie offers based on movies you've watched in the past [2, 12]. Another example is Amazon and its algorithms, which recommend books based on books you've previously purchased [3, 6, 11].

This work is based on one of these algorithms, the FP-Growth algorithm. FP-Growth (Frequent Pattern Growth) is a reasonably young algorithm. They were first described in 2000. FP-Growth offers a radical approach - to abandon the generation of candidates (generation of candidates is the basis of the Apriori algorithm) [1, 2, 9]. Theoretically, this approach will further increase the algorithm's speed and use even less memory. This is achieved by storing the prefix tree in memory, not from combinations of candidates but the transactions themselves [8, 16, 21].

There are several research methods. It is accepted to allocate two basic levels of scientific knowledge: empirical and theoretical. This division is because the subject can acquire knowledge experimentally (empirically) and through complex logical operations, i.e. theoretically [1, 3, 8].

The practical level of knowledge includes [13, 17]:

¹ The Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022), May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: nataliya.i.boyko@lpnu.ua (N. Boyko); oleksandr.a.tkachyk@lpnu.ua (O. Tkachyk)
ORCID: 0000-0002-6962-9363 (N. Boyko); 0000-0002-0728-4208 (O. Tkachyk)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org) Proceedings

- Observation of phenomena.
- Accumulation and selection of facts.
- Establishing ties between them.

The practical level is collecting data (facts) about social and natural objects. The object under study is reflected by external connections and manifestations at the functional level.

The theoretical level of cognition is associated with the predominance of mental activity, with the understanding of empirical material, it's processing.

The theoretical level reveals [2, 9]:

- Internal structure and patterns of development of systems and phenomena.
- Their interaction and conditionality.

2. Materials and methods

FP-Growth is an improvement on the Apriori method.

Disadvantages of the Apriori algorithm [11, 20]:

- To use Apriori, you need to generate sets of candidates. These elements can be significant if the set of elements in the database is also significant.
- Apriori requires multiple database scans to verify support for each generated set of items, resulting in high costs.

Unlike the Apriori algorithm, the alternative algorithm search algorithm avoids the multi-stage process of generating candidates while reducing the number of movements in the multimodal data set. This algorithm is an improved method of Apriori. Frequent regularity is formed without the need to generate candidates. The FP-Growth algorithm represents a database in a regular pattern tree or FPG tree [21, 13].

This tree structure will keep the link between the set of elements. A single frequent element fragments the database. This fragmented part is called a "pattern fragment". Sets of elements of these fragmentary patterns are analyzed. Thus, the search for frequent sets of items is relatively reduced with this method.

2.1. Analysis of the existing informational system

A Frequent Pattern-Growth tree is a tree-like structure made from initial sets of database elements. The purpose of the FPG tree is to remove the most common pattern. Each node of the FPG tree is an element of a set of elements [9, 14, 20].

To explain the tree structure of the alternative algorithm, the root node must be denoted by zero, and the leaves must acquire values from the data sets. The association of nodes with lower nodes, which are sets of elements, with other sets of elements is preserved during the formation of the tree.

FP-Growth Algorithm steps [7, 17]:

1) In the first stage, an FPG-tree is built that shows the data that is most common in the database. It is similar to the first step of the algorithm Apriori. FPG-tree is built based on ordered data sets, which are written in descending order of their support values. There is a rule for constructing an FPG tree. It should be formulated as follows: if a node of the same name occurs in a tree, a new node is not created, and the index of the corresponding node is increased by 1. Otherwise, a new node with index one is created.

2) In the second stage, the elements that are often found in the multimodal data set are removed from the tree.

3) In the third step, another transaction check is performed in the database. In this step, the components of the same name are removed from the multimodal data set with repeated force. The set of elements with the maximum number starts at the top, the next set of elements with the smaller number below. This means that a tree branch comprises sets of transaction elements in descending order.

4) In the fourth stage, any element from the original data set is selected, and all paths leading to this element are found. Then count the number of encounters of a given element on each way. Then the element itself (set suffix) is removed from the paths leading to it. As a result, the coincidence of

elements in the prefixes of the ways is calculated and recorded in descending order. Thus a new set of elements of groups is formed. On its basis, a new conditional tree is built, which is associated with one object. All nodes whose support is equal to or greater than the minimum value are present on this tree. Node indices are summed if there is an element with a frequency greater than 2. starting from the root, the paths to the nodes whose support is greater than or equal to the specified.

5) The paths from the root to the nodes are fixed in the next stage. the item that was deleted in the previous step is returned, and the final support values are calculated. From it is the path along the tree FPG. This path is called the conditional template base.

A conditional template database is a database that consists of configuration loops in the FPG tree that meet the lowest node (suffix).

6) A conditional tree FPG is constructed, which is formed by counting sets of elements on the path. The FPG tree considers the sets of elements corresponding to the threshold support.

7) Frequent templates are generated from the FPG tree.

Advantages of the FP-Growth algorithm [8, 19]:

- An alternative algorithm search algorithm does not require scanning each transaction on all iterations as it does in executing the algorithm Apriori. It must be performed twice.
- During the execution of the alternative algorithm, the same name elements are removed, and they are not combined as in the classical algorithm Apriori. It is this process that makes it work faster.
- In the process of building a FP-tree, it is reduced, and accordingly, the number of elements is smaller, which allows you to optimize the work of the database and optimize its dimensionality.
- The effectiveness of the alternative algorithm for finding alternative rules is due to the removal of long and short partial chains.

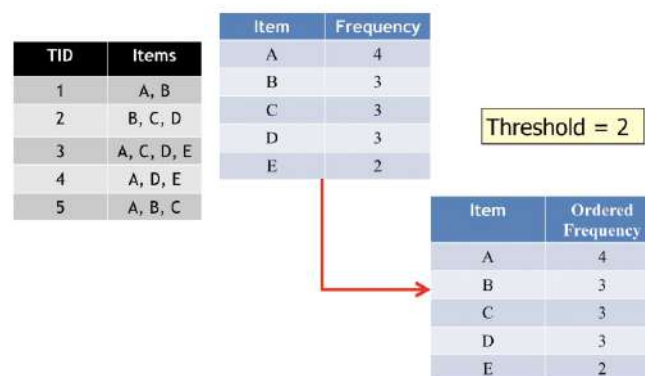
Disadvantages of the FP-Growth algorithm [5, 10, 16]:

- The process of building a FP-tree when working with an alternative algorithm for finding associative rules is cumbersome and complex, rather than the algorithm Apriori.
- In performing each stage of the alternative algorithm for finding associative rules, many resources are used because, algorithmically, it is more complex than Apriori.
- When building a tree using an alternative algorithm for finding associative rules on large sets of multimodal data, many nodes can occur. This can overflow the database.

3. Results of research and experiments

Let's start with creating a tree [1, 12, 15]:

Select the table headings to build a FP-tree and determine its first instance. It speeds up access to all elements of the tree. You need to create an index plugin to keep track of the total number of certain types of items in the tree. You can use the createTree () function to build a tree based on minimal support and dataset arguments. During the first pass, all elements are checked. In the process, the number of identical on each path is counted. As a result in Figure 1 shows the title table.



TID	Items
1	A, B
2	B, C, D
3	A, C, D, E
4	A, D, E
5	A, B, C

Item	Frequency
A	4
B	3
C	3
D	3
E	2

Threshold = 2

Item	Ordered Frequency
A	4
B	3
C	3
D	3
E	2

Figure 1: Result of the Frequent pattern tree

The updateTree () function increases the FP tree with a set of elements (Figure 2a)).

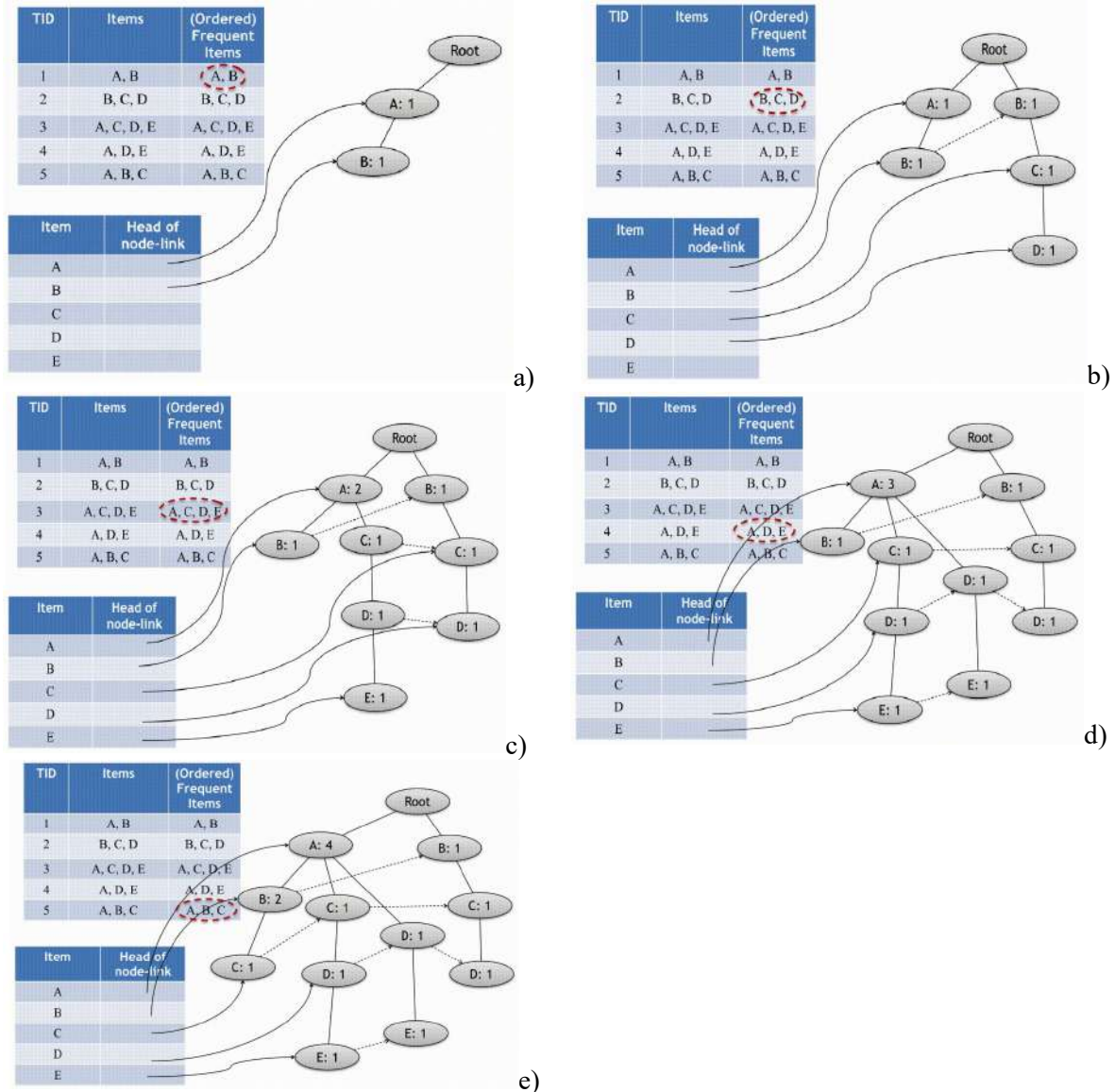


Figure 2: a) Sets of items found in the database after scanning; b) The root is represented as a zero value; c) A tree branch is built from sets of transaction elements in descending order; d) The branch of the transaction has a common prefix to the root; e) Increase by 1 total node and the number of new nodes.

The updateHeader () function allows you to fix connections with nodes, so Fig. 2 b shows an FP-tree with instances of the required elements.

In fig. 2c shows the process when it is impossible to obtain list entries using the createTree () function. In the process of processing, a dictionary is expected, in which there is a set of elements and the frequency of their use.

The createInitSet () function does this transformation for us (Figure 2d)).

Finally, we can easily generate all frequent sets (Figure 3).

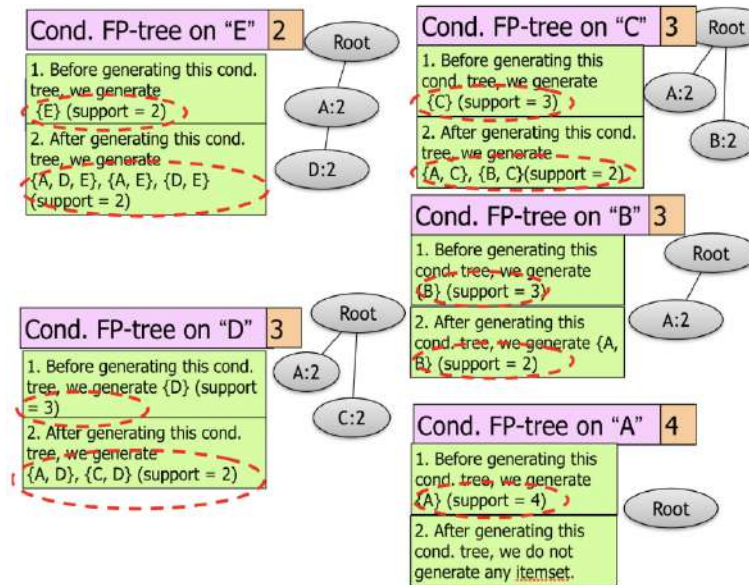


Figure 3: Generated frequent sets

3.1. Development of algorithms for solving a functional problem

Let's start creating the algorithm itself. There are three main steps to finding frequent sets of elements in an FP tree:

- 1) Get a database of conditional templates from the FP-tree.
- 2) The next step is to create a conditional tree based on conditional templates.
- 3) To build a node, repeat steps 1 and 2 until one element remains.

The ascendTree () function, which climbs our tree and collects the names of the elements it encountered (Figure 4):

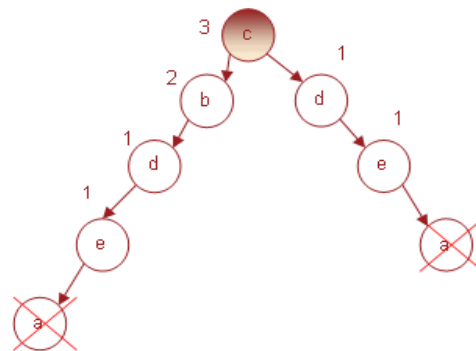


Figure 4: AscendTree () function results

The findPrefixPath () function sorts through the list until it reaches the end. For each element it encounters, it calls the ascendTree () function.

Figure 5 shows the process when condPats is returned to the list and added to the template dictionary.

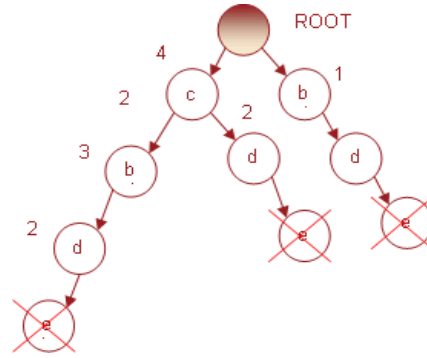


Figure 5: FindPrefixPath () function results

Giving search parameters for our algorithm, it finds the necessary data for us.

Command for the algorithm: indPrefixPath('x', myHeaderTab['x'][1]).

Results: {frozenset({'z'}): 3}.

3.2. Definition and estimation of qualitative indicators of algorithms, comparing with existing ones

Data preprocessing is required to apply any algorithms, including the classical Apriori algorithm and the alternative search algorithm. The data must be without extra attributes, precise, not zeros, etc. Reprocessing should be treated with all responsibility because the best result depends on their effective processing.

Two data sets will be used for the experimental study. Datasets were obtained from the UCI machine learning databases (Table 1).

Table 1

Data set details

Dataset name	Number of Instances	Number of Attributes
Mushroom	8124	22
Super Market	4627	217

We are analyzing the data in the Table 1, according to the sets of input elements Mushroom and Supermarket, we can conclude that the efficiency of the alternative algorithm for finding associative rules is better than the classical algorithm a priori. This difference in the execution time of transactions is presented in Fig. 6-9. Execution time is when to mine frequent data patterns with different transactions (Figures 6-9).

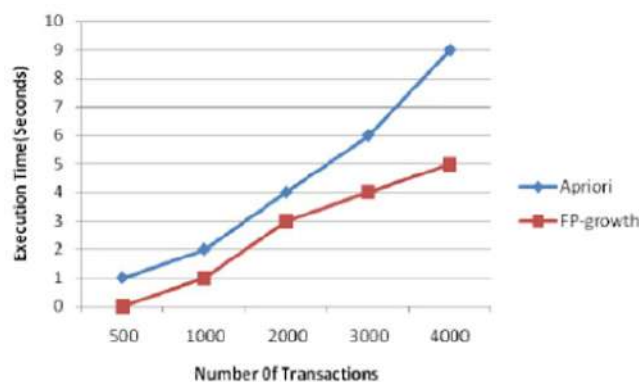


Figure 6: Execution time with increased transactions for the Mushroom date set (support = 0.1)

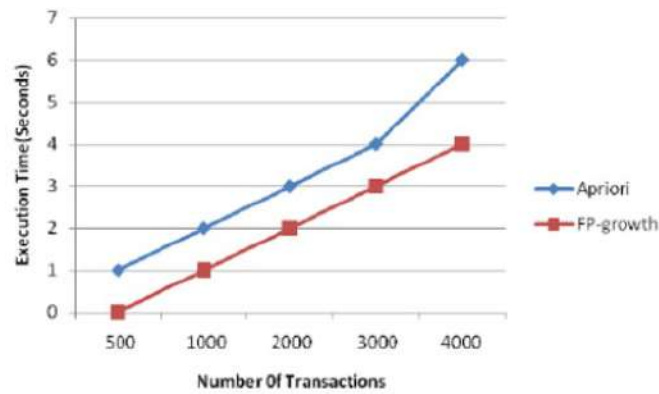


Figure 7: Execution time with increased transactions for the Mushroom date set (support= 0.5)

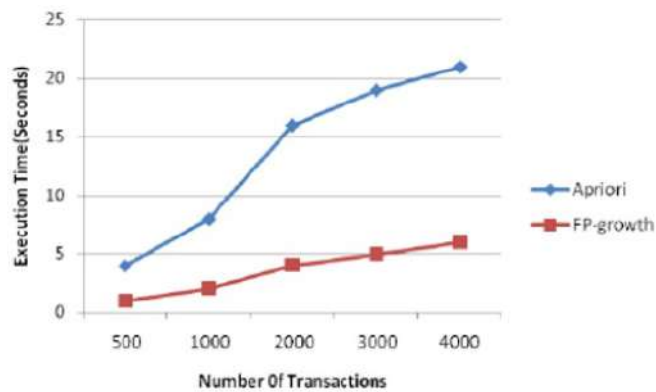


Figure 8: Execution time with increased transactions for the SuperMarket date set (support = 0.1)

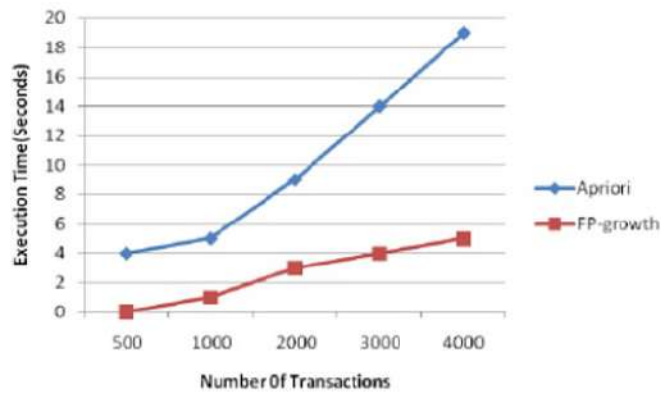


Figure 9: Execution time with increased transactions for the SuperMarket date set (support = 0.5)

A comparison of the curves in Figures 10-13 showed that when using an alternative algorithm search algorithm, as the number of transactions increases, the execution rate also increases, but support levels are minimal. This means that the FPG- algorithm is more efficient in terms of saving time than the Apriori.

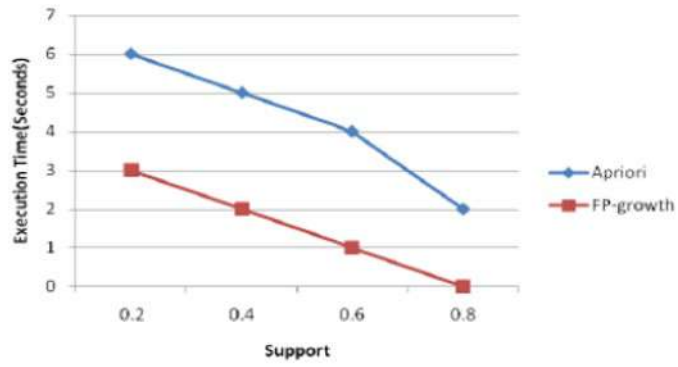


Figure 10: Execution time with different levels of support for the Mushroom set date (2000 transactions)

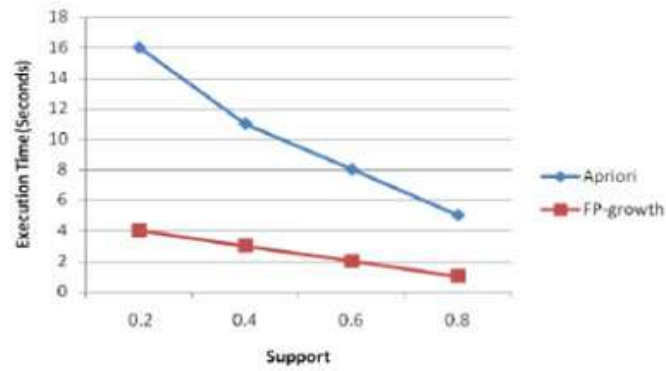


Figure 11: Execution time with different levels of support for the Mushroom set date (4000 transactions)

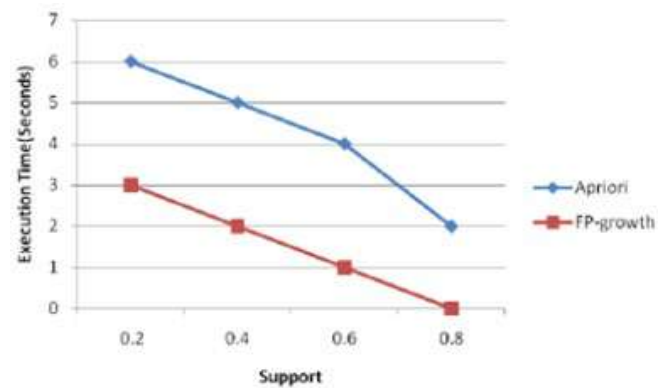


Figure 12: Execution time with different levels of support for the SuperMarket date set (2000 transactions)

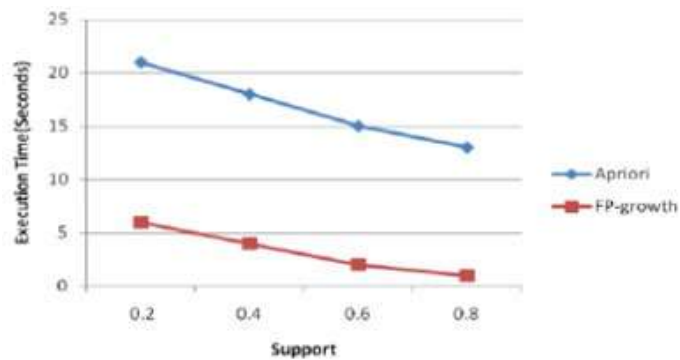


Figure 13: Execution time with different levels of support for the SuperMarket date set (4000 transactions)

Figures 10-13 show the performance analysis of execution time Apriori algorithm and FP-Growth algorithm with different levels of support. This shows the time required to be completed.

The FP-Growth algorithm is significantly smaller compared to the Apriori algorithm.

4. Conclusions

So, we got acquainted with the basic theory of ARL ("who bought x, also bought y") and the basic concepts and metrics (support, confidence, lift and conviction).

The principle of operation of the FP-Growth algorithm was described in detail, both its mathematical part and the software part on the example of the Python language.

FPG is more effective than a priori. It is convenient because the set of multimodal source data is represented as a tree. If each element occurs several times, then on the FP-tree, it is defined as a node, and its index indicates the number of its repetitions in the set.

A comparison of the work of the two algorithms showed that with the increase in the size of the original data set, the time spent searching for the most common element. (Figure 14).

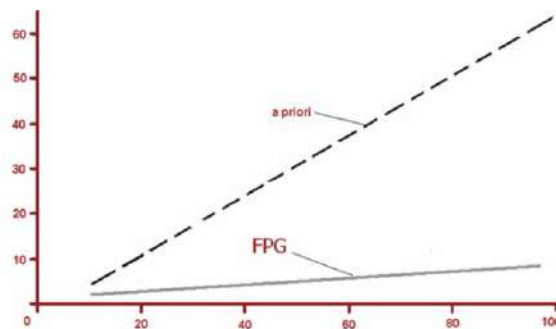


Figure 14: Graph of comparison results of Apriori and FP-Growth algorithms

The work of an alternative algorithm search algorithm provides the most efficient and complete extraction of frequent subject sets to compress database transactions. After all, the divide and conquer technique is used when constructing an FP tree. With its help, it is allowed to perform decomposition of one complex task into several simple ones. in the process of completing the stages of the algorithm, it is possible to avoid the procedure of losses in the generation of candidates.

As a result of the study of the two algorithms on the sets of sizeable multimodal data, the efficiency of searching for associative rules in favor is proved by the FP-Growth algorithm. In the example, the efficiency was determined in the accelerated search and informativeness of the found associative regulations and the ability to predict decisions.

5. References

- [1] H. Alshammari, S.A. El-Ghany, A. Shehab, Big IoT Healthcare Data Analytics Framework Based on Fog and Cloud Computing. *Journal of Information Processing Systems*, Vol. 16, 2020, pp. 1238–1249.
- [2] J. Qi, P. Yang, M. Hanneghan, S. Tang, B. Zhou, A Hybrid Hierarchical Framework for Gym Physical Activity Recognition and Measurement Using Wearable Sensors. *IEEE Internet of Things Journal*, Vol. 6, 2019, pp. 1384–1393.
- [3] T. Wang, L. Qiu, A.K. Sangaiah, A. Liu, Z. Alam Bhuiyan, Y. Ma, Edge-Computing-Based Trustworthy Data Collection Model in the Internet of Things. *IEEE Internet of Things Journal*, Vol. 7, 2020, pp. 4218–4227.
- [4] N. Boyko, K. Kmetyk-Podubinska and I. Andrusiak, Application of Ensemble Methods of Strengthening in Search of Legal Information, in: *Lecture Notes on Data Engineering and*

Communications Technologies, Vol. 77, 2021, pp. 188-200. https://doi.org/10.1007/978-3-030-82014-5_13

- [5] A. Al-Shargabi, F. Siewe, A Lightweight Association Rules Based Prediction Algorithm (LWRCCAR) for Context-Aware Systems in IoT Ubiquitous, Fog, and Edge Computing Environment, in: Proceedings of the Proceedings of the Future Technologies Conference (FTC), 2020.
- [6] C.H. Lee, J.S. Park, An SDN-Based Packet Scheduling Scheme for Transmitting Emergency Data in Mobile Edge Computing Environments. *Journal Series: Human-centric Computing and Information Sciences*, Vol. 11, 2021.
- [7] Y. He, Z. Tang, Strategy for Task Offloading of Multi-user and Multi-server Based on Cost Optimization in Mobile Edge Computing Environment. *Journal of Information Processing Systems*, Vol. 17, 2021, pp. 615–629.
- [8] J.-H. Lee, Next Task Size Prediction Method for FP-Growth Algorithm. *Journal Series: Human-centric Computing and Information Sciences*, Vol. 11, 2021.
- [9] Y.A. Ünvan, Market basket analysis with association rules. *Journal Communications in Statistics - Theory and Methods*, Vol. 50, 2021, pp. 1615–1628.
- [10] Z. Mar, K.K. Oo, An Improvement of Apriori Mining Algorithm using Linked List Based Hash Table, in: Proceedings of the 2020 International Conference on Advanced Information Technologies (ICAIT), 4–5 November 2020.
- [11] Z. Pan, P. Liu, J. Yi, An Improved FP-Tree Algorithm for Mining Maximal Frequent Patterns, in: Proceedings of the 2018 10th International Conference on Measuring Technology and Mechatronics Automation (ICMTMA), 10–11 February 2018.
- [12] N. Boyko, A look trough methods of intellectual data analysis and their applying in informational systems, in: XIth International Scientific and Technical Conference Computer Sciences and Information Technologies (CSIT), 2016, pp. 183-185.
- [13] M. Yin, W. Wang, Y. Liu and D. Jiang, An improvement of FP-Growth association rule mining algorithm based on adjacency table, in: 2nd International Conference on Material Engineering and Advanced Manufacturing Technology (MEAMT 2018), Vol. 189, 2018. <https://doi.org/10.1051/mateconf/201818910012>.
- [14] J. Hao, M. He, A Parallel FP-Growth Algorithm Based on GPU, in: IEEE, International Conference on E-Business Engineering. IEEE Computer Society, 2017, pp. 97-102.
- [15] H.Y. Chang, J.C. Lin, M. L. Cheng, A Novel Incremental Data Mining Algorithm Based on FP-growth for Big Data, in: International Conference on NETWORKING and Network Applications, 2016, pp. 375-378.
- [16] J. Heaton, Comparing dataset characteristics that favor the Apriori, Eclat or FP-Growth frequent itemset mining algorithms, in: Southeastcon. IEEE, 2017, pp. 1-7.
- [17] J. Hao, M. He, A Parallel FP-Growth Algorithm Based on GPU, in: IEEE, International Conference on E-Business Engineering. IEEE Computer Society, 2017, pp. 97-102.
- [18] B. Subbulakshmi, B. Dharini, C. Deisy, Recent weighted maximal frequent itemset Mining, in: International Conference on I-Smac. IEEE, 2017, pp. 391-397.
- [19] T. Willhalm, FPTree: A Hybrid SCM-DRAM Persistent and Concurrent B-Tree for Storage Class Memory, in: International Conference on Management of Data, 2016, pp. 371-386.
- [20] Y. Zeng, Y. Shiqun, L. Jiangyue, M. Zhang, Research of Improved FP-Growth Algorithm in Association Rules Mining. *Scientific Programming*, Vol. 5, 2015, pp.124-136.
- [21] C. Jun, G. Li, An improved FP-growth algorithm based on item head table node. *Journal of Information Technology*, Vol. 12, 2013, pp. 34–35.

Aggregate Parametric Representation of Image Structural Description in Statistical Classification Methods

S. Gadetska¹, V. Gorokhovatskyi², N. Stiahlyk³ and N. Vlasenko⁴

¹ Kharkiv National Automobile and Road University, Kharkiv, Ukraine

² Kharkiv National University of Radio Electronics, Kharkiv, Ukraine

³ Educational and Scientific Institute "Karazin Banking Institute" V.N. Karazin Kharkiv National University, Kharkiv, Ukraine

⁴ Simon Kuznets Kharkiv National University of Economics, Kharkiv, Ukraine

Abstract

Finding effective classification solutions based on the study of the processed data nature is one of the important tasks in modern computer vision. Statistical distributions are a perfect tool for presenting and analyzing visual data in image recognition systems. They are especially effective when creating new feature spaces, particularly, by aggregating descriptor sets in some appropriate way, including bits. For this purpose, it is natural to apply the number of criteria designed to compare the distribution parameters of the analyzed samples. The article develops a speed-efficient method of image classification by introducing aggregate statistical features for the composition of the description components. The metric classifier is based on the use of statistical criteria to assess the significance of the classification decision. The developed classification method based on the aggregation of the feature image set is implemented; the workability of the proposed classifier is confirmed. On the examples of the application of variants of the method for the system of the real images features, its effectiveness was experimentally evaluated.

Keywords

Computer vision, key point, descriptor, data aggregation, metric classifier, statistical distribution, processing speed

1. Introduction

The use of statistical data science tools in computer vision systems to build classifiers for visual object images aims to provide the necessary performance indicators based on the study of properties, content, the structure of the etalons and implementation of obtained knowledge in the classification process [1-5]. The finite set of descriptors of the image key points (KP) is considered here as an element of the image space in the environment of vector data with real or binary components in the implementation of structural recognition methods [2]. Recently, BRISK and ORB descriptors with binary components have become popular due to low computational costs [3-5, 11].

Statistical data distributions are a perfect tool for presenting and analyzing the data of visual object using image recognition systems. The number of statistical methods can be considered as fundamental apparatus for making a classification decision if the description of the recognized object is given by a set of vectors. The study of data distributions in the set of KP descriptors confirmed its effectiveness in sense of providing the indicators of the classification quality and processing speed [2].

It may be considered necessary to study in depth the statistical properties of the descriptor set in terms of the main issue of distinguishing multidimensional data to solve the classification problem. This task is

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: svgadetska@ukr.net (S. Gadetska), gorohovatsky.vl@gmail.com (V. Gorokhovatskyi); natastyaglick@gmail.com (N. Stiahlyk), gorohovatskaja@gmail.com (N. Vlasenko)

ORCID: 0000-0002-9125-2363 (S. Gadetska); 0000-0002-7839-6223 (V. Gorokhovatskyi); 0000-0002-0573-3508 (N. Stiahlyk); 0000-0002-2560-9854 (N. Vlasenko)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

especially important in the construction of new effective space features, particularly, by aggregating a descriptor set by their vector components [3, 9].

For this purpose, it is natural to have developed the use of the apparatus of statistical criteria aimed at comparing the distribution parameters of the studied samples. The classifier based on the aggregate features organizes a new data space as a set of descriptors to evaluate the similarity of the feature vectors of the recognized object and the etalon, and the classification is done by optimizing the degree of such similarity.

The probabilistic model of generating vector data of visual object description is a practical approach to formalizing the process of classifier construction, the essence of which is to build and study statistical distributions of object components based on the aggregation and optimization procedures on the set of classes [1, 5]. In spite of the widespread use and applied effectiveness of KP descriptors for the classification of visual objects [2-4], the issue of the statistical nature of these methods and the choice of effective ways to analyze their effectiveness for real data sets remains unexplored [1, 2].

The main task of the paper is to provide a statistical apparatus of data analysis to build and confirm the effectiveness of the image classifier on the aggregate data representation based on a set of key point descriptors.

2. Problem Statement

Let's consider a multidimensional space B^n of binary n - dimension vectors, where descriptions of the object and etalons will be constructed. Description Z is defined on the basis of the KP descriptor set of the visual object in the form of a finite set of s binary vectors:

$$Z = \{z_v\}_{v=1}^s, z_v \in B^n$$

In a more detailed form, let's consider and analyze the description

$$Z = \left\{ \{z_v\}_{v=1}^s \right\}_{i=1}^n, \quad (1)$$

as a matrix of binary values of $s \times n$ size.

We will traditionally consider classification as a transformation

$$K: Z \rightarrow [1, \dots, m],$$

where each class is represented by etalon descriptions E_j , $j = 1, \dots, m$, which are available for analysis [2].

The visual objects classification as assigning their description to one of the etalon classes based on the aggregate representation of the description data using the tools and criteria of mathematical statistics will be studied. Generally, the problem of classification is formally reduced to determining the degree of relevance of two vector sets with binary components.

We will build in a certain way a secondary integrated systems of features

$$P = \{p_k\}_{k=1}^n$$

on the basis of descriptions Z and $\{E_j\}_{j=1}^m$ and implement them in the classification solution. We use a metric approach to determine the degree of similarity of feature values P for the object and etalons.

The introduction of aggregate features contributes to a significant acceleration of the classifier decision process, the gain in comparison with the traditional method of voting descriptors reaches hundreds of times [4, 5]. Also the separation properties of the newly created system of features using traditional statistical criteria will be investigated. The research is a development of the authors' works [2, 3, 5] in the sense of implementing a generalized parametric bitwise representation of a descriptor set and implementation of new variants of classifiers using statistical analysis for transformed data. Data of structural descriptions of etalons are taken from the article [2].

3. Literature review

The formal definition of the classification problem with the description of the image as a set of KP descriptors is formulated in papers [2, 3], which also study the advantages of implementing a structural

description model in the methods of statistical classification [1, 4-8]. It is noted that the primary problem is the excessive computational costs for processing large spatial data sets.

Articles [3-5, 8, 13, 20] investigate statistical models for the synthesis of feature space modifications to reduce the amount of computation, in particular, the application of data aggregation methods by forming distributions and defining statistical data centers. Works [1, 7, 12] are devoted directly to the analysis of learning models for the fixed base of descriptions used in computer vision and the definition of the function of belonging to a fixed system of classes.

Articles [8, 11, 15, 16, 19] discuss the principle of construction feature detectors for the binary descriptors of KP. Studies [1, 2, 7] contain results on the applied implementation of statistical approaches to the visual images classification using an ensemble processing. In [1, 5-7, 13-15] methods of evaluating the effectiveness of intelligent systems using statistical and metric measures of similarity are described. The advantages of statistical solutions such as high processing speed, sufficient resistance to distortion and ensuring the required level of classification efficiency are discussed.

Work [10] is used as sources of traditional and modern methods of statistical evaluation, the book [12] contains a description of applied features of software modeling, and sources [2-5] include the results of authors' research in implementing statistical approaches to develop structural methods image classification. In particular, [2] proposed technologies of component analysis and spatial processing for the classification of visual objects using statistical characteristics of the structural description of the image.

4. Proposed Approaches

We consider a transformation $Z \rightarrow P$, $Z \subset B^n$, from a fixed set Z of binary vectors – KP descriptors for a given object into a numerical vector $P = \{p_k\}_{k=1}^n$, which components are calculated by some rule. This approach will give a possibility to identify and distinguish visual objects on the basis of smaller data, as set of vectors is transformed into a single vector [3].

We will carry out classification on the basis of estimating the differences in the values of vectors P for different descriptions. The representations of them are considered in two different ways which take into account the structural features of the studied data and, as a result, ensure the efficiency of the recognition process.

The first of the proposed ways for constructing a vector P is to find the average sum of binary values (number of units) consecutively for each bit with the number i separately, based on the full set of object description Z . For a fixed description we obtain vectors of the form:

$$P^{(1)} = \{p_i^{(1)}\}_{i=1}^n, p_i^{(1)} = \frac{1}{s} \sum_{v=1}^s z_{vi}, 0 \leq p_i^{(1)} \leq 1 \quad (2)$$

The vector (2) can be represented as an aggregate parameter formed on a set of descriptors by bitwise analysis of data by adding the values of the corresponding bits (columns of the matrix (1)) and dividing by the dimension s of columns.

We believe it is possible to consider the distribution of values of the i -th bit of the object description close to binomial, which is determined by the Bernoulli formula. According to it the probability of occurrence of v units at the i -th bits in the description of s vectors can be found in the following way:

$$P_s^{(i)}(v) = C_s^v p_i^v (1 - p_i)^{s-v}, \quad (3)$$

where p_i is the distribution parameter. From the traditional point of view it is equal to the probability of occurrence of a bit equal 1 for the i -th component of the set Z . Also, from the applied point of view this probability is equal to the i -th component of the vector P determined by the formula (2).

Note that the process of comparing two objects can be based on bitwise comparison of values calculated for each of these objects by formula (3), which can be considered as aggregated by bits. But such an idea for a classification rule is not justified enough for application due to cumbersome calculations.

In this paper, it is proposed to classify the studied objects on the basis of the distribution parameter p_i . Based on general considerations, we will consider the tuple of values $P^{(1)} = (p_1^{(1)}, \dots, p_n^{(1)})$ obtained on the basis of the description Z , like its aggregated parametric representation.

Let's consider $P^{(1j)}, j = 1, \dots, m$ as a vector aggregated by columns of the matrix for the binary description of the etalon E_j with the number $j = 1, \dots, m$, according to (2) and $P^{(1O)}$ as an aggregate vector for the description of the studied object O .

To compare the aggregate descriptions of objects of type (2) and, accordingly, to solve the classification problem, we introduce the classifier

$$K^{(1)} : k^{(1)} = \arg \min_{j \in [1; m]} D^{(1j)} \quad (4)$$

$$D^{(1j)} = \frac{1}{s} \sum_{i=1}^n |p_i^{(1j)} - p_i^{(1O)}|, j = 1, \dots, m \quad (5)$$

where $D^{(1j)}, j = 1, \dots, m$ can be considered as a normalized measure of the similarity of the corresponding vectors, and the expression $|p_i^{(1j)} - p_i^{(1O)}|, i = 1, \dots, n, j = 1, \dots, m$ as the Manhattan distance between them.

Classifier (4) implements the principle of analysis "object – etalon" based on the aggregate vector representation P . We emphasize that expression (4) can be considered as a decisive rule, which is formulated in terms of metrics, in particular, Manhattan.

To confirm the significance of the decision, as well as to control the obtained result of the classifier (4) with the involvement of aggregate vectors, we use methods of mathematical statistics, namely, a paired two-sample t-test for averages [10], which provides pairwise comparison of the studied objects as vectors P for a statistically significant difference in their average values. When using this test, two samples of the same volume are considered, in which the elements have a fixed location (as coordinates).

In the process of testing the null hypothesis regarding the equality of the averages in these samples, Student's statistics is used [10]; a level of significance α is established, equal to the probability of making an error of the I type, i.e. rejecting the null hypothesis if it was correct; based on the initial data, the p -value is calculated as the maximum possible probability of error of the I type. Then, if the p -value is less than the established α , then the null hypothesis is rejected, and an alternative hypothesis is accepted regarding the significant difference of the means (at the level of significance α). Otherwise, there is no reason to reject the null hypothesis of no statistically significant differences between the means [10].

Note that for the sake of universality of the study, it may be appropriate to perform analysis of variance of aggregate vectors constructed from etalons $E_j, j = 1, \dots, m$, in order to ensure that the reduction of data dimensionality did not affect the difference in the etalon set. Also, since the structure of the vectors aggregated by formula (2) requires the consideration of paired samples, in this case it is possible to use only nonparametric analysis of variance, for example, in the form of the Friedman test [10].

The second way for constructing a vector P is to represent the components of the vector as the average sum of unit bits separately for each binary descriptor of the description. In this case we obtain aggregate description vectors in the form:

$$P^{(1)} = \{p_i^{(1)}\}_{i=1}^n, p_i^{(1)} = \frac{1}{s} \sum_{v=1}^s z_{vi}, 0 \leq p_i^{(1)} \leq 1 \quad (6)$$

Then $P^{(2j)}, j = 1, \dots, m$ is a vector aggregated by the matrix's rows of the etalon description E_j with the number $j = 1, \dots, m$ by expression (6); and $P^{(2O)}$ is an aggregated vector according to the description of the studied object O .

When aggregating vectors in the form (6) we obtain the independent samples and directly coordinate comparison of them is impossible. Therefore, for the correct application of the classifier according to scheme (4) - (5), we propose to perform pre-ranking (for example, in ascending order) aggregated according to (6) vectors of the two descriptions being compared. In this case, the classifier takes the form:

$$K^{(2)} : k^{(2)} = \arg \min_{j \in [1; m]} D^{(2j)} \quad (7)$$

$$D^{(2j)} = \frac{1}{n} \sum_{i=1}^s |p_i^{(2j)} - p_i^{(2O)}|, j = 1, \dots, m \quad (8)$$

Due to the independent nature of the data testing for significant differences by statistical methods also in this case requires the use of appropriate approaches. The essence of two-sample t-test for averages for two independent samples which is appropriate in this situation is to compare the averages of two sets of disordered elements, which are calculated separately for each binary descriptor of the full description. The test procedure is the same as for the case of dependent samples, and differs only in the formula, according to which the statistics are calculated on the basis of the given data [10].

Also note that in order to confirm the fact of statistically significant differences of the aggregate vectors (6) constructed on the etalons $E_j, j=1, \dots, m$, it can be appropriate to provide analysis of variance of them. As the structure of aggregate vectors requires the consideration of independent samples, it is possible to use the methods of both parametric and nonparametric variance analysis (for example, in the form of the Kruskal-Wallis test [10]).

5. Experiments

Let's consider an example with experimental descriptions of three fixed etalons E_1, E_2, E_3 , and E_4 , obtained from E_1 by rotation. Examples of images based on the results of software modeling with the formed coordinates of the BRISK KP descriptors are shown in Figure 1 [2, 11, 16]. For the descriptions of these images in the form of a set of descriptors, our calculations are performed. In the example, $n = 512$ is the dimension of the descriptor, $s = 500$ is the number of descriptors in the description, m is the number of the etalons ($m = 3$).

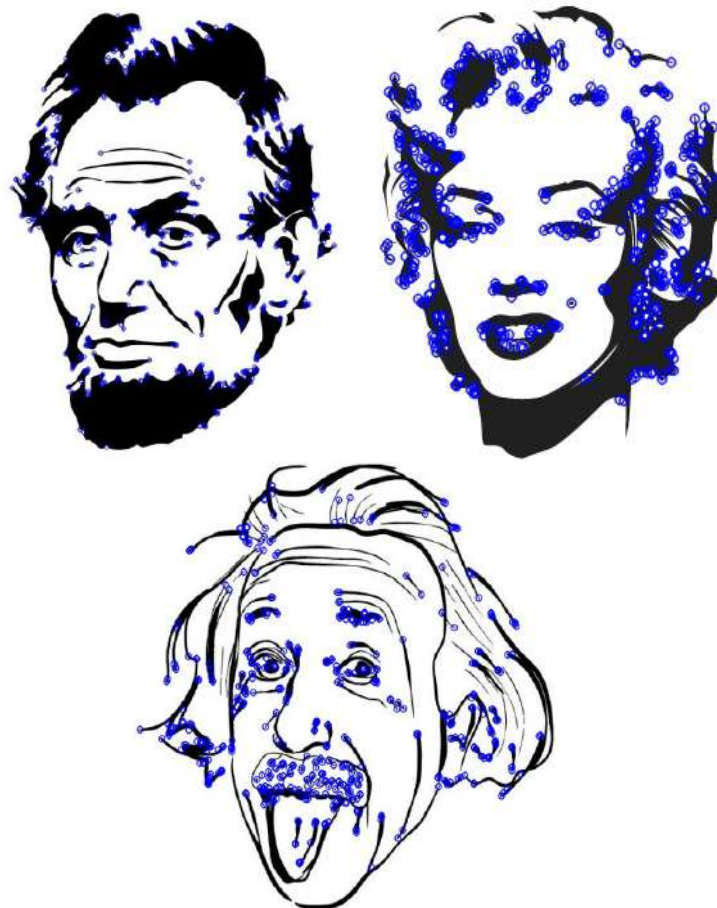


Figure 1: Etalons with marked KP

We present the result of implementation of the proposed classification approach based on the values of the parameters P calculated on the database of three etalons E_1, E_2, E_3 (Fig. 1) and the image E_4 , transformed by rotation of E_1 . Note that the representation using KP descriptors provides invariance to the transformations of displacement, rotation and scale of the analyzed object [3].

Fragments of the calculation results for aggregate vectors $P^{(1j)}, j = 1, 2, 3, 4$ of the form (2) for objects E1, E2, E3, E4 are given in the table 1.

Table 1

Fragments of vectors $P^{(1j)} = \{p_i^{(1j)}\}_{i=1}^{512}, j = 1, 2, 3, 4$

Component number	$P^{(11)}$	$P^{(12)}$	$P^{(13)}$	$P^{(14)}$
1	0.896	0.794	0.738	0.860
2	0.636	0.610	0.788	0.624
3	0.820	0.742	0.682	0.790
4	0.782	0.730	0.650	0.756
5	0.580	0.500	0.770	0.556
6	0.612	0.612	0.586	0.610
...	...	...	...	...
506	0.246	0.398	0.412	0.254
507	0.426	0.484	0.530	0.440
508	0.474	0.494	0.552	0.470
509	0.410	0.584	0.606	0.420
510	0.290	0.470	0.514	0.308
511	0.440	0.598	0.590	0.480
512	0.390	0.598	0.596	0.406

According to the formulas (4, 5) using the data of E4 we have:

$$D^{(11)} = 0.017; D^{(12)} = 0.071; D^{(13)} = 0.105.$$

As we can see, the application of the classifier (4) gives the correct recognition of the object E4 as a transformed etalon E1 because it has the best similarity with the first etalon.

To confirm the fact of statistically significant closeness of E4 to E1 as well as statistically significant difference of E4 from other etalons E2, E3 we use a paired two-sample t-test for averages

applied to samples represented by aggregate vectors $P^{(1j)} = \{p_i^{(1j)}\}_{i=1}^{512}, j = 1, 2, 3, 4$, formed by (2). The results are shown in the table 2.

Table 2

The results of the application of a paired two-sample t-test

Samples	E1, E4	E2, E4	E3, E4
p – value	0.252	0.002	0.0000000006
Significance	no	yes	yes

As we can see the first p - value is much higher the significance level $\alpha = 0.05$ (equal to the probability of error of the first type). That indicates the absence of statistically significant differences between $P^{(11)}$ and $P^{(14)}$ (i.e. between the first etalon E1 and object E4 transformed from E1). Other obtained p - values are less than 0.05, confirming a statistically significant difference in pairs between $P^{(14)}, P^{(12)}$ and $P^{(14)}, P^{(13)}$ (i.e. between E4, E2 and E4, E3).

Note, that for a large sample size (in the example we have $n = 512$), checking the data for compliance with the normal distribution law when using a paired t-test is not mandatory [10].

Note also that the visual comparison of bar charts which is a graphical representation of aggregated vectors by formula (2) is a clear confirmation of the results obtained on the difference of

objects (Fig. 2) (similarity between $P^{(14)}$, $P^{(11)}$ and significant difference between pairs $P^{(14)}$, $P^{(12)}$ and $P^{(14)}$, $P^{(13)}$).

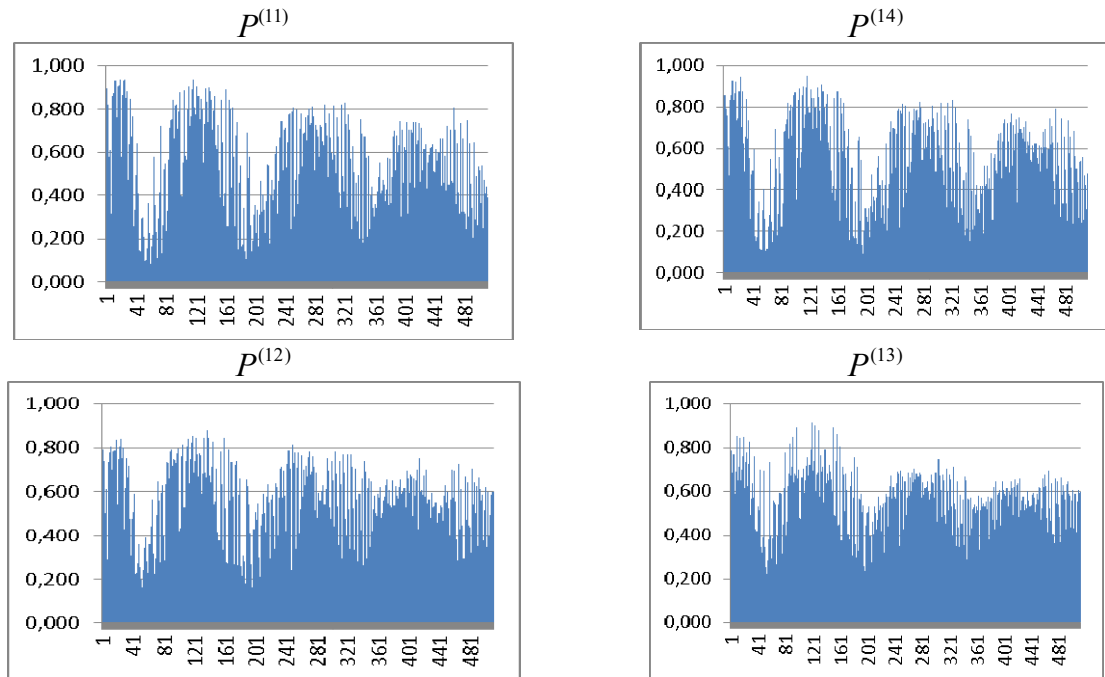


Figure 2: Graphic representation of vectors $P^{(1j)}$, $j = 1, 2, 3, 4$

Fragments of the calculation results for aggregate vectors $P^{(2j)}$, $j = 1, 2, 3, 4$ of the form (6) for objects E1, E2, E3, E4 are given in the table 3.

Table 3

Fragments of vectors $P^{(2j)} = \{p_v^{(2j)}\}_{v=1}^{500}$, $j = 1, 2, 3, 4$

Component number	$P^{(21)}$	$P^{(22)}$	$P^{(23)}$	$P^{(24)}$
1	0.588	0.648	0.484	0.379
2	0.531	0.611	0.625	0.533
3	0.467	0.607	0.457	0.592
4	0.531	0.639	0.672	0.488
5	0.590	0.494	0.463	0.268
6	0.471	0.621	0.602	0.432
...	...	...	...	...
494	0.480	0.590	0.523	0.570
495	0.391	0.541	0.383	0.494
496	0.582	0.525	0.463	0.523
497	0.549	0.568	0.523	0.479
498	0.416	0.334	0.551	0.521
499	0.547	0.383	0.559	0.545
500	0.459	0.385	0.523	0.576

According to the formulas (7, 8) we obtain, taking into account the preliminary ranking of the components of the vectors $P^{(2j)}$, $j = 1, 2, 3, 4$ and using the data of E4:

$$D^{(21)} = 0.005, D^{(22)} = 0.012, D^{(23)} = 0.033.$$

As a result of the analysis, we see again that the application of the classifier (7) leads to the correct recognition of the object E4 as a transformed etalon E1, as its similarity with the first etalon is the greatest.

To confirm the fact of statistically significant closeness of E4 to E1 as well as statistically significant difference of E4 from other etalons E2, E3 we use a two-sample t-test for averages with different variances

applied to samples represented by aggregate vectors $P^{(2j)}, j = 1, 2, 3, 4$ formed by (6). The results are shown in the table 4.

Table 4

The results of the two-sample t-test

Samples	E1, E4	E2, E4	E3, E4
p - value	0.843	0.024	0.0000000005
Significance	no	yes	yes

We also see that the first p - value is much higher the significance level $\alpha = 0.05$. That indicates the absence of statistically significant differences between E1 and E4. Other p - values are less than 0.05, confirming a statistically significant difference in pairs between E4, E2 and E4, E3.

Similar to the previous one, the visual comparison of bar diagrams, which are a graphical interpretation of the vectors aggregated by formula (6), confirms the results obtained (Fig. 3).

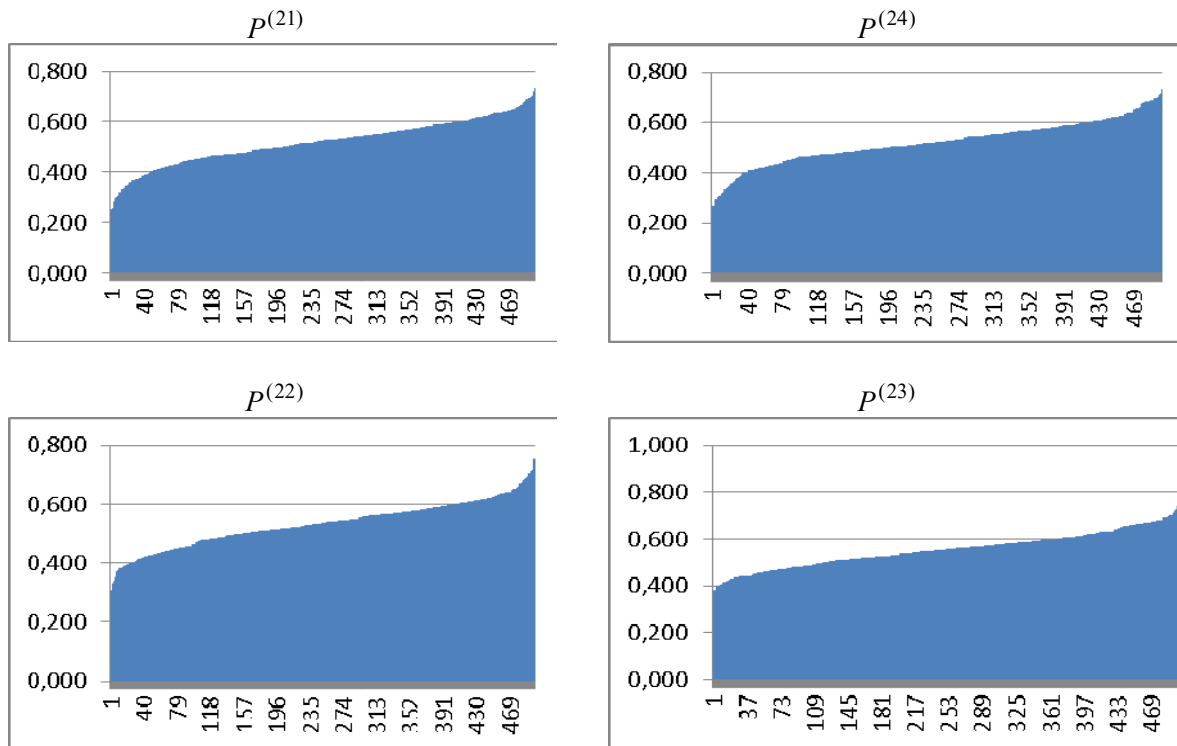


Figure 3: Graphic representation of vectors $P^{(2j)}, j = 1, 2, 3, 4$

6. Results

The main result of this research is the development of models for the image classification based on the apparatus of statistical analysis of component sets of the object description and metric means of classification deciding. The synthesis of an aggregate feature system based on KP descriptor set makes possible to build a classifier that works successfully for the real images database.

The proposed approaches of data analysis models are based on the degree of similarity between the object and etalons are workable and quite effective. Computational simulation performed on the example with 3 etalons confirmed efficiency of the proposed method using statistical criteria of significant data differences.

7. Conclusion and Future Work

Statistical data analysis remains a powerful research factor for intelligent decision making, machine learning, and data science. The conducted research makes it possible to evaluate the applied efficiency of the application of the feature aggregate system for the effective implementation of the visual object classification by a set of key point descriptors. The research has shown that the available information in the form of a bit representation of the object description is quite sufficient for statistical differentiation of data for different visual objects.

The novelty of the investigation is the further development of the image classification method using an integrated statistical feature system for structural description, confirmation of its effectiveness and the significance of this system for classification within the given image database. The proposed classifier construction method allows further generalization in terms of fragment size aggregation that implies reduction of processing time.

8. Acknowledgements

The work was performed within the framework of the state budget research of Kharkiv National University of Radio Electronics "Deep hybrid systems of computational intelligence for data flow analysis and their rapid learning" (№ ДР0119U001403).

9. References

- [1] P. Flach, Machine learning. The Art and Science of Algorithms that Make Sense of Data, New York, NY , Cambridge University Press, 2012.
- [2] S.V. Gadetska, V.O. Gorokhovatskyi, N. I. Stiahlyk, N.V. Vlasenko, Statistical data analysis tools in image classification methods based on the description as a set of binary descriptors of key points. Radio Electronics, Computer Science, Control, №4 (2021) 58-68. doi: 10.15588/1607-3274-2021-4-6
- [3] V.O. Gorokhovatskyi, S. V. Gadetska, Determination of Relevance of Visual Object Images by Application of Statistical Analysis of Regarding Fragment Representation of their Descriptions. Telecommunications and Radio Engineering, 78 (3) (2019) 211–220.
- [4] Y.I. Daradkeh, V. Gorokhovatskyi, I. Tvoroshenko, S. Gadetska, M. Al-Dhaifallah, Methods of Classification of Images on the Basis of the Values of Statistical Distributions for the Composition of Structural Description Components. IEEE Access, 9 (2021) 92964-92973, doi: 10.1109/ACCESS.2021.3093457
- [5] S. Gadetska, V. Gorokhovatskyi, N. Stiahlyk, Image structural classification technologies based on statistical analysis of descriptions in the form of bit descriptor set. In CEUR Workshop Proceedings: Computer Modeling and Intelligent Systems (CMIS-2020), 2608, 2020, pp.1027-1039. URL: <http://ceur-ws.org/Vol-2608/>
- [6] A. Oliinyk, S. Subbotin, V. Lovkin, The System of Criteria for Feature Informativeness Estimation in Pattern Recognition. Radio Electronics, Computer Science, Control. № 4 (2017) 85–96. doi: 10.15588/1607-3274-2017-4-10
- [7] J.F. Peters, Foundations of computer vision: Computational Geometry. Visual Image Structures and Object Shape Detection, Cham, Switzerland: Springer International Publisher (2017) 417 p.
- [8] A. Hietanen and et al., A comparison of feature detectors and descriptors for object class matching, Neurocomputing, 184 (2016) 3-12.
- [9] O. Gorokhovatskyi, V. Gorokhovatskyi, O. Peredrii, Analysis of Application of Cluster Descriptions in Space of Characteristic Image Features. Data, 3(4), 52. (2018). doi: 10.3390/data3040052. URL: <https://www.mdpi.com/2306-5729/3/4/52>

- [10] N. Berk Kenneth, P. Carey. Data Analysis with Microsoft Excel, Brooks/Cole, 2009.
- [11] S. Leutenegger, M. Chli, Y. Roland, Siegwart. BRISK, Binary Robust Invariant Scalable Keypoints, Computer Vision (ICCV) (2011) 2548–2555.
- [12] L. Shapiro, J. Stockman, Computer vision, New Jersey, USA: Prentice Hall, 2001.
- [13] F. Malik, B. Baharudin, Analysis of distance metrics in content-based image retrieval using statistical quantized histogram texture features in the DCT domain. URL: <https://www.sciencedirect.com/science/article/pii/S1319157812000444>.
- [14] M. Muja, D.G. Lowe, Fast Matching of Binary Features. Conference on Computer and Robot Vision, 2012, pp. 404–410. doi: 10.1109/CRV.2012.60.
- [15] S. Kim, I.-S. Kweon, Biologically motivated perceptual feature: Generalized robust invariant feature. In Asian Conf. of Comp. Vision (ACCV-06), 2006, pp. 305–314.
- [16] Open CV Open Source Computer Vision. URL: <https://docs.opencv.org/master/index.html>.
- [17] P. Wei, J. Ball, and D. Anderson, Fusion of an ensemble of augmented image detectors for robust object detection, Sensors, 18(3), (2018). doi:10.3390/s18030894.
- [18] I.S. Tvoroshenko, O.O. Kramarenko, Software determination of the optimal route by geoinformation technologies, Radio Electronics, Computer Science, Control, 3, (2019) 131–142.
- [19] P. Gayathiri, M. Punithavalli, Partial Fingerprint Recognition of Feature Extraction and Improving Accelerated KAZE Feature Matching Algorithm. International Journal of Innovative Technology and Exploring Engineering (IJITEE), Volume-8 Issue-10, 2019, pp. 3685–3690. doi: 10.35940/ijitee.J9653.0881019
- [20] A. Oliinyk, S. Subbotin, A stochastic approach for association rule extraction. Pattern Recognition and Image Analysis, Vol. 26, No. 2 (2016) 419–426. doi: 10.1134/S1054661816020139

Using the Temporal Data and Three-dimensional Convolutions for Sign Language Alphabet Recognition

Serhii Kondratiuk^{a,b}, Iurii Krak^{a,b}, Vladislav Kuznetsov^a and Anatoliy Kulias^a

^a Glushkov Cybernetics Institute, Kyiv, 40, Glushkov ave., 03187, Ukraine

^b Taras Shevchenko National University of Kyiv, Kyiv, 64/13, Volodymyrska str., 01601, Ukraine

Abstract

Research is the development of the communication technology for detection of gestures of Ukraine sign alphabet. Basis of neural network with mobilenet architecture was improved with a new deep learning architecture – mobilenetv3. Another improvement lies in the field of the dataset – additional portion of train dataset was collected, and test dataset became more diverse, also due to improved data augmentation techniques. Deep learning model was improved with three dimensional convolution and data processing technique in a form of temporal frames. A lot of experiments to have a consistent improvement in model performance with better data augmentation, novel mobilenetv3 architecture as a basis and spatio-temporal frame are demonstrated the effectiveness proposed approach.

Keywords ¹

gesture detection, 3d convolutions, mobilenetv3, temporal data, data augmentation, sign language

1. Introduction

The research is based on previous work in the field of gesture detection, specifically Ukrainian sign language [1]-[3]. The previous work introduced a new communication technology, both for studying and testing of Ukrainian sign language, with possibility to scale with other languages. One of the most prominent features was cross-platform of the technology [4]-[6].

Sign language is a frequently used method of communication among people with special communication needs. Those with hearing disabilities may utilize supplemental software to connect with society and within their own group. The dactyl alphabet should be learned utilizing gesture recognition technology.

With development of deep learning network and hardware which is a sufficient fit to train such a network, the gesture detection was implemented using a convolutional neural network using novel architecture mobilenet, which allowed to deploy the technology on mobile phones and for it to be a truly cross-platform.

Further improvement to the approach is present in the research, specifically the three-dimensional convolutions with the use of spatio-temporal data frame, which in combination with the newer mobilenetv3 architecture allows to achieve higher quality of sign recognition, due to ability to detect not only static but also dynamic features. Configuration and optimization of the approach are also part of the research.

¹CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL: sergey.kondrat1990@gmail.com (S.Kondratiuk); yuri.krak@gmail.com (I.Krak); kuznetsow.wlad@gmail.com (V.Kuznetsov);
anatoly016@gmail.com (A.Kulias)
ORCID: 0000-0002-5048-2576(S.Kondratiuk); 0000-0002-8043-0785(I. Krak); 0000-0002-1068-769X(V. Kuznetsov);
0000-0003-3715-1454 (A.Kulias)

2. Existing research

Sign gesture detection is a complex task, due to human hand being a highly dynamic object, able to represent signs with different peculiarities. Also, signs can be observed from different angles in different scales and in both axes. Another issue is lighting condition and physical parameters of the hand itself and surrounding environment.

Various algorithms based on conventional computer vision with hand-crafted features, such as the histogram of oriented gradients [7], bag-of-features [8], hyperplanes separation [9]. Unfortunately, due to mentioned before, peculiarities of the hand and its high-dynamic nature, such approaches are not robust enough and suffer a lot from changes in environment, background, or quality of the input images. Also, they are not scalable with the enlargement of image database of hand gestures samples.

Current state-of-the-art hand gesture recognition architectures [10], [11], [12] are based on convolutional neural networks, because they can overcome the issues with the environment and become robust to changes in background and surrounding, scale of the image and hand within it, and quality of the image itself. Another prominent feature is the ability of the convolutional neural network to improve as the training dataset grows and becomes more diverse.

3D convolutions allowed the same models to benefit from sequence of pictures, e.g., videos with activities. They show high performance with large amounts of data, even in the form of pre-recorded videos. (AlexNet [13], Sports-1M [14], Kinetics [15], Jester [16]). No overfitting happens with such huge and diverse datasets.

In order to train and deploy deep learning models on low performance devices, such as smartphones, a subset of lightweight architectures with typically smaller amount of parameters and more effective structure were developed (SqueezeNet[17], MobileNet[18], MobileNetV2 [19], ShuffleNet [20] and ShuffleNetV2 [21], MobileNetV3 [22]). [23] presents Sequential Pattern Mining for tree topologies based recognition in their research.

3. Proposed approach

The proposed approach lies in two domains:

- Improving techniques for data augmentation
- Presenting, improving, and configuring approach of spatio-temporal data.

Such development allows to make the trained model become more robust to changes in input data, without the need to collect more of the datasets manually. On the other hand, it allows to improve the approach of detecting a sign on an image sequence (or video), and gestures are mostly a highly dynamical object, whilst transitions from one gesture to another could be highly diverse, and a model would benefit a lot from training on such transitions. All dataset processing and model training was performed using cross-platform framework [24].

Another improvement is in using a more novel, advanced, and optimized architecture as a basis – mobilenetv3.

3D convolution is used to improved detection with sequence of images (video). Such developments [25]-[27] show significant improvement of architectures with such enhancements in tasks with dynamic activities. Combination of lightweight architecture with such 3d convolutions allows to improve performance with dynamic sign detection on videos and maintain a possibility to deploy cross-platform, even on a low-performance device such as a smartphone.

Spatio-temporal detectors can build spatial descriptors that include both spatial and temporal information. Both single images and sequences of images can be utilized as input for the model in the research.

It is possible to train the model to be more resistant to change in such a dynamic object as hand by analyzing multiple adjacent images in the sequence at once. This allows the network to be taught to consider the temporal aspect, i.e., the dynamics of change in movements in multiple input images. If an image contains artifacts, bad lighting, is fuzzy, or is obstructed in some manner, spatio-temporal approach can be used to smooth things out by utilizing surrounding frames in sequence.

4. Spatio-temporal frame

The idea of the of the spatio-temporal frame lies in concepts: to collect all the spacious information from the image, needed for the convolutional network to detect gestures, and second, to merge multiple frames information, to collect temporal component of the data, meaning dynamic changes from image to image. Only three-dimensional convolutions would be able to process such input and thus be able to detect patterns in temporal component over a sequence of input images.

A single image sequence could be treated as a single training sample for a convolutional neural network with 3d convolutions. However, is it unclear in such case, which would be the required length of the video. Moreover, it would limit somehow the format of the input data, requiring it to be of a specific predefined length or duration.

As a solution to such issues, in the research, instead of setting a requirement to input video, a concept of spatio-temporal frame was introduced.

It is similar to a floating window concept, which is widely used in object detection, when the image is passed in a predefined order with a window of a size significantly smaller than size of the image. Similarly, the spatio-temporal frame has a predefined size (for instance, 8 frame) and goes through the video of arbitrary length in a chronological order. Another major concept is the overlapping of the spatio-temporal frame, in other words, how many last frames in previous frame and first frames in next frame are the same. Having this parameter set too high, the result will be in excessive amount of data which will be generated by the preprocessing of training dataset and, finally, the neural network will train redundantly huge amount of data, which is mostly the same, also taking longer amount of time to train until convergence.

Thus, we have two hyper-parameters, define at the step of preprocessing of the dataset (splitting videos into spatio-temporal frames of fixed length) – size of the frame and number of overlapped frames. These parameters affect architecture, speed, and quality of the trained model, so they such be tuned carefully, and such tuning process was a part of the research, and optimal parameters were defined based on model performance. These hyperparameters also affect model architecture (size of layers).

Therefore, single spatio-temporal frame can be presented as:

$$D = \{d_{i-k}, \dots, d_i, \dots, d_{i+k}\}, \quad i = \overline{1, n-k}, \quad (1)$$

where k is the number of previous and subsequent frames from the current, from which a sequence of images is formed (Fig. 1).

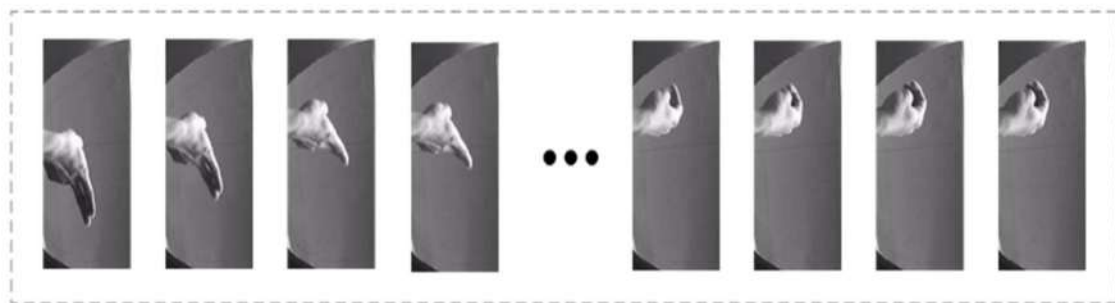


Figure 1: Two subsequences created from a single video stream

Figure 2 shows the schema in which one video is divided into spatio-temporal frames with optimal parameters.

Also, during split into train set and test set, a distribution of the data should be maintained in terms of size, quality, focal length, lighting, background, artifacts, blur and etc.

A uniform data processing approach was designed to convert them to a generic form for further computations inside the specified recognition model, both at the training and recognition stages.

As a result, from one video with a sign it is possible to obtain multiple sub sequences.

The process of tuning such hyper-parameters requires a pipeline, which consists of

- dataset preprocessing

- model architecture selection
- model training
- testing on a predefined test set (for the testing to be in fair conditions)

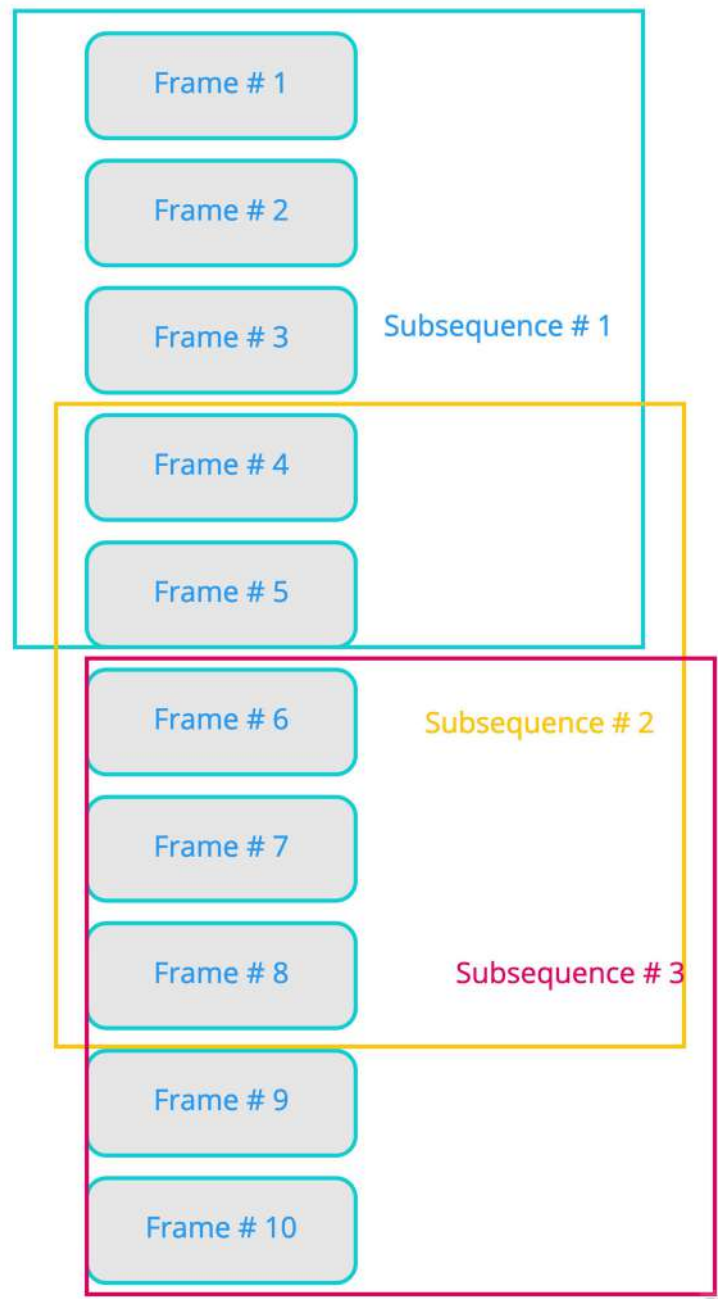


Figure 2: Schema in which one video is divided into spatio-temporal frames with optimal parameters.

There are three steps of dataset preprocessing:

- denoising
- resizing
- normalization

The schema of a pipeline which was used in the research is shown at Figure 3.

The new MobileNetV3 architecture (Fig. 4) is an improvement of its previous versions – MobileNet and MobileNetV2. Previous work used MobileNetV2 as a basis for convolutional neural net architecture, which was afterwards enhanced with 3d convolutions. As a part of previous work development, novel

mobilenetv3 architecture is used in the research, also enhanced with 3d convolutions. Newest mobilenet reincarnation also has two versions – large and small, oriented on high and low hardware resource respectively. Also, a redesign of expensive layers took place, which allowed to further improve performance speed on all platforms.

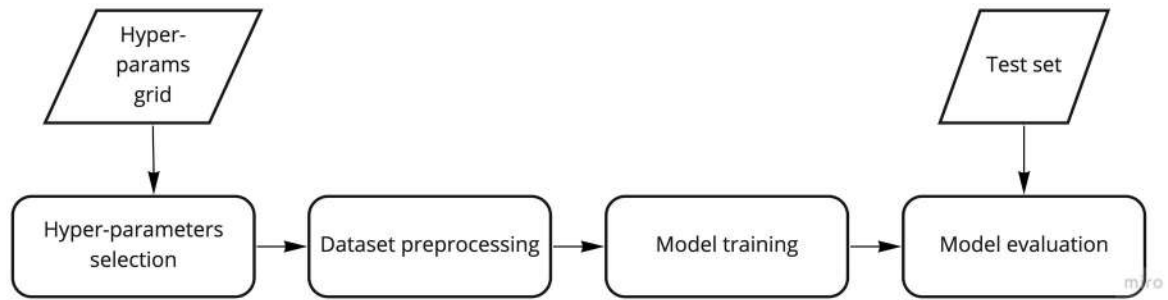


Figure 3: Schema of a pipeline for tuning spatio-temporal frame parameters.

Input	Operator	exp size	#out	SE	NL	s
$224^2 \times 3$	conv2d, 3x3	-	16	-	HS	2
$112^2 \times 16$	bneck, 3x3	16	16	✓	RE	2
$56^2 \times 16$	bneck, 3x3	72	24	-	RE	2
$28^2 \times 24$	bneck, 3x3	88	24	-	RE	1
$28^2 \times 24$	bneck, 5x5	96	40	✓	HS	2
$14^2 \times 40$	bneck, 5x5	240	40	✓	HS	1
$14^2 \times 40$	bneck, 5x5	240	40	✓	HS	1
$14^2 \times 40$	bneck, 5x5	120	48	✓	HS	1
$14^2 \times 48$	bneck, 5x5	144	48	✓	HS	1
$14^2 \times 48$	bneck, 5x5	288	96	✓	HS	2
$7^2 \times 96$	bneck, 5x5	576	96	✓	HS	1
$7^2 \times 96$	bneck, 5x5	576	96	✓	HS	1
$7^2 \times 96$	conv2d, 1x1	-	576	✓	HS	1
$7^2 \times 576$	pool, 7x7	-	-	-	-	1
$1^2 \times 576$	conv2d 1x1, NBN	-	1024	-	HS	1
$1^2 \times 1024$	conv2d 1x1, NBN	-	k	-	-	1

Figure 4: Architecture of MobileNetV3-small

In MobileNetV3 there have been two strategies to improve the network:

- Platform-aware NAS for block-wise search
- NetAdapt for layer-wise search

Introducing a dynamic and temporal component into the technology presented new challenges into the results interpretation. One of the issues which appeared in the process became instability of the prediction on dynamic data. In other words, a sequence of frames produces a sequence of predictions, but there can be wrong predictions, or the prediction could start frantically change from frame to frame.

As a part of improvement of the technology for such cases, an approach for stabilization of predictions was presented and implemented. The are two components of such solution: smoothing of the prediction probabilities and accumulation of probabilities from previous spatio-temporal frames. Also additional anomaly-detection approach is used to neglect predictions within a frame which are highly different from the context.

With such approach, first prediction of the first frame is considered as baseline. Further predictions on next frames start accumulating with the previous prediction, an in case if it's highly different from a history of multiple predictions before – such an anomaly is neglected.

Predictions from earlier subsequences are used to build up the model, which then uses that information to update the current recognition result only when the total number of predictions surpasses a certain threshold.

$$\sum_{t=t-ni=t-k}^{t+n} \sum_{i=t-k}^{t+k} p_i > threshold, \quad (2)$$

where: p_i - the probability of a gesture in the frame; i - frame number on the current subsequence; t - number of the current subsequence; k - the size of the subsequence in both directions; n - number of accumulated subsequences.

5. Dataset of sign images and their augmentations

In previous work a dataset of 50000 images was collected, with different sign, corresponding to Ukrainian sign alphabet (Fig. 5). During creation of such a dataset, it was aimed at to make it diverse in terms of lighting conditions (20 % of data in bad of light conditions, 30 % in mediocre light conditions and 50 % in good quality lighting). Almost 10% of data was is poor quality, with noise and very blurry.



Figure 5: Dataset subsample

In order to increase the dataset size, a list of data augmentation techniques from previous work was enlarged, thus next data augmentation techniques were used:

- rotation
- random cropping
- flipping
- brightness adjustment
- noise addition
- salt and pepper (replaces pixels in images with salt/pepper noise (white/black-ish color))
- coarse dropout - sets rectangular areas within images to zero.
- gamma contrast - adjust image contrast by scaling pixel values to
- affine
- blur
- emboss

As a result, the train dataset was increased in 5 times and a final amount of 250,000 pictures was generated. 20% of the data was used as testing subset for pipeline for spatio-temporal hyperparameter tuning. Some additional data augmentation techniques were used exclusively to the test dataset (of size 50,000 pictures). This was done in order to verify that the technology will stay robust even in unseen before conditions and environment.

Figure 6 shows examples of original images and their augmented counterparts.

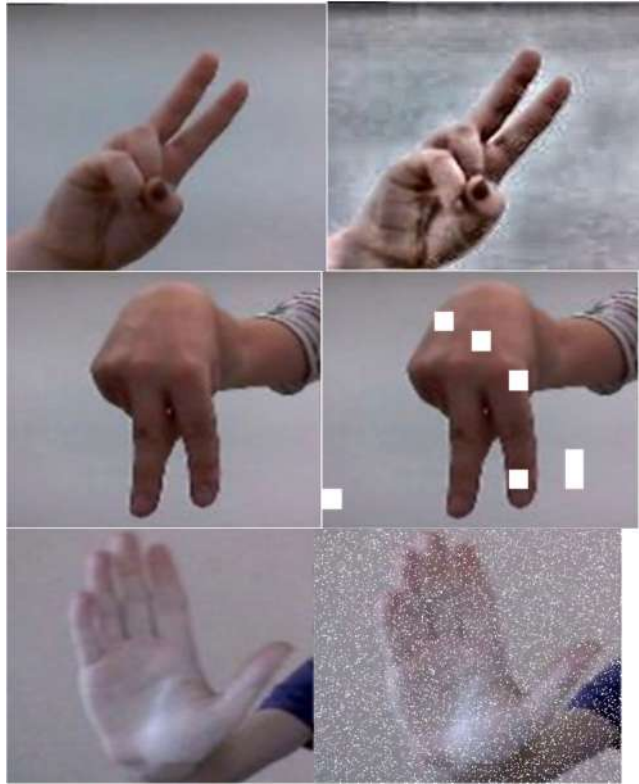


Figure 6: Original images (left) and augmented images (right). Topmost – heavy augmentation (multiple at once), middle – dropout, bottom – salt and pepper.

Figure 7 shows percentage of images with different light conditions in the train and test split datasets, divided into such conditions: bad, mediocre and good.

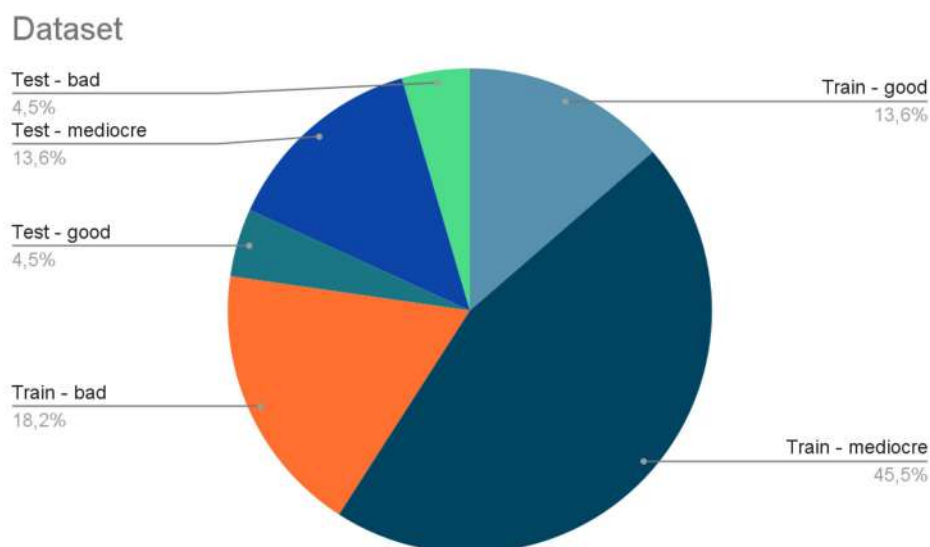


Figure 7: Distribution of light quality on the train and test datasets.

It is also important not to overfit the model with artificial patterns present in the dataset, thus it must maintain statistically significant variety. Also, there should not be major shifts in train dataset comparing to test dataset. All these conditions were met in the train dataset conducted and augmented as a part of research.

6. Experiments

During dataset augmentation of test split with techniques, which were not used in the train split, performance of the model dropped, naturally, due to higher complexity of the testing data. However, still it showed performance comparable with state-of-the-art approaches (for instance, squeeze-nets). Figure 8 shows performance increase in using more novel mobilenetv3 model as a basis, and also increase in performance with three dimensional convolutions. All measurements show macro-averaged f1-score.

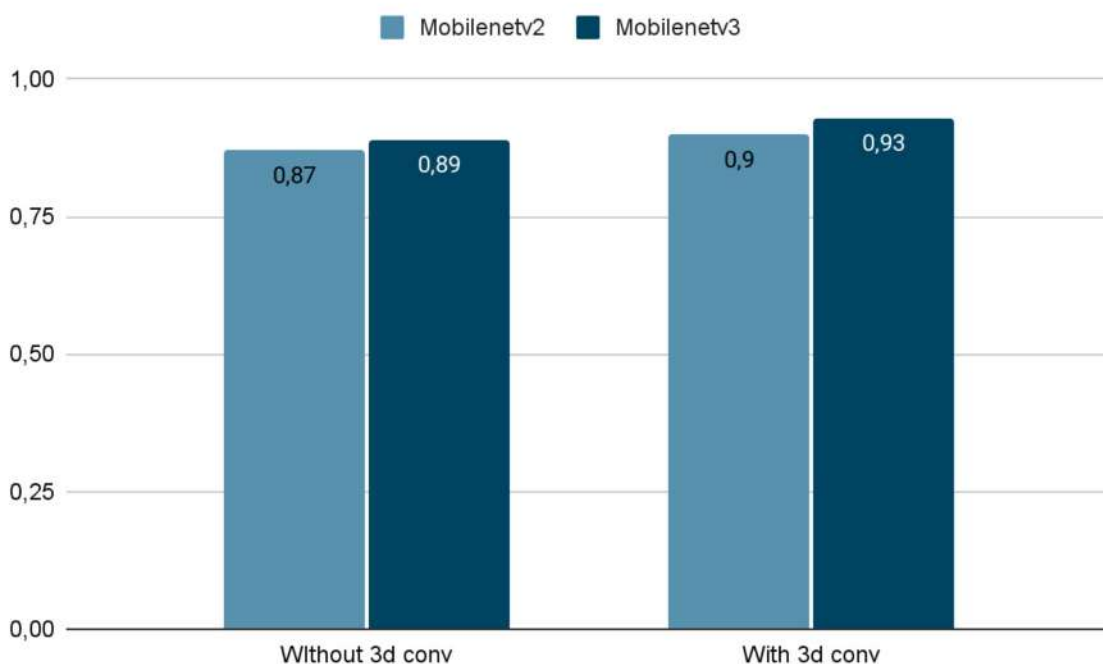


Figure 8: Comparison of models

During models training, hyperparameters of spatio-temporal frame were tuned first. After those parameters fixed, multiple model architectures varieties were considered, with mobilenetv3 as a basis for all of them. Mobilenetv3-small was preferred over mobilenetv3-large due to cross-platform considerations, with ability of such model to perform with relatively same gesture recognition performance but on a lower performance hardware, such as smartphones. Experimental results proved small version to be sufficient for gesture model. At least five different modifications of mobilenetv3-small architecture were tuned and evaluated during the experiments. It is important to analyze confusion matrix of model prediction. This also helped to select the best approach and best configuration. All of these measures allowed to design a balanced neural network that was both small and effective on the test data set.

Hyperparameters form a grid, which contains a configuration for model architecture, training hyperparameters and spatio-temporal hyperparameters.

Example grid:

$$\left\{ \begin{array}{l} \text{learning_rate: [0.001, 0.0001],} \\ \text{batch_size: [8, 16, 32],} \\ \text{layers_config: [config1, config2, config3],} \\ \text{overlapping_frames: 2,} \\ \text{frame_size: 5,} \\ \text{decay: 1} \end{array} \right. \quad (3)$$

Each model train was performed with common techniques for fighting overfitting. Model's prediction time is sufficient for real-time (24 fps) performance using Nvidia K80 GPU.

7. Conclusions

As a result of research, a technology for gesture communication was developed and improved within multiple domains. A novel architecture of mobilenetv3 was used as a lightweight neural network model, which allows both to get a high level of gesture recognition performance with ability to run on a low-performance hardware, such as smartphones. To overcome recognition issues with such a high dynamic object as hand and gesture animation, a three-dimensional convolution technique was used to enhance the neural network architecture. As a solution to use as input videos of different length, a spatio-temporal frame concept was introduced and implemented as a part of research. Having its own hyperparameters, such as number of overlapping frames and size of the frame, it needed configuration. As a part of research, best hyperparameters were tuned, both for spatio-temporal frame and for the neural network itself. A special pipeline was built for tuning hyperparameters for spatio-temporal frame. Additional data augmentation techniques were used to increase the train dataset, also some techniques were used only for test split of dataset, to prove robustness of the model to unseen conditions.

It was proved with experiments to have a consistent improvement in model performance with better data augmentation, novel mobilenetv3 architecture as a basis and spatio-temporal frame approach, which resulted in 0.93 macro-score f1.

8. References

- [1] Kondratiuk S. Gesture recognition using cross platform software and convolutional neural networks // Artificial Intelligence. – b.83-84 – 2019 – p.94-100.
- [2] S. Kondratiuk, I. Krak, A. Kylias, V. Kasianiuk. Fingerspelling Alphabet Recognition using Cnns with 3d Convolutions for Cross Platform Applications. Advances in Intelligent Systems and Computing. Vol. 1246 AISC. 2021, pp.585-596. doi:10.1007/978-3-030-54215-3_37.
- [3] Yu.V.Krak, Yu.V. Barchukova, B.A. Trotsenko. Human hand motion parametrization for dactylemes modeling, Journal of Automation and Information Sciences, 43(12) (2011):1-11. doi:10.1615/JAutomatInfScien.v43.i12.10
- [4] The Linux Information Project, Cross-platform Definition. www.linfo.org
- [5] Y.V. Krak, A.A. Golik, V. S. Kasianiuk. Recognition of dactylemes of Ukrainian sign language based on the geometric characteristics of hand contours defects. Journal of Automation and Information Sciences, 48(4)(2016):90-98. doi:10.1615/JAutomatInfScien.v48.i4.80
- [6] J. Smith, N. Ravi. The Architecture of Virtual Machines. Computer. IEEE Computer Society. 38 (5) (2005): 32–38.
- [7] L. Prasuhn, Y. Oyamada, Y. Mochizuki, H. Ishikawa. A hog-based hand gesture recognition system on a mobile device. In 2014 IEEE International Conference on Image Processing (ICIP), pages 3973-3977. IEEE, 2014.
- [8] N. H. Dardas, N. D. Georganas. Real-time hand gesture detection and recognition using bag-of-features and support vector machine techniques. IEEE Transactions on Instrumentation and measurement, 60(11) (2011): 3592-3607.
- [9] I.V. Krak, G.I. Kudin, A.I. Kulias. Multidimensional Scaling by Means of Pseudoinverse Operations, Cybernetics and Systems Analysis, 55(1) (2019):22-29. doi: 10.1007/s10559-019-00108-9

- [10] O. Kopuklu, N. Kose, G. Rigoll. Motion fused frames: Data level fusion strategy for hand gesture recognition. arXiv preprint arXiv:1804.07187, 2018.
- [11] P. Molchanov, S. Gupta, K. Kim, K. Pulli. Multi-sensor system for driver's hand-gesture recognition. In Automatic Face and Gesture Recognition (FG), 2015 11th IEEE International Conference and Workshops, volume 1, pages 1-8. IEEE, 2015.
- [12] P. Molchanov, S. Gupta, K. Kim, J. Kautz. Hand gesture recognition with 3d convolutional neural networks. In The IEEE Conference on Computer Vision and Pattern Recognition (CVPR) Workshops, pages 1-7. June 2015.
- [13] A. Krizhevsky, I. Sutskever, G.E. Hinton. Imagenet classification with deep convolutional neural networks. In Advances in neural information processing systems, pages 1097-1105, 2012.
- [14] A. Karpathy, G. Toderici, S. Shetty, T. Leung, R. Sukthankar, L. Fei-Fei. Large-scale video classification with convolutional neural networks. In Proceedings of the IEEE conference on Computer Vision and Pattern Recognition, pages 1725-1732, 2014.
- [15] J. Carreira, A. Zisserman. Quovadis, action recognition a new model and the kinetics dataset. In Computer Vision and Pattern Recognition (CVPR), 2017 IEEE Conference, pages 4724-4733. IEEE, 2017.
- [16] T. B. N. GmbH. The 20bn-jester dataset v1. <https://20bn.com/datasets/jester>, 2019.
- [17] F. N. Iandola, S. Han, M. W. Moskewicz, K. Ashraf, W. J. Dally, K. Keutzer. Squeezenet: Alexnet-level accuracy with 50x fewer parameters and 0.5 mb model size. arXiv preprint arXiv:1602.07360, 2016.
- [18] A.G. Howard, M. Zhu, B. Chen, D. Kalenichenko, W. Wang, T. Weyand, M. Andreetto, H. Adam. Mobilenets: Efficient convolutional neural networks for mobile vision applications. arXiv preprint arXiv:1704.04861, 2017.
- [19] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, L.-C. Chen. Mobilenetv2: Inverted residuals and linear bottlenecks. In 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 4510-4520. IEEE, 2018.
- [20] X. Zhang, X. Zhou, M. Lin, J. Sun. Shufflenet: An extremely efficient convolutional neural network for mobile devices. In 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition, pages 6848-6856. IEEE, 2018.
- [21] N. Ma, X. Zhang, H.-T. Zheng, J. Sun. Shufflenet v2: Practical guidelines for efficient cnn architecture design. arXiv preprint arXiv:1807.11164, 5, 2018.
- [22] A. Howard, M. Sandler, G. Chu, L.-C. Chen, B. Chen, M. Tan, W. Wang. Searching for MobileNetV3. arXiv: 1905.02244, 5, 2019
- [23] Eng-Jon Ong et al. Sign language recognition using sequential pattern trees. In: Computer Vision and Pattern Recognition (CVPR), 2012 IEEE Conference on. IEEE. 2012, pp. 2200-2207
- [24] Tensorflow framework documentation [<https://www.tensorflow.org/api/>]
- [25] K. Hara, H. Kataoka, Y. Satoh. Can spatiotemporal 3d cnns retrace the history of 2d cnns and imagenet. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, pages 18-22, 2018.
- [26] J. Hu, L. Shen, G. Sun. Squeeze-and-excitation networks. In Proceedings of the IEEE conference on computer vision and pattern recognition, pages 7132-7141, 2018.
- [27] K. He, X. Zhang, S. Ren, J. Sun. Deep residual learning for image recognition. In Proceedings of the IEEE conference on computer vision and pattern recognition, pages 770-778, 2016.

Neuroevolutionary mechanisms in the synthesis of spiking neural networks

Serhii Leoshchenko^a, Andrii Oliinyk^a, Sergey Subbotin^a, Matviy Ilyashenko^a and Artem Borovikov^a

^a National university "Zaporizhzhia polytechnic", Zhukovskogo street 64, Zaporizhzhia, 69063, Ukraine

Abstract

Since the advent and start of active research on the results of IBM True North, attention to spiking neural networks (SNN) has increased significantly. Such neural networks have even greater potential in the field of artificial intelligence than deep, recurrent and other modern artificial neural network (ANN) architectures. This is easily explained by the ease of their arrangement into neuromorphic systems. However, in most cases, it is difficult to synthesize and train such neuromodels, because classical methods cannot be applied to such complex models. A number of works suggest the use of Composite Pattern Producing Networks. This approach is more similar to another method of ANN synthesis, namely a group of neuroevolution methods that, together with the meta-parameters of the network, evolutionarily modify its structure: Topology and weight Evolving Artificial Neural Network (TWEANN). Therefore, the aim of this work is to investigate the possibility of using neuroevolution methods and its individual mechanisms during SNN synthesis.

Keywords 1

machine learning, artificial intelligence, artificial neural networks, neuroevolution, synthesis, topology, spiking neural network, recurrent neural network, deep neural network, deep learning

1. Introduction

Quite often, the operation of ANN is compared to the work of the brain. However, such networks have taken only the most superficial form and essence from real, living organisms. Due to the too large difference between real neural networks and ANNs, it is necessary to invent different surrogate methods of network training. Naturally, nowhere in nature does not exist something like backpropagation network training, just as there is no unsupervised learning in its pure form [1-4].

SNNs were created with a greater reference to the real work of the brain, and use a method of transmitting information in the likeness of biological neurons (Fig. 1). So, for example, in brain neurons, an impulse is generated at the moment when the current sum of changes in the membrane potential crosses the threshold [5-7]. The pulse occurrence rate and time model of pulse bundle carry information about the external stimulus and ongoing calculations. SNN is based on a similar method of pulse generation and information transmission: neurons use differentiated, nonlinear activation functions, the use of which allows you to create structures with a thickness of more than one layer.

The first SNN model was proposed back in 1952 by Alan Hodgkin and Andrew Huxley [8]. The main feature of such a model is the generation and propagation of action potentials in neurons [8]. After that, with biological refinements and high computational costs, various models of neurons were proposed: Jolivet, Timothy and Gerstner [9]; Izhikevich [10]; Delorme [11]. The latter is very

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: sergio.zntu@gmail.com (S. Leoshchenko); olejnikaa@gmail.com (A. Oliinyk); subbotin@zntu.edu.ua (S. Subbotin);
matviy.ilyashenko@gmail.com (M. Ilyashenko); atemiks@gmail.com (A. Borovikov)
ORCID: 0000-0001-5099-5518 (S. Leoshchenko); 0000-0002-6740-6078 (A. Oliinyk); 0000-0001-5814-8268 (S. Subbotin); 0000-0003-4624-4687 (M. Ilyashenko); 0000-0003-1429-5930 (A. Borovikov)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

popular, as it takes into account the properties of an external stimulus, accumulating the flow of charge through the cell membrane, when crossing a certain threshold.

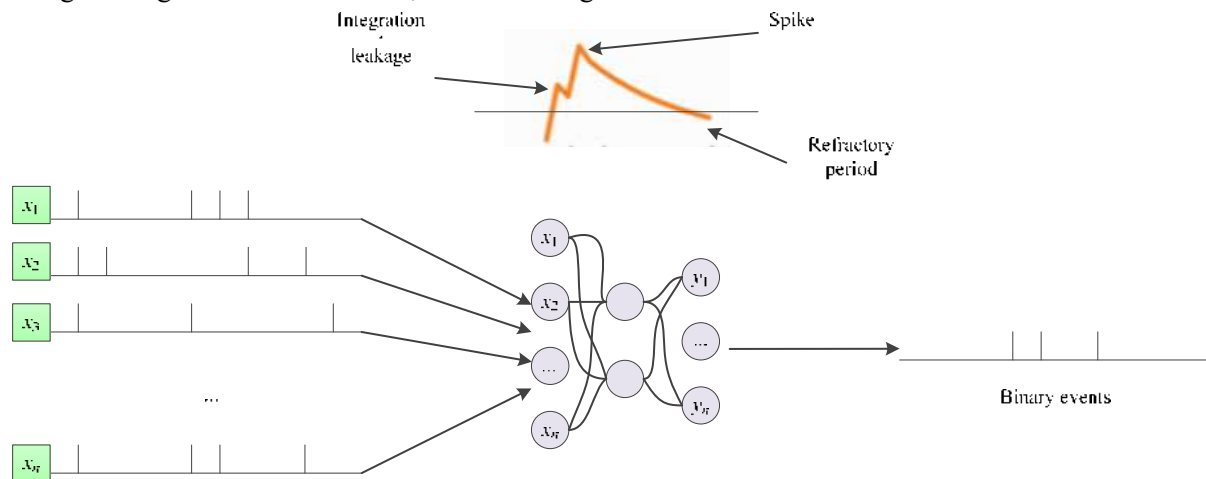


Figure 1: The simple example of SNN architecture

Later, thanks to the research of Kohonen, Grossberg and Anderson, a powerful theoretical foundation was formed, with the help of which it became possible to further develop ANN, namely the design and implementation of multilayer structures [5-7]. But learning is still a huge challenge. Since activation functions are derived, there is room for using gradient optimization methods to train neural networks. With the proliferation of available large labeled data sets for training neural networks, with the increasing computing power of GPUs, and advanced regularization techniques, neural networks are becoming incredibly multi-layered, allowing you to generalize large amounts of invisible data. Layering is a huge advantage in the performance of neural networks [1-3].

It is well known that the brain's ability to recognize complex visual patterns or identify a speaker in a noisy environment is the result of several consecutive processing steps and a variety of learning mechanisms that are also built into SNN with deep learning [8-11]. Compared to ANN and deep learning, SNN learning is at the earliest stages of development. An important feature of this type of system is the neural architecture. It is much better suited for processing space-time data, especially in online mode. The representation of data in time and space that SNNs possess allows such neural networks to perform calculations at the level of the human brain, as well as to understand brain activity in the spatiotemporal structure. It is very important to understand in the coming years how to train such neural networks to perform various tasks [1-7].

In case of looking SNN from an engineering point of view, it can also be found interesting tasks here. This type of neural network has a number of advantages in hardware implementation over conventional ANNs. The pulse bundle in the SNNs is scattered over time, each of them contains a huge amount of information, which can significantly reduce power consumption. Thus, you can create a hardware platform with moderate resource intensity indicators that adapts its work to pulse activity.

It is worth noting that SNNs inspired by the biological example, in principle, recognize images much better and faster than conventional ANNs [12]. Moreover, SNNs allow the use of training methods that depend on the time of occurrence of impulses between pairs of directly connected neurons, in which information for changing the weight of edges is available locally. This training method is very similar to what happens in many parts of the brain.

The resulting pulse bundle is represented as undifferentiated sums of Delta functions. Accordingly, it is difficult to apply derivative-based optimization methods to SNN training, although various types of approximate derivatives have recently been actively studied. Despite the fact that in theory pulse neural networks have equivalent computing power in turing [13], it is still problematic to train SNNs, especially those with a multi-layered architecture. In many existing pulse neural networks, only one layer can be trained. If it is possible to provide pulse systems with multi-layer training, then productivity in various tasks will increase tenfold.

The architecture of pulsed neural networks consists of pulsed neurons and interconnected synapses. Pulse bundle in neural networks of impulse neurons propagate through synaptic

connections. A synapse can be either excitatory, when the membrane potential of a neuron increases after receiving a signal, or slowing down. The amount of rib weight can change as a result of training. Deep learning of multilayer SNNs is a real mystery, since the undifferentiation of pulse bundle does not allow the use of popular methods, such as the backpropagation method (Fig. 2) [14-19].

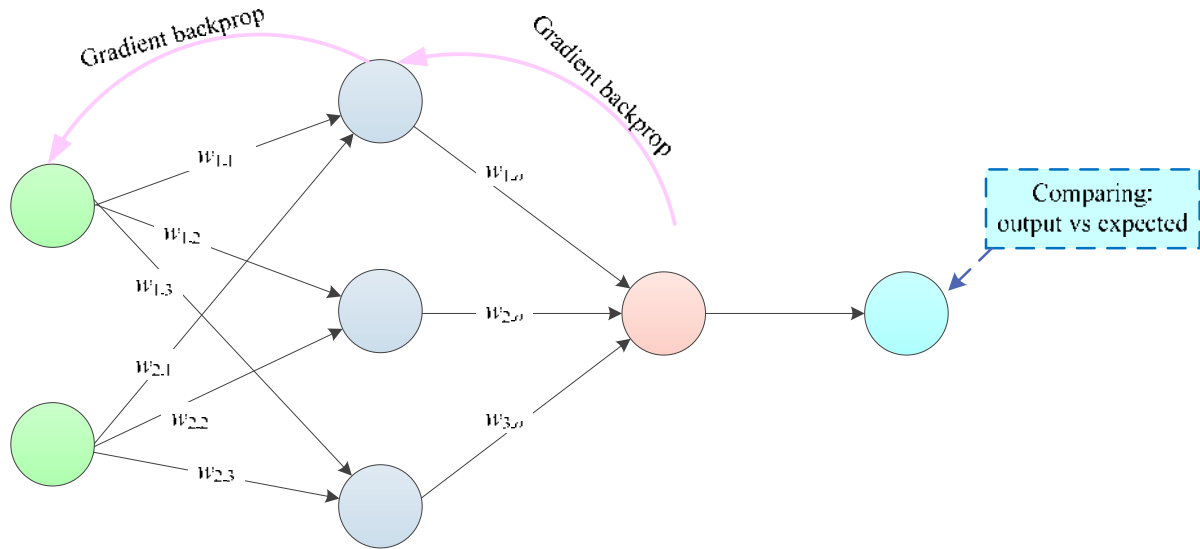


Figure 2: Using the backpropagation method for correcting weights coefficients

2. Related Works

As mentioned earlier, the training of all kinds of artificial neural networks occurs by adjusting scalar synaptic weights. In SNNs, can be used training methods that are close to those used by the brain. Scientists have identified many variants of this training method, but all of them fall under the general term dendritic plasticity (STDP) [20]. A key feature of dendritic plasticity is that the weight of the rib connecting pre- and postsynaptic neurons is regulated with their pulse time in the range of approximately tens of milliseconds in duration.

2.1. Method SpikeProp

This is the first method of training SNNs by error propagation [21], [22]. The cost function takes into account the fluctuation period, and thus this method can classify nonlinearly separated data for a time-encoded XOR problem using a 3-level architecture. The main decision at the stage of development of the method was the choice of the Gerstner neuron model (so-called, SRM) [23]. Using this model, the question of taking derivatives at the oscillation output was circumvented, since the required result was directly modeled as a continuous value. One of the limitations of this method is that each original unit was forced to generate exactly one oscillation. In addition, the values of continuous variables, such as in XOR problems, had to be encoded as delays between fluctuations, which can be quite long.

2.2. Method Remote supervised learning (ReSuMe)

This training method consists of a single oscillating neuron, which receives vibrations from many other oscillating presynaptic neurons [24]. The goal is to train the synapse to trigger a postsynaptic neuron to generate oscillation waves with the desired oscillation period. ReSuMe adapted the delta rule used for non-pulse linearized units to SNN. In the delta rule, the weights change proportionally:

$$\Delta\omega = (y^d - y^{real})x = y^d x - y^{real} x, \quad (1)$$

where x is presynaptic input;

y^d is desired output value;

y^{real} is real value at the output.

By reformulating the equation above, you can get the sum of STDP s anti- STDP:

$$\Delta\omega = \Delta\omega^{STDP}(S^{in}, S^d) + \Delta\omega^{\alpha STDP}(S^{in}, S^{real}). \quad (2)$$

The above $\Delta\omega^{STDP}$ is a function of correlation between presynaptic and desired oscillation periods, where $\Delta\omega^{\alpha STDP}$ is depends on presynaptic and real oscillation periods. Since this method uses a correlation between a set of presynaptic neurons and a wavering neuron, there is no physical connection between them. This is why this method is called remote.

2.3. Method Chronotron

This method was developed based on the Temporal method [25], which could train individual neurons to recognize the encoding at the exact time of arrival of the oscillation. A limitation of the Temporal method was the ability to output 0 or 1 at the output during the selected time interval. Because of this, it was impossible to encode information about the time of arrival of the oscillation in the output data. The creation of Chronotron was influenced by the success of SpikeProp and its successors. The idea of Chronotron was to use more complex units of distance measurement: VP [26] between two vibrations for the training. They adapted the VP distance so that it was piecewise differentiated and suitable as a function of cost to perform gradient descent with respect to weights.

2.4. Problems of existing methods

Among the existing methods, there are several main disadvantages. First, each channel needs a separate quick convergence scheme. That is, there is a risk of multiplying derivatives on each channel, which in the future can impose large computational needs [21-26].

Secondly, it becomes quite difficult to get a large number of channels due to pulse distortions that occur when connecting matching circuits to coaxial cables [21-26].

But the key problem is always the question of topology (the structure of such a network). There are many approaches that apply reverse propagation to SNNs. But such methods cannot use traditional reverse propagation and change gradient updates or SNNs themselves to overcome this complexity. These types of methods often impose architectural constraints on the resulting SNNs, for example, requiring a feedforward connections SNN architecture. Therefore, there are many SNN topologies with desired properties, such as repeatability, that cannot be optimized using these methods [21-26].

That is why it is an urgent task to develop evolutionary methods for the full synthesis of the SNN model. When implementing the methods of the evolutionary approach, SNN population assessment will allow us to further use the processes of stochastic variation and selection to create the next population or control these processes. During successive iterations of mutations and selection of individual networks with increasing fitness, which can be broadly defined as accuracy in the classification problem or maximum reward from the environment, the new solutions will be qualitatively different from the previous ones. One of the significant advantages of evolutionary approaches is their flexibility. As for SNNs, evolutionary optimization can potentially affect any network topology, as well as optimize any network parameter.

3. Proposed method

In general, the method proposed in this paper is similar to a modified genetic algorithm (MGA) for synthesizing ordinary ANNs [27]. So, at the beginning, it should be noted that direct encoding of the structure of individuals is used to perform synthesis, each of which is a separate neural network, but the basis of such encoding is inter-neural connections. This makes it possible to track the origin of each parameter (node or connection) in the genome in more detail. In the subsequent stages of the

method, this greatly simplifies the processes of crossing and mutation, allowing you to build and change genetic information about individuals [27].

The main innovation is the use of certain template mechanisms, namely neural patterns (NP) based on Composite Pattern Producing Networks (CPPN) technology [28]. This mechanism provides abstraction of the processes of natural evolution [28]. NPs will consist of neurons with various activation functions, including periodic functions such as sinusoidal and symmetric functions such as absolute value. The basic idea is based on the fact that a CPPN-based NP with multiple entry points can determine the relationship between a pair of points. A fixed set of neurons is introduced as a specific substrate to the NP, which further simplifies neurosynthesis using MGA.

In general, at the beginning of evolution, we obtain a population with NPs: $P = \{Ind_1, Ind_2, Ind_3, \dots, Ind_n\} = \{NP_1, NP_2, NP_3, \dots, NP_n\}$, that determine the connection patterns of the resulting SNNs $P = \{NP_1, NP_2, NP_3, \dots, NP_n\} = \{SNN_1, SNN_2, SNN_3, \dots, SNN_n\}$. Each NP is defined by encoded genetic information with the definition of its internal weights and activation functions. However, it is important to note that each individual neuron in an NP can have any activation function from a specific set (sin, tanh, gauss, relu, identity): $NP = SNN = \{struct, param\} = \{Node, Connections, weights, Function\}$ [29-31]. To decipher SNN from NP, the coordinates for a pair of neurons are passed to the network that creates the NP. The network output determines the weight and delay of the connection between two neurons in the dive neural network. This process is repeated for each pair of neurons to build a complete network.

After preparing the population, it is possible to perform further synthesis. The main steps for synthesizing and decoding SNN from NP are shown in Fig. 3.

However, several nuances of this synthesis should be noted:

- the positions of SNN neurons, in the form of NP, are pre-recorded during the evolutionary search for each individual. That is, any mutation and modification is possible only in the form of a new individual;
- each solution is decoded and evaluated to complete a given task (classification or management). Therefore, the suitability of the solution is determined by the performance of the generated SNN for the specified task;
- speciation is used to protect innovations in new solutions and can stop if they are not available.

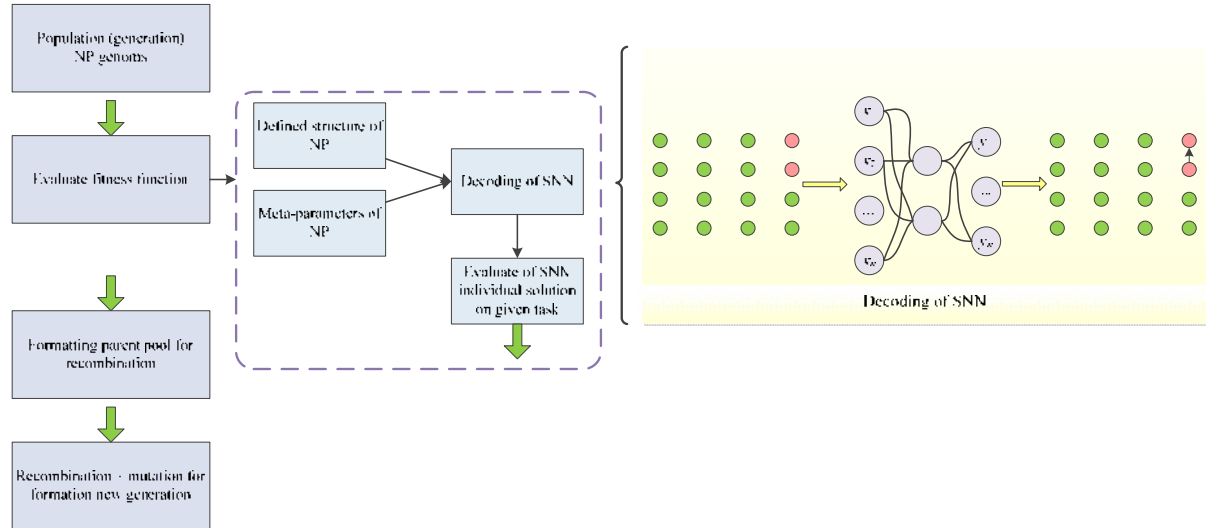


Figure 3: The general scheme of the method

4. Experimental research

To study the results of the method, the task of medical diagnostics was chosen. A data set was selected as input about characteristics of patients with pneumonia, which was recently presented by authors M.-A. Kim, J. Seok Park, C. W. Lee, and W.-I. Choi [32]. Total sample size: 77490 values. Table 1 shows the characteristics of the data set.

For this task, the development of neuromodels will make it much easier to determine the further diagnosis of a person after collecting data on their well-being. Given that pneumonia is one of the most important signs and complications of COVID-19 [33], [34], after additional training on advanced data, this model can be used to diagnose patients or predict the further development of disease dynamics.

The effectiveness of SNN models will be compared with neuromodels based on recurrent neural networks (RNNs) (Fig. 2.a)) and deep neural networks (DNNs) (Fig. 2.b)) [35].

RNN is a type of ANN and is used in natural language processing and speech recognition applications. The RNN model is designed to recognize consistent data characteristics and then use templates to predict the future scenario [35]. A DNN is a neural network with a certain level of complexity, a neural network with more than two layers. DNNs use sophisticated mathematical modeling to process data in complex ways [35].

For the experiments, a workstation with the following characteristics was used: Intel Core i5-12600 central processor (3.30-4.80 GHz (Intel Turbo Boost 2.0), 6 cores and 12 threads), 16 Gb RAM (dual-channel mode), SK hynix SC308 128 GB SSD (M.2), the Python programming language.

Table 2 shows the results of synthesis methods with different typologies of neural networks.

Table 1

General characteristics of the data set

Total number of values	77490	Number of attributes	54
The type of the data	Numeric (after consideration of)	Number of instances'	1435

Table 2

General results on the data set

Method of synthesis	Synthesis Time, s	Error at the training sample	Error at the test sample
MGA SNN	9352	0.01	0.137
MGA RNN	6793	0.018	0.148
MGA DNN	7625	0.014	0.138

5. Discussions of results

From the analysis of the obtained experimental results, we can conclude that the proposal to use neuroevolution methods for SNNs synthesis shows good results for its further deployment and use. After all, for networks of this class, the issue of a complex and sometimes not fully clear structure is always relevant. The use of neuroevolutionary approaches with NP encoding and decoding greatly facilitates this task, because in this case structural synthesis is reduced to well-known approaches to neuroevolutionary synthesis using TWEANNs. Parametric synthesis, on the other hand, is fine-tuning at exclusively defined stages without having to iteratively calculate fitness functions or its derivatives over and over again.

On the other hand, we can conclude that the practical use of SNN-class neural networks is still extremely promising, because the results obtained for the accuracy of work cannot fully demonstrate the feasibility of significant time costs. Thus, the values of errors in both cases (training and testing) do not differ in significance, which could cover the time spent on the synthesis of such networks and their subsequent calculations. After all, it is worth noting that in the absence of specialized equipment

for emulating the operation of SNN, the execution of this process on a regular workstation is also accompanied by time overlays.

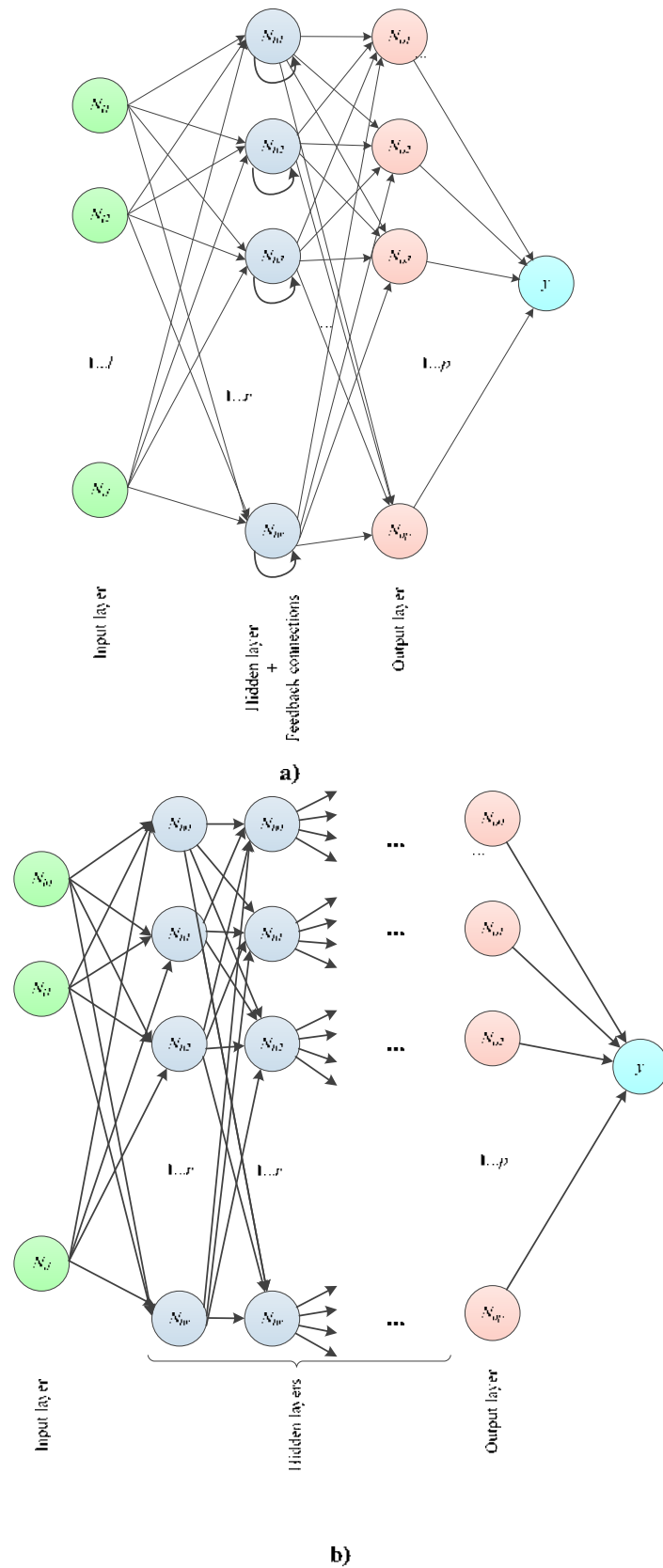


Figure 4: The comparing of RNN and DNN topologies

It is also worth analyzing the work with complex topologies of classical ANNs separately. Thus, the differences in synthesis time are explained by a simpler calculation and correction of recurrent connections in contrast to the calculation of nodes and weights of interneuronal connections in a set of hidden layers. And the analysis of the accuracy of the synthesized solutions based on training and test data shows that the difference in accuracy is not critical and is acceptable within the task. That is why we can talk about a specific indication, about the use of a particular ANN topology, depending on which of the resources is more important: the time of synthesis and operation of the neuromodel or its accuracy.

6. Conclusion

This paper presents an approach to neuroevolution synthesis of complex ANN topologies, namely SNN. The aim of the work was to demonstrate the feasibility and effectiveness of the method. The experiments shown demonstrate the strengths and weaknesses of the approach, as well as identify ways for future research.

One of the main problems that is tracked from the very beginning is the sensitivity to mutations of rigidly defined neuromodel structures. Because mutation of even one structural unit (for example, between neural connections) leads to the creation of a completely new individual, which can sometimes be perceived as undesirable computational difficulties.

However, the main theoretical advantage of this approach to SNN synthesis is the possibility of controlled network development to a large size. In this work, only networks with fewer than a hundred neurons were tested. Therefore, network scaling properties are a further area of research.

Another important research perspective is the introduction and use of methods for indirect encoding of NP and ANN individuals in general during synthesis. Such approaches can help to study in more detail the many variants of relationships between unique SNN characteristics. For example, different input encoding strategies can lead to variability in the configuration of neurons for NP.

Future work may use more of the relationships between neuroevolution and SNNs and extend this method to more complex tasks.

7. Acknowledgements

The work was carried out with the support of the state budget research projects of the state budget of the National University "Zaporozhzhia Polytechnic" "Development of methods and tools for analysis and prediction of dynamic behavior of nonlinear objects" (state registration number 0121U107499) and "Intelligent methods and tools for diagnosing and predicting the state of complex objects" (state registration number 0122U000972).

8. References

- [1] D. Ping, The Machine Learning Solutions Architect Handbook: Create machine learning platforms to run solutions in an enterprise setting, Packt Publishing, Birmingham, 2022.
- [2] A. Burkov, Machine Learning Engineering, True Positive Inc., Quebec City, 2020.
- [3] A. Park, Machine Learning: 2 Books in 1 – The Complete Guide for Beginners to Master Neural Networks, Artificial Intelligence, and Data Science with Python, 2nd. ed., Independently published, 2022.
- [4] S. Alkhalifa, Machine Learning in Biotechnology and Life Sciences: Build machine learning models using Python and deploy them on the cloud, Packt Publishing, Birmingham, 2022.
- [5] C.C. Aggarwal, Neural Networks and Deep Learning: A Textbook, Springer, Berlin, 2018.
- [6] P.R. Hiesinger, The Self-Assembling Brain: How Neural Networks Grow Smarter, Princeton University Press, Princeton, 2021.
- [7] D. Graupe, Principles of Artificial Neural Networks: Basic Designs to Deep Learning (4th Edition) (Advanced Circuits and Systems), 4th ed., WSPC, Singapore, 2019.

- [8] A.Hodgkin, A. Huxley, A quantitative description of membrane current and its application to conduction and excitation in nerve, *J. Physiol*, vol. 117(4), (1952) 500-544. doi: 10.1113/jphysiol.1952.sp004764
- [9] R. Jolivet, J. Timothy, W. Gerstner, The Spike Response Model: A Framework to Predict Neuronal Spike Trains, in: *Proceedings of the 2003 joint international conference on Artificial neural networks and neural information processing, ICANN/ICONIP'03*, ACM Digital Library, 2003, pp. 846-853. doi: 10.1007/3-540-44989-2_101
- [10] E. Izhikevich, Simple model of spiking neurons, *Transactions of neural networks*, vol. 14, (2003) 1569-1572. doi: 10.1109/TNN.2003.820440
- [11] A. Delorme et al., SpikeNET: A simulator for modeling large networks of integrate and fire neurons, *Neurocomputing*, vol. 26(7), (1999) 989-996. doi: 10.1016/S0925-2312(99)00095-8
- [12] M. Pfeiffer, T. Pfeil, Deep learning with spiking neurons: Opportunities and challenges, *Frontiers in Neuroscience*, vol. 12: 774 (2018). doi: 10.3389/fnins.2018.00774
- [13] W. Maass, Lower bounds for the computational power of networks of spiking neurons, *Neural Computation*, vol. 8, (1996) 1-40.
- [14] J.A.J. Alsayaydeh, A. Aziz, A.I.A. Rahman, S.N.S. Salim, M. Zainon, Z.A. Baharudin, M.I. Abbasi, A.W.Y. Khang, Development of programmable home security using GSM system for early prevention, vol. 16(1) (2021) 88-97.
- [15] J.A.J. Alsayaydeh, W.A.Y. Khang, W.A. Indra, V. Shkarupylo, J. Jayasundar, Development of smart dustbin by using apps, *ARPN Journal of Engineering and Applied Sciences*, vol. 14(21) (2019) 3703-3711.
- [16] J.A.J. Alsayaydeh, W.A.Y. Khang, W.A. Indra, J.B. Pusppanathan, V. Shkarupylo, A.K.M. Zakir Hossain, S. Saravanan, Development of vehicle door security using smart tag and fingerprint system, *ARPN Journal of Engineering and Applied Sciences*, vol. 9(1) (2019) 3108-3114.
- [17] J.A.J. Alsayaydeh, W.A.Y. Khang, A.K.M.Z Hossain, V. Shkarupylo, J. Pusppanathan, The experimental studies of the automatic control methods of magnetic separators performance by magnetic product, *ARPN Journal of Engineering and Applied Sciences*, vol. 15(7) (2020) 922-927.
- [18] J.A.J. Alsayaydeh, W.A. Indra, W.A.Y. Khang, V. Shkarupylo, D.A.P.P. Jkatisan, Development of vehicle ignition using fingerprint, *ARPN Journal of Engineering and Applied Sciences*, vol. 14(23) (2019) 4045-4053.
- [19] J.A. Alsayaydeh, M. Nj, S.N. Syed, A.W. Yoon, W.A. Indra, V. Shkarupylo, C. Pellipus, Homes appliances control using bluetooth, *ARPN Journal of Engineering and Applied Sciences*, vol. 14(19) (2019) 3344-3357.
- [20] S. Tazerart, D. E. Mitchell, S. Miranda-Rottmann, R. Araya, A spike-timing-dependent plasticity rule for dendritic spines, *Nature Communications*, vol. 11 (2020). doi: 10.1038/s41467-020-17861-7
- [21] R.V. Florian, Supervised Learning in Spiking Neural Networks. In: Seel N.M. (eds) *Encyclopedia of the Sciences of Learning*. Springer, Boston, MA, 2012. doi: 10.1007/978-1-4419-1428-6_1714
- [22] S. M. Bohte, Error-backpropagation in temporally encoded networks of spiking neurons, in: *Proceedings of the Artificial Neural Network and Machine Learning, ICANN 2011*, ACM Digital Library, 2011, pp. 60-68.
- [23] W. Gerstner, W. M. Kistler, *Spiking neuron models: Single neurons, populations, plasticity*, Cambridge University Press, Cambridge, 2002.
- [24] F. Ponulak, A. Kasinski, Supervised learning in spiking neural networks with ReSuMe: Sequence learning, classification, and spike shifting, *Neural Computation*, vol. 22, (2009) pp. 467-510. doi: 10.1162/neco.2009.11-08-901
- [25] R. Gütig, H. Sompolinsky, The tempotron: A neuron that learns spike timing-based decisions, *Nature Neuroscience*, vol. 9, (2006) 420-428. doi: 10.1038/nn1643
- [26] T. Serre, Hierarchical models of the visual system, *Encyclopedia of Computation Neuroscience*, Springer, New York, NY, 2014. doi: 10.1007/978-1-4614-7320-6_345-1
- [27] S. Leoshchenko, A. Oliinyk, S. Subbotin, N. Gorobii, T. Zaiko, Synthesis of artificial neural networks using a modified genetic algorithm, in: *Proceedings of the 1st International Workshop on Informatics & Data-Driven Medicine, IDDM 2018, CEUR-WS*, 2018, pp. 1-13

- [28] D. Elbrecht, C. Schuman, Neuroevolution of Spiking Neural Networks Using Compositional Pattern Producing Networks, in: Proceedings of the ICONS 2020: International Conference on Neuromorphic Systems 2020, ICONS 2020, ACM Digital Library, 2020, pp. 1-5. doi: 10.1145/3407197.3407198
- [29] A.Oliinyk, S. Skrupsky, S. Subbotin, Experimental research and analysis of complexity of parallel method for production rules extraction. Automatic Control and Computer Sciences, 52 (2) (2018) 89-99. doi: 10.3103/S0146411618020062
- [30] S. Leoshchenko, A. Oliinyk, S. Skrupsky, S. Subbotin, T. Zaiko, Parallel Method of Neural Network Synthesis Based on a Modified Genetic Algorithm Application, in: Proceedings of the 8th International Conference on Mathematics. Information Technologies. Education, MoMLeT&DS-2019, CEUR-WS, 2019, pp. 11–23.
- [31] A.Oliinyk, S. Leoshchenko, V. Lovkin, S. Subbotin, T. Zaiko, Parallel data reduction method for complex technical objects and processes, in: Proceedings of the 2018 IEEE 9th International Conference on Dependable Systems, Services and Technologies, DESSERT-2018, IEEE, 2018, pp. 496-501. doi: 10.1109/DESSERT.2018.8409184
- [32] M.-A. Kim, J. Seok Park, C. W. Lee, W.-I. Choi, Pneumonia severity index in viral community acquired pneumonia in adults, PLoS One, vol. 14(3) (2019). doi: 10.1371/journal.pone.0210102
- [33] G. Raghu, K. C Wilson, COVID-19 interstitial pneumonia: monitoring the clinical course in survivors, The Lancet Respiratory Medicine, vol. 8(9) (2020) 839-842. doi: 10.1016/S2213-2600(20)30349-0
- [34] A. Hani, N.H. Trieu, I. Saab, S. Dangeard, S. Bennani, G.Chassagnon, M.-P. Revel, COVID-19 pneumonia: A review of typical CT findings and differential diagnosis, Diagnostic and Interventional Imaging, vol. 101(5) (2020), 263-268. doi: 10.1016/j.diii.2020.03.014
- [35] S. Leoshchenko, A. Oliinyk, S. Subbotin, T. Zaiko, A Using modern architectures of recurrent neural networks for technical diagnosis of complex systems. Proceedings of the 5th International Scientific-Practical Conference Problems of Infocommunications. Science and Technology (PICST2018), IEEE, Kharkiv, Ukraine, 2018, pp. 411–416. doi: 10.1109/INFOCOMMST.2018.8632015

Method of Forming the Context of Advertising and Target Audience based on Associative Rules Learning

Khrystyna Lipianina-Honcharenko ¹, Taras Lendiuk ¹, Anatoliy Sachenko ^{1,2}, Jacek Wołoszyn ²

¹ West Ukrainian National University, Lvivska Str., 11, Ternopil, 46000, Ukraine

² Kazimierz Pulaski University of Technology and Humanities in Radom, ul. Malczewskiego 29, 26-600 Radom, Poland

Abstract

Nowadays, important mechanisms of study are content and techniques of its creation, the problem of influencing the target audience, which itself seeks to shape communication processes. Internet content occupies a position of powerful communication technology, which continues to grow rapidly and gain influence. Creating a large number of advertisements, especially texts, is extremely expensive. Therefore, it is worth considering how generate these texts automatically. In this regard, it is possible to assume that the development of a method of forming the context of advertising and target audience based on learning associative rules is relevant and can increase the effectiveness of advertising, and thus reduce the cost of online advertising of higher education institutions. The input data used a survey of students majoring in Computer Science, regarding admission. The 152 students took part in the survey and answered 10 questions. The experimental results confirmed, the proposed method enabled to increase the effectiveness of advertising on social networks at least in 23%, and reduce the price in 90%.

Keywords

Data analysis, advertising content, Associative Rules Learning, Apriori algorithm, Facebook.

1. Introduction

The value of advertising is crucial for a company, because only this can make people aware of the company's product and, doing so, can create a good opportunity to sell it to customers [2]. The manual work of a large number of advertisements, especially texts, is extremely expensive. Therefore, it is worth considering how to generate these texts automatically.

Therefore, we can assume that the development of a method of forming the context of advertising and target audience based on learning associative rules is relevant and can increase the effectiveness of advertising, and thus reduce the cost of online advertising of higher education institutions.

This work is devoted to this topic and it is distributed as follows. Section 2 discusses the analysis of related work; section 3 presents a method of forming the context of advertising and the target audience based on the associative rules learning. Section 4 presents the implementation of the method. Section 5 presents the conclusions of the study.

2. Related Work

Given the complexity of the modern digital advertising ecosystem, there are many studies describing the impact of advertising content on social networks on attracting customers [19], using data from Facebook in: medicine [20], psychology [21], sociology [22], politics [23] and others. Research [24,

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL: xrustya.com@gmail.com (Kh. Lipianina-Goncharenko); tl@wunu.edu.ua (T. Lendiuk); as@wunu.edu.ua (A. Sachenko);
jacek.woloszyn@uthrad.pl (J. Wołoszyn)

ORCID: 0000-0002-2441-6292 (Kh. Lipianina-Goncharenko); 0000-0001-9484-8333 (T. Lendiuk); 0000-0002-0907-3682 (A. Sachenko)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

25] provides an understanding of social networks for higher education institutions such as Instagram, Pinterest, Snapchat and WhatsApp.

In [13], a model is proposed that is able to use its targeting strategy in accordance with the received feedback. The model uses the Thompson Sampling algorithm applied to the user space. [1] a set of models of regression, cluster analysis and analysis of association rules to find patterns of user behavior in relation to marketing campaigns, taking into account user characteristics and financially significant indicators.

The aim of the article [12] is to study the analysis of social media data using machine learning tools; new approach to social media marketing strategy uses the Waikato Knowledge Analysis Environment (WEKA). In [5] the article implemented the level of the aspect of mood analysis, based on machine learning algorithms of SVM and NB classification.

In [17] the study analyzed the various needs of customers, focusing on online advertising based on methods of classification, segmentation and clustering. In [16] the model of selection of the optimal amount of advertising on various Internet resources is analyzed in order to achieve the desired coverage of the target audience. In addition, the method of multi-criteria optimization with the definition of the obtained objective function is considered, which allows considering simultaneously the various aspects of media selection issue and optimal budgeting and budget allocation. The proposed approach [2] draws on knowledge that can support several solutions, ranging from marketing campaigns for each customer segment, redesign of the store layout to product recommendations. In [15] the development of advertising art design on the basis of information technologies is mainly investigated.

In a study [4, 26], it is proposed to apply the training of association rules to find influential bloggers in time using the Apriori algorithm. An improved Apriori algorithm has been proposed to identify the relationship between the TECP and the thirty-five factors, which cover four categories of household characteristics, including housing, socio-demographic, household appliances and heating, and energy attitudes [9]. In [10] a new effective system of recommendations based on the Apriori algorithm for user requirements was proposed.

Article [18] proposes the generation of advertising texts based on keywords that take into account product information. Ins [3] was developed a web-based system of recommendations for choosing a property using the method of content-based filtering. The referral system provides information about properties based on user behavior by searching for advertising content that the user previously searched for. [11] presents an intelligent management system for advertising on social networks, based on data analysis techniques to automatically create ads.

The above-mentioned works mostly analyze the actions of users on online advertising. There are also a number of works that present research on methods of associative rules learning. There are also works that use different approaches to the formation of content and target audience of Internet advertising (analogues).

Therefore, goal of this paper is to develop a method of forming the context of advertising and the target audience based on the associative rules learning.

The novelty of the work is the formation of the most profitable (in the economic aspect) text of the Free Economic Zone advertising and the corresponding target group, which will increase the effectiveness of the advertising campaign based on learning the rules of association.

3. Proposed Method

To reduce the time spent on the formation of advertising content for the target audience, the authors have developed a method of forming the context of advertising and the target audience based on the associative rules learning. The proposed method is illustrated schematically (Fig. 1) and is represented by the following steps:

Step 1. Conduct a survey of students (Block 1). It is important to determine the gender characteristics of the respondents, as this will make it possible to determine the target audience in the future. Convert to csv format (Block 2).

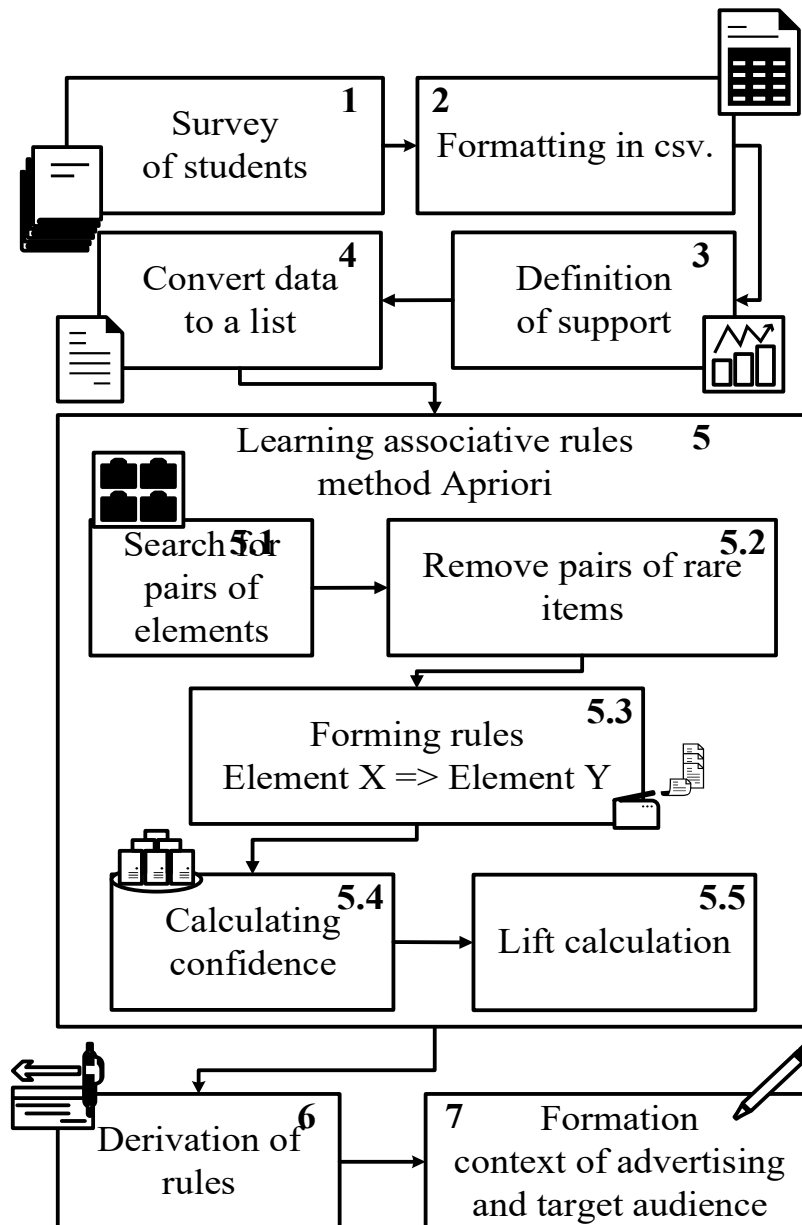


Figure 1: Algorithmic structure of forming the context of advertising and target audience based on learning associative rules

Step 2. Calculation of support for each individual element (Block 3). Support is simply the number of transactions during which a particular product (or combination of products) occurs.

Step 3. Conversion of data into a list (Block 4).

Step 4. Start (Block 5) learning associative rules based on the Apriori method [6].

Step 4.1. Finding support for frequent sets of elements (Block 5.1). Search for pairs of items that appear most often in pairs. The Apriori algorithm ignores all pairs that contain any of the rare elements (Block 5.2).

Step 4.2. Formation of rules (Block 5.3). The most frequent sets of elements, converted into association rules, in the format: Element X => Element Y.

Step 4.3. Calculation of confidence (Block 5.4). Confidence shows the percentage of cases in which this rule applies.

Step 4.3. Elevator calculation (Block 5.5). Raising a rule is an indicator of effectiveness, which indicates the strength of the relationship between the products in the rule. Raising the rule is determined by the following formula:

$$lift = \frac{P(X \cap Y)}{P(X) \times P(Y)}$$

where, P is the probability of the frequency of combination of elements in the generated rule.

Step 5. Output of results-rules (Block 6).

Step 6. Creating the context of advertising and target audience (Block 7), for higher education institutions on the basis of the received rules.

4. Experimental Results

The Python language was chosen to conduct the method of forming the context of advertising and the target audience based on the associative rules learning.

The input data used a survey of students majoring in Computer Science, regarding admission. 152 students took part in the survey and answered 10 questions. All students' feedbacks are in .csv format. The sample is representative, as the survey was conducted among students majoring in Computer Science, who themselves have recently been faced with the choice of where to enter, so their feedback is the most informative for this type of advertising.

Before starting the analysis of the rules of the association, first determine the frequency distribution of the elements (Table 2). The chart shows that the majority of male respondents, as well as the highest number of answers, said that they found information on their own in social networks, majoring in Computer Science.

Table 1

Response rate

Item name	Count
Male	67
independently found information	59
I saw a lot of interesting information on social networks	41
advised familiar relatives	39
Female	37
Interest in computers	35
The presentation of the specialty by the representatives was interesting	33
Interest in design	32
Interest in technology	30
Passed here for public training	29
Augmented reality	29
I got a call and was convinced that it would be interesting to study here	29
Representatives of the specialty came to our school	20
Internet of Things	17
Robotics	15
Internet of Things. Augmented reality	12
Software control of drones	11
The cost of training suited	11
From social networks	10
Due to quarantine, I decided to choose a place closer to home	8
Interest in IT	6
Internet of Things. Augmented reality. Robotics. Software control of drones	5
Augmented reality. Robotics.	4
Internet of Things. Robotics.	4
Augmented reality. Software control of drones	4
Augmented reality. Robotics. Software control of drones	2
Internet of Things. Augmented reality. Robotics.	2

Item name	Count
Internet of Things. Robotics. Software control of drones	2
Quality presentation of the material	2
A friend studied in the same specialty	1
A friend recommended	1
Web	1
I wanted to study to be a programmer	1
Because a large variety of specialties could not be determined	1
Robotics. Software control of drones	1
Chose occasionally	1
I like working with computers and I am interested in the specialty	1
Prospect	1
My chosen specialty is the most relevant	1
Because this specialty covers many areas that I really like	1
Many acquaintances study in Ternopil	1
A relative studied at the university	1
I have wanted to become a programmer for a long time, so I chose this specialty	1
I found information from social networks on my own	1
I really wanted to enroll in FCIT and the most advanced training is on CS	1
I was interested in computer science	1

In addition, it is important to note that even the most frequent response of more than 11% is male (Table 2). Next, we will use this information as a guide when setting a minimum support threshold.

Table 2

The most common answers

Index	Item	Count	Percentage
13	augmented reality	29	4.93
40	passed here for public training	29	4.93
47	interest in technology	30	5.10
45	interest in design	32	5.44
37	interested in the presentation of the specialty by representatives	34	5.78
46	interest in computers	35	5.95
34	female	37	6.29
39	advised familiar relatives	39	6.63
31	saw a lot of interesting information on social networks	41	6.97
41	independently found information	59	10.03
43	male	67	11.39

After converting the dataset into the desired list, we will display the results, namely the rules, which in the future will allow to form the context of advertising and target groups.

Therefore, after starting the algorithm was generated:

1. Counting sets of items of length 1:
 - 48 candidates were found for sets of length 1;
 - Found 15 large sets of items of length 1.
2. Counting sets of items of length 2:
 - 105 candidates were found for sets of length 2;
 - 32 large sets of items of length 2 were found.

Based on the experimentally significant parameters of the algorithm, the generated rules are filtered. Parameters:

min_support = 0.11 – defined as the percentage of most frequent responses;
min_confidence = 0.65 – this probability of determining the rules is sufficient and in experimental studies showed quite good results for this sample.

Rules:

- {male, interest in design} -> {augmented reality} (conf: 0.737, supp: 0.135, lift: 1.965, conv: 2.375);
- {advised by familiar relatives, passed here for public education} -> {male} (conf: 1.000, supp: 0.115, lift: 1.552, conv: 355769230.769);
- {saw a lot of interesting information on social networks, advised familiar relatives} -> {male} (conf: 0.857, supp: 0.115, lift: 1.330, conv: 2.490);
- {saw a lot of interesting information on social networks, interest in technology} -> {found information on his own} (conf: 0.706, supp: 0.115, lift: 1.244, conv: 1.471);
- {independently found information, interest in computers} -> {male} (conf: 0.800, supp: 0.154, lift: 1.242, conv: 1.779);
- {interest in technology} -> {independently found information} (conf: 0.700, supp: 0.202, lift: 1.234, conv: 1.442);
- {passed here for public training} -> {male} (conf: 0.793, supp: 0.221, lift: 1.231, conv: 1.720);
- {advised familiar relatives, interest in design} -> {male} (conf: 0.778, supp: 0.135, lift: 1.207, conv: 1.601);
- {interest in computers} -> {male} (conf: 0.771, supp: 0.260, lift: 1.197, conv: 1.556);
- {Augmented Reality} -> {male} (conf: 0.690, supp: 0.192, lift: 1.071, conv: 1.146);
- {saw a lot of interesting information on social networks} -> {male} (conf: 0.659, supp: 0.260, lift: 1.022, conv: 1.042).

Almost all of the received rules mention the gender of the respondent (male), which indicates the main target audience for the specialty “Computer Science”. Rules that are more than two elements of the answer, allow to create content for advertising. Let's form some examples of content (Table 3).

Table 3

Formation of advertising content in relation to the generated rules

# of variant	Rule	Content
1	{male, interest in design} -> {Augmented reality}	Interested in Design, Try Yourself in Augmented Reality
2	{advised by familiar relatives, passed here for public education} -> {male}	We Are Recommended When Public Education Is Important
3	{saw a lot of interesting information on social networks, advised familiar relatives} -> {male}	We are recommended after browsing our social networks
4	{saw a lot of interesting information on social networks, interest in technology} -> {male}	Interested in technology, visit our social networking pages, there's lots of interesting information
5	{self-found information, interest in computers} -> {male}	Interested in Technology, Visit our Social Networking Pages

To compare the effectiveness of the generated advertising content based on the associative rules learning, a comparative experiment was conducted on Facebook on the business page “Computer Science of ZUNU”. The first option (Option 0) advertising (Fig. 2), developed on the basis of the rules identified in previous research, namely:

- The greatest interaction with the video advertising of Facebook “Computer Science ZUNU” had males in the age category 18-25, 35-55 [7];
- Male and female clients in the age category of 40-55 had the greatest interaction with the ZUNU Computer Science business page.



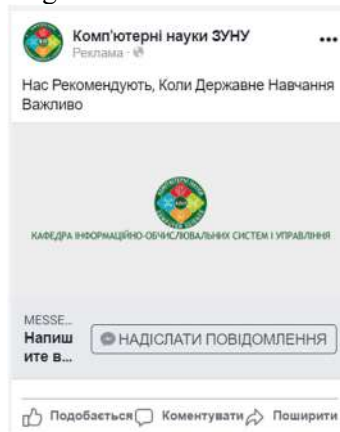
Variant 0

Figure 2: Previous version of the advertisement “Computer Science” of the Western Ukrainian National University on Facebook

Fig. 3 presents an advertisement formed on the basis of learning associative rules, a comparative experiment was conducted on Facebook on the business page “Computer Science of ZUNU”. The content used is from Table 3 and the target audience is male for all age groups.



Variant 1



Variant 2



Variant 3



Variant 4



Variant 5

Figure 3: New variants of advertising "Computer Science" of the Western Ukrainian National University on Facebook

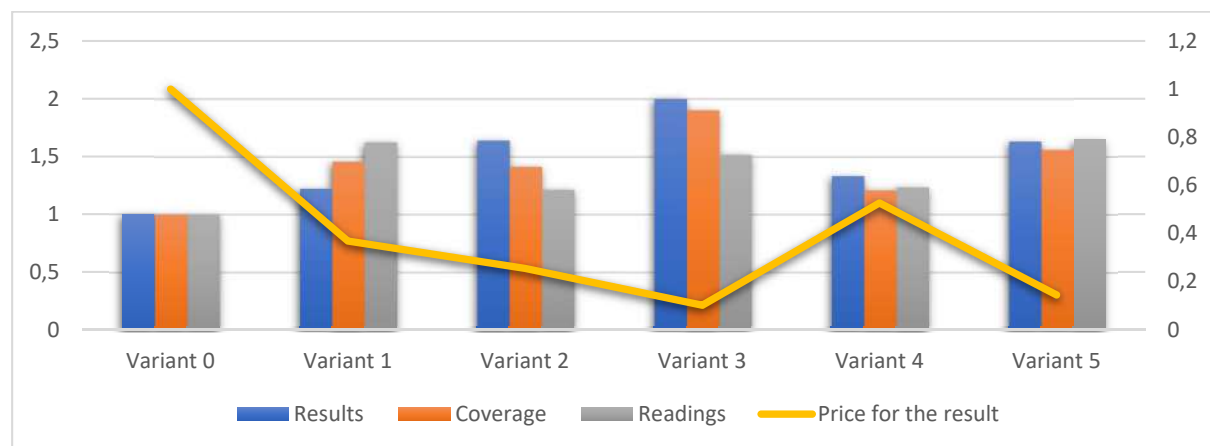
Table 4 presents a comparison of the effectiveness of the generated advertising content based on the associative rules learning, in the period from May 4, 2021 – May 31, 2021 with all content options, including the old.

Table 4

Comparison of the effectiveness of the generated advertising content

Advertising variant	Results		Coverage		Readings		Price for the result	
	Index	Changes	Index	Changes	Index	Changes	Index	Changes
Variant 0	120	100%	4498	100%	5395	100%	0,23	100%
Variant 1	147	123%	6561	146%	8776	163%	0,08	37%
Variant 2	197	164%	6364	141%	6560	122%	0,06	25%
Variant 3	240	200%	8561	190%	8192	152%	0,02	10%
Variant 4	160	133%	5442	121%	6671	124%	0,12	53%
Variant 5	196	163%	7024	156%	8867	164%	0,03	14%

Table 4 shows that all variants of the new ad gave improved results. The result indicator shows how many times customers have been in contact with the ad. Advertising performed best in Option 3 (Fig. 4), as it performed 100% better than Option 0. Option 3 performed best in all indicators. Namely, the coverage is 90% better than option 0 and 52% better in terms of the number of impressions. It also reduced the price for the result by 90% in option 3.

**Figure 4:** Comparison of the effectiveness of the generated advertising content

Thus, the proposed method enabled to increase the effectiveness of advertising on social networks at least in 23%, and reduce the price in 90%. Moreover, the authors believe that growing the number of student surveys can increase the quality of associative rules learning as well as get better keywords for content formation. That will lead to increasing the effectiveness of advertising and reducing its costs respectively.

5. Conclusions

Developed method of forming the context of advertising and target audience based on the associative rules learning makes it possible to increase the effectiveness of advertising, and thus reduce the cost of online advertising of higher education institutions. Also, the developed method will allow to form rules between the answers of respondents for the formation of advertising content and the definition of the target group.

To implement the developed method, we used a survey of students majoring in Computer Science, regarding admission. 152 students took part in the survey and answered 10 questions.

The proposed method enabled to increase the effectiveness of advertising on social networks at least in 23%, and reduce the price in 90%. Moreover, the authors believe that growing the number of student surveys can increase the quality of associative rules learning as well as get better keywords for content formation. That will lead to increasing the effectiveness of advertising and reducing its costs respectively.

Unlike analogues [3, 11, 18] the developed a method enables to form rules between the answers of respondents, construct the advertising content and determine the target group.

In the future, the authors plan to develop the information system that will automatically generate the content of the advertising message and select the target audience using the deep learning methods [27, 28].

6. References

- [1] M. M. Monastyrskaya, V. I. Soloviev, Improving customer relationship management based on intelligent analysis of user behavior patterns, in: Proceedings of the 2020 13th International Conference “Management of Large-Scale System Development” (MLSD), 2020, pp. 1-4. doi:10.1109/mlsd49919.2020.92477 doi: 10.1109/MLSD49919.2020.9247718.
- [2] A. Griva, C. Bardaki, K. Pramatar, Katerina, D. Papakyriakopoulos, Retail business analytics: Customer visit segmentation using market basket data, *Expert Systems with Applications*, 100 (2018), 1-16. doi: 10.1016/j.eswa.2018.01.029.
- [3] T. Badriyah, S. Azvy, W. Yuwono, I. Syarif, Recommendation system for property search using content based filtering method, in: Proceedings of the 2018 IEEE International Conference on Information and Communications Technology (ICOIACT), 2018, 25–29. doi: 10.1109/ICOIACT.2018.8350801.
- [4] B. Shazad, H. U. Khan, M. Farooq, A. Mahmood, I. Mehmood, S. Rho, Y. Nam, Finding temporal influential users in social media using association rule learning, *Intelligent Automation And Soft Computing* 26 (2020) 87–98. doi: 10.31209/2019.100000130.
- [5] S. Vanaja, M. Belwal, Aspect-level sentiment analysis on e-commerce data, in: Proceedings of the 2018 International Conference on Inventive Research in Computing Applications (ICIRCA), 2018, pp. 1275-1279. doi: 10.1109/ICIRCA.2018.8597286.
- [6] R. Agrawal, R. Srikant, Fast algorithms for mining association rules, in: Proceedings of the 20th International Conference on Very Large Data Bases VLDB, September 1994, vol. 1215, pp. 487-499.
- [7] H. Lipyanina, S. Sachenko, T. Lendyuk, A. Sachenko, Targeting model of HEI video marketing based on classification tree, in: Proceedings of the 16th International Conference on ICT in Education, Research and Industrial Applications. Integration, Harmonization and Knowledge Transfer. Volume II: Workshops, ICTERI 2020, Kharkiv, Ukraine, 6-10 October 2020, CEUR Workshop Proceedings, vol. 2732, pp. 487-498. <http://ceur-ws.org/Vol-2732/20200487.pdf>.
- [8] H. Lipyanina, A. Sachenko, T. Lendyuk, S. Nadvynychny, S. Grodskyi, Decision tree based targeting model of customer interaction with business page, in: Proceedings of the third International Workshop on Computer Modeling and Intelligent Systems (CMIS-2020), April 27 – May 1, 2020. CEUR Workshop Proceedings, vol. 2608, pp. 1001–1012. <http://ceur-ws.org/Vol-2608/paper75.pdf>.
- [9] F. Wang. K. Li. N. Duić. Z. Mi. B.-M. Hodge. M. Shafie-khah. J. P. S. Catalão, Association rule mining based quantitative analysis approach of household characteristics impacts on residential electricity consumption patterns, *Energy Conversion and Management* 171 (2018) 839–854. doi: 10.1016/j.enconman.2018.06.017.
- [10] S. Alzu'bi, B. Hawashin, M. Eibes and M. Al-Ayyoub, A novel recommender system based on apriori algorithm for requirements engineering, in: Proceedings of the 2018 Fifth International Conference on Social Networks Analysis, Management and Security (SNAMS), 2018, pp. 323-327. doi: 10.1109/SNAMS.2018.8554909.
- [11] J. Aguilar, G. Garcia, An adaptive intelligent management system of advertising for social networks: A case study of facebook, *IEEE Transactions on Computational Social Systems* 5 (2018) 20-32. doi: 10.1109/TCSS.2017.2759188.
- [12] B. S. Arasu, B. J. B. Seelan, N. Thamaraiselvan, A machine learning-based approach to enhancing social media marketing, *Computers & Electrical Engineering* 86 (2020) 106723. doi: 10.1016/j.compeleceng.2020.106723.
- [13] A. Popov, D. Iakovleva, Adaptive look-alike targeting in social networks advertising, *Procedia computer science* 136 (2018) 255-264. doi: 10.1016/j.procs.2018.08.264.

- [14] N. Shah, S. Engineer, N. Bhagat, et al., Research trends on the usage of machine learning and artificial intelligence in advertising, *Augment Hum Res* 5 (2020). doi: 10.1007/s41133-020-00038-8.
- [15] Z. Liu, Development of advertising art design based on information technology, in: J. Jansen B., Liang H., Ye J. (eds) *International Conference on Cognitive based Information Processing and Applications (CIPA 2021)*, volume 85 of *Lecture Notes on Data Engineering and Communications Technologies*, Springer, Singapore, 2021, pp. 3-10. doi: 10.1007/978-981-16-5854-9_1.
- [16] O. Barabash, G. Shevchenko, N. Dakhno, O. Neshcheret, A. Musienko, Information technology of targeting: optimization of decision making process in a competitive environment, *International Journal of Intelligent Systems and Applications* 9 (2017) 1-9. doi: 10.5815/ijisa.2017.12.01.
- [17] R. Saito, K. Otake, T. Namatame, Analysis of fashion market trend using advertising data of shopping information site, in: Meiselwitz G. (eds) *Social Computing and Social Media. Participation, User Experience, Consumer Experience, and Applications of Social Computing. HCII 2020*, volume 12195 of *Lecture Notes in Computer Science*, Springer, Cham, 2020, pp. 389-400. doi: 10.1007/978-3-030-49576-3_28.
- [18] K. Wakimoto, S. Kawamoto, P. Zhang, Keyword-based text generation for internet advertisement, in: *Proceedings of the 34th Annual Conference of the Japanese Society for Artificial Intelligence*, 2020, pp. 1-4.
- [19] D. Lee, K. Hosanagar, H. S. Nair, Advertising content and consumer engagement on social media: Evidence from Facebook, *Management Science* 64 (2018) 5105-5131. doi: 10.1287/mnsc.2017.2902.
- [20] A. M. Jamison, D. A. Broniatowski, M. Dredze, Z. Wood-Doughty, D. A. Khan, S. C. Quinn, Vaccine-related advertising in the Facebook Ad Archive, *Vaccine* 38, (2019) 512-520. doi: 10.1016/j.vaccine.2019.10.066.
- [21] S. Youn, S. Kim, Understanding ad avoidance on Facebook: Antecedents and outcomes of psychological reactance, *Computers in Human Behavior* 98 (2019) 232-244. doi: 10.1016/j.chb.2019.04.025.
- [22] C. L. White, B. Boatwright, Social media ethics in the data economy: Issues of social responsibility for using Facebook for public relations, *Public Relations Review* 46, (2020) 101980. doi: 10.1016/j.pubrev.2020.101980.
- [23] A. Gitomer, P. V. Oleinikov, L. M. Baum, et al., Geographic impressions in Facebook political ads, *Appl Netw Sci* 6, (2021). doi: 10.1007/s41109-020-00350-7.
- [24] A. Peruta, A. B. Shields, Social media in higher education: understanding how colleges and universities use Facebook, *Journal of Marketing for Higher Education* 27 (2017) 131-143. doi: 10.1080/08841241.2016.1212451.
- [25] S. Manca, Snapping, pinning, liking or texting: Investigating social media in higher education beyond Facebook, *The Internet and Higher Education* 44 (2020) 100707. doi: 10.1016/j.iheduc.2019.100707.
- [26] A. Alhegami, H. Alsaedi, A framework for incremental parallel mining of interesting association patterns for big data, *International Journal of Computing* 19 (2020) 106-117. URL: <http://computingonline.net/computing/article/view/1699>
- [27] M. Komar, A. Sachenko, V. Golovko and V. Dorosh, Compression of network traffic parameters for detecting cyber attacks based on deep learning, in: *Proceedings of the 2018 IEEE 9th International Conference on Dependable Systems, Services and Technologies (DESSERT)*, 2018, pp. 43-47, doi: 10.1109/DESSERT.2018.8409096.
- [28] Zh. Hu, Ye. V. Bodyanskiy, N. Ye. Kulishova, O. K. Tyshchenko, A multidimensional extended neo-fuzzy neuron for facial expression recognition, *International Journal of Intelligent Systems and Applications (IJISA)* 9 (2017) 29-36. DOI: 10.5815/ijisa.2017.09.04.

Methodology for Control of Helicopters Aircraft Engines Technical State in Flight Modes Using Neural Networks

Serhii Vladov^a, Yurii Shmelov^a and Ruslan Yakovliev^a

^a *Kremenchuk Flight College of Kharkiv National University of Internal Affairs, vul. Peremohy, 17/6, Kremenchuk, Poltavaska Oblast, Ukraine, 39605*

Abstract

In this work, the creation of a methodology for control the technical state of helicopters aircraft engines using recurrent neural networks, which, unlike other types of neural networks, have a more optimal method for training a neural network, and also after training allows you to get a lower percentage of errors in comparison with a convolutional neural network and multilayer perceptron. A mathematical description of the methodology for control the technical state of helicopters aircraft engines in flight mode using neural networks based on the numerical solution of a discrete optimal control problem has been implemented. The quality of the trained neural network is assessed, including the calculation of the testing error, which is no more than 1.2 % of the deviation from the a priori correct result on the vectors of the test set of sets. With an allowable value of 2 %, this allows us to speak about the efficiency of the method.

Keywords

Aircraft engine, recurrent neural network, testing error, optimal control, backpropagation method.

1. Introduction

In recent decades, the development and improvement of aircraft gas turbine engines (GTE), including helicopters, has been accompanied by toughening requirements for the reliability and efficiency of their automatic control systems (ACS). Modern helicopters gas turbine engines are complex technical devices that differ in the variety of physical processes occurring in them and are characterized by multidimensionality, multi-connectivity, nonlinearity, non-stationary work processes, a significant influence of operating modes and external conditions on the characteristics of their functioning. The listed features lead to the formation of a stable trend in the development of ACS for gas turbine engines of helicopters, characterized by a constant increase in the complexity and the number of tasks solved with their help. One of the important tasks is to improve the methods and algorithms for controlling the helicopters gas turbine engine in the conditions of a helicopter flight in terms of thermogasdynamic parameters, which is due to the presence of strict requirements for ensuring the safety and efficiency of flights [1, 2].

2. Literature review

Modern approaches to control of aircraft gas turbine engines technical state are described in [3–5]. The main factors that must be taken into account in the helicopter flight mode include uncertainty

The Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022), May 12-th, 2022, Zaporizhzhia, Ukraine
EMAIL: ser26101968@gmail.com (S. Vladov); nviddil.klk@gmail.com (Yu. Shmelov); ateu.nv.klk@gmail.com (R. Yakovliev)
ORCID: 0000-0001-8009-5254 (S. Vladov); 0000-0002-3942-2003 (Yu. Shmelov); 0000-0002-3788-2583 (R. Yakovliev)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

factors, such as incompleteness of a priori and operational information, inaccuracy of mathematical models of GTE, errors of sensors and actuators, changes in engine characteristics during operation, occurrence of possible failures of GTE functional elements [6, 7].

In recent years, their construction in the class of intelligent control systems providing robustness, adaptability and fault tolerance of GTE control processes in conditions of uncertainty have been considered as a promising direction in solving problems of monitoring the technical state of gas turbine engines.

Despite a significant amount of research in the development of GTE control algorithms, the existing information technologies for controlling GTE parameters are not perfect for a number of reasons. On the one hand, this is a weak informational "linkage", the absence of elements of "intelligence", allowing to quickly, efficiently and effectively support decision-making in emergency situations in flight.

On the other hand, it is the complexity of the processes occurring in GTE ACS, the complexity of their mathematical description, the limited composition of the measured parameters, their technological spread, etc. These lead to the need for comprehensive automation and intellectualization of decision-making processes of GTE ACS technical state [7] under conditions of uncertainty, including using the control method according to the FDI (Fault Detection and Identification) model [8–11].

Thus, the development of intelligent algorithms for automatic control of helicopters gas turbine engine technical state, as well as the study of the features of their practical application, taking into account the limitations on the available computing resources of an onboard digital computer, is an urgent scientific and practical task.

3. The use of neural networks in the problems of control of aircraft gas turbine engines technical state

The implementation of the FDI method [8–11] allows maximum consideration of the individual characteristics of a gas turbine engine by using a mathematical model that adapts (adjusts) to the individual characteristics of the latter [12]. When using neural networks to solve the problems of monitoring the technical state of a gas turbine engine of helicopters, the available a priori information is presented to the neural network in the form of ready-made solutions (problem books), on the basis of which the process of its training (additional training) is carried out. When assessing the quality of the network, data from the test sample are fed to its input, on the basis of which it calculates the vector of deviations (the difference between the output of the neural network and the desired characteristics) [11, 12].

Let us consider the problem of control of helicopters GTE technical state in flight mode based on neural networks in the following formulation [13, 14]. We will assume that all possible states of the gas turbine engine can be divided into two classes S_0 (all serviceable states of the gas turbine engine) and $\bar{S}_0$ (all faulty states, characterized by the presence of at least one defect in the operation of the gas turbine engine), uniting related states that are close to each other according to certain integral indicators. Required based on the results of a limited number of measurements of the vector of engine output parameters $Y(t_i)$, $t_i \in T$ (where t_i – discrete moments of time; T – observation interval), make a decision on whether the GTE belongs to one of the specified classes of states. The solution of this problem in general form is reduced to finding a certain separating function (hypersurface) in the space of controlled parameters of the gas turbine engine. To solve this problem, in this work, an approach is implemented based on the construction of the specified decision rule using neural networks.

4. Development of a methodology for control of helicopters aircraft gas turbine engines technical state in flight mode using neural networks

Based on the above, within the framework of the task of control of helicopters gas turbine engine technical state, in this work, neural network algorithms are investigated, a formalized formulation of

problems is given, recommendations for solving this spectrum of problems are formed, and an engineering methodology is proposed.

The general scheme of the methodology is shown in fig. 1. The main thermogasdynamic parameters of the aircraft engine (recorded by sensors and calculated using a mathematical model) are fed to the input.

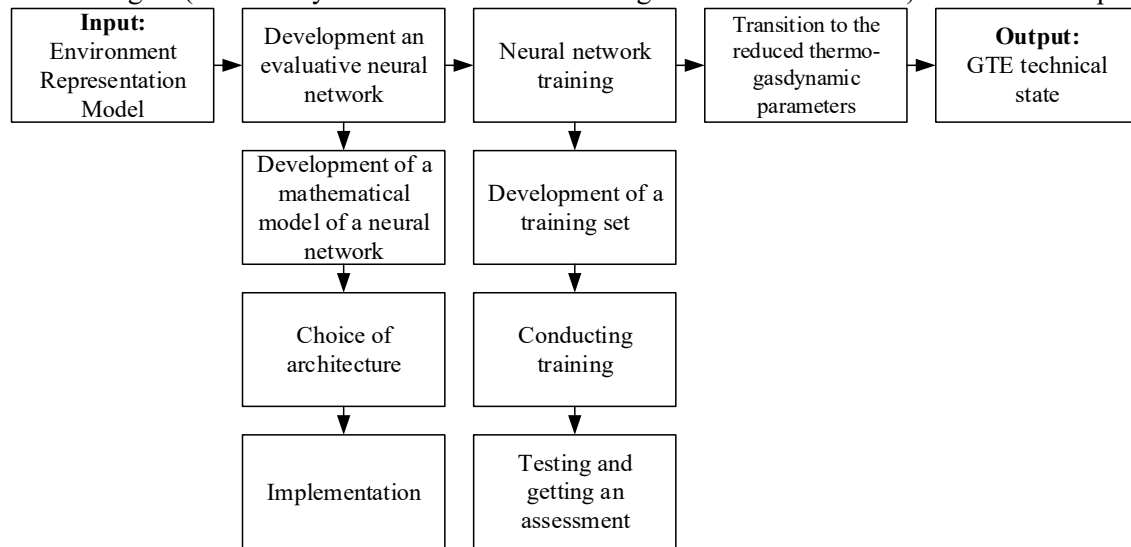


Figure 1: Block diagram of the methodology for control of helicopters aircraft gas turbine engines technical state in flight mode using neural networks

At the first stage, a model of an aircraft engine is formed by systematizing the data. Then a neural network is created – in view of the fact that neural networks have the property of non-universality, it is necessary to carry out procedures for choosing an architecture for each new task. For each task, it is necessary to select several types of architectures, conduct testing in the context of a specific task. Accordingly, it becomes necessary to train / retrain the newly created neural network. Then, for the convenience of further use and displaying the graphical representation, the transition to the reduced thermogasdynamic parameters [15, 16]. This stage is necessary for human control of the neural network, in the case when there is an additional stage of control.

5. Review and selection of neural network architecture

The architecture of neural networks most often used to solve problems of monitoring the technical state of complex dynamic objects is a feedforward network, the input neurons of which are fed with the values of the attributes of the classified object, and the output is a label or a numeric code of the class. Multilayer perceptrons are commonly used. In such networks, the elements of the feature vector arrive at the input neurons and are distributed to all neurons of the first hidden layer of the neural network, and as a result, the dimension of the problem changes [17, 18].

A convolutional neural network processes the data not entirely, but in fragments, but the data is not split into parts, but a kind of sequential run is carried out. Then the data is transferred further along the layers [19, 20]. In addition to convolutional layers, union layers are also used. Combine layers are compressed with depth (usually a power of two). Several perceptrons (feedforward network) are added to the final layers for further data processing.

In the process of implementing a convolutional neural network, technical problems often arise related to the format of the input data and the model used, which did not allow the implementation of this network and adequately assess the capabilities of this architecture for solving the problem of monitoring the technical condition of an aircraft engine.

The use of a recurrent neural network can significantly reduce the computational complexity of the backpropagation method, which will be used to train the neural network [21, 22]. To train the neural network, we will use the error backpropagation algorithm, this algorithm is optimal for the classification problem using a neural network. It is also worth noting that this neural network needs long-term training on a larger number of training sets than a multilayer perceptron.

Comparing the proposed architectures, we can conclude that for the task at hand, the optimal option would be a recurrent neural network, which, after training, makes it possible to obtain a lower percentage of errors in comparison with a convolutional neural network and a multilayer perceptron (comparison graphs are shown in fig. 2), in addition, it possesses a simpler implementation than a convolutional neural network.

This architecture assumes a more optimal method for training a neural network. As a disadvantage, there will be a larger number of training sets and a longer training time, which is compensated by the higher accuracy of calculations (fig. 3).

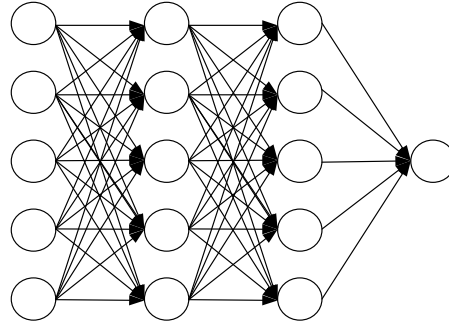


Figure 2: Final neural network architecture



Figure 3: Comparison of the percentage of errors of classical neural network architectures designed for control complex dynamic objects: 1 – recurrent neural network; 2 – multilayer perceptron; 3 – convolutional neural network

It is also worth noting that a recurrent neural network of the LSTM structure [23] has been successfully applied to solve the problem of identification of an aviation GTE. Therefore, in this work, a recurrent neural network of the LSTM structure is used to solve the problem of monitoring the technical state of aircraft GTEs of helicopters in flight mode. At the same time, increasing the accuracy of the GTE model as a control object by modifying the architecture of the LSTM network.

6. Mathematical description of the methodology for control of helicopters gas turbine engine technical state in flight mode using neural networks

Let us solve the optimal control problem that simulates the dynamics of the considered artificial neural network in the notation given below. The dynamics of a network of n neurons is described by a system of differential equations with delay:

$$\dot{x}_i(t) = -\gamma_i x_i(t) + \sum_{j=1}^n [\omega_{ij}(t) g(x_j(t)) + u_i(t)]; \quad i, j = \overline{1, N}; \quad (1)$$

where $\dot{x}_i(t) = -\gamma_i x_i(t) + g_i(z_i(t)) + u_i(t); \quad i, j = \overline{1, N}; \quad z_i(t) = \sum_{j=1}^n \omega_{ij}(t) x_j(t-h);$

$$g_i(z_i(t)) = \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}}; \text{ thus, } \dot{x}_i(t) = -\gamma_i x_i(t) + \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}} + u_i(t); \quad i, j = \overline{1, N}.$$

From the point of view of a biological prototype, i.e. neural network, we can say that the equation describes the accumulated potential (electrical impulse) of a neuron at a given time, as well as its change over time. The potential of a neuron is formed and changes under the influence of many factors: $x_i(t)$ – own potential of a neuron at a given time; γ_i – intrinsic attenuation of the i -th neuron, describes the effect on the neuron of its own forces, which negatively affect the potential, as well as signal attenuation during transmission from one neuron to others. Sum $\sum_{j=1}^n \omega_{ij}(t) x_j(t-h)$ can be called the sum of the potentials of an ensemble of neurons. This is the sum of the impact of all neighboring neurons on the i -th neuron. Sum $\sum_{j=1}^n \omega_{ij}(t) x_j(t-h)$ is the main element in the formation of the potential of the i -th neuron, so this sum will be called the body of the i -th neuron $\omega_{ij}(t) x_j(t)$. Element $x_j(t-h)$ shows the lag of the neural network signal. Thus, the potential of the i -th neuron is sufficiently influenced by the residual impulse of neurons at the previous moment in time.

Control functions $\omega_{ij}(t)$ describe the axons of neurons – an electrical or chemical impulse that is transmitted from one neuron to another, thereby changing the potentials, is the most important connecting element of the neural network, since responsible for the interaction and performance of the entire network. In this case, this is the impact on the i -th neuron, j -th neuron.

The activation function $g_i(z_i(t))$ transforms the accumulated potential of the neuron according to some functional dependence. The prototype is the processes occurring in the body of a neuron, caused, for example, by signals or impulses from the peripheral nervous system of the body (due to any changes in the external environment).

The intrinsic potential of neurons should not go beyond the limits:

$$x_i(t) \leq B_i; \quad i = \overline{1, n}. \quad (2)$$

The characteristics of neurons at the initial moment of time are known:

$$x_i(0) = a_i; \quad i = \overline{1, n}; \quad x_i(t) = \varphi_i(t); \quad i = \overline{-h, 0}. \quad (3)$$

The control function $u_i(t)$ characterizes the external influence on the i -th neuron. These can be any changes in the environment to which the body reacts by changing the rate of transmission of electrical and chemical impulses of the nervous system. Restrictions on control functions are known:

$$|\omega_{ij}| \leq b; \quad |u_i| \leq c. \quad (4)$$

The goal of controlling the dynamics of a neural network is to train the network, which implies the following tasks (criteria):

- 1) At the final moment of time, the characteristics of the neurons must coincide with the input data A_i ;
- 2) During the execution of the process, the characteristics of neurons should not go beyond the specified range of B_i values.

3) Control u_i should strive for the minimum value for the given process.

4) Control ω_{ij} should also tend to the minimum value for this process.

Control tasks can be formalized in the form of the following target functional:

$$I([x], [\omega], [u], [t]) = S \sum_{i=1}^n (x_i(T) - A_i)^2 + \sum_{i=1}^n \int_0^T M_i \left(\max(0; x_i(t) - B_i) \right)^2 dt + L \sum_{i=1}^n \int_0^T u_i^2(t) dt + K \sum_{j=1}^n \sum_{i=1}^n \int_0^T \omega_{ij}^2(t) dt \rightarrow \inf; \quad i = \overline{1, n}. \quad (5)$$

The main practical goal is to obtain optimal process controls, with the help of which the minimum of the target functional is achieved.

We obtain the optimal process controls by the gradient descent method (back propagation of the error) [24].

To solve the optimal control problem (2)–(5), we use the Pontryagin maximum principle [25]. We introduce the Pontryagin function:

$$H([x],[\omega],[u],[t]) = -\lambda_0 \sum_{i=1}^n \left(M_i \left(\max(0; x_i - B_i) \right)^2 + Lu_i^2 \right) - \lambda_0 K \sum_{j=1}^n \sum_{i=1}^n \omega_{ij}^2 + \sum_{i=1}^n p_i(t) \times \left(-\gamma_i x_i + \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}} + u_i \right); \quad i, j = \overline{1, n}. \quad (6)$$

Let's write down the maximum principle:

$$\begin{aligned} & -\lambda_0 \sum_{i=1}^n M_i \left(\max(0; \bar{x}_i - B_i) \right)^2 - \lambda_0 \sum_{i=1}^n Lu_i^2 - \lambda_0 K \sum_{j=1}^n \sum_{i=1}^n \bar{\omega}_{ij}^2 + \lambda_0 \sum_{i=1}^n L \bar{v}_i^2 + \sum_{i=1}^n p_i(t) \left(-\gamma_i \bar{x}_i + g_i(z_i(t)) + \bar{u}_i \right) = \\ & = \max \left(-\lambda_0 \sum_{i=1}^n M_i \left(\max(0; \bar{x}_i - B_i) \right)^2 - \lambda_0 \sum_{i=1}^n L \bar{v}_i^2 - \lambda_0 K \sum_{j=1}^n \sum_{i=1}^n \bar{\omega}_{ij}^2 + \sum_{i=1}^n p_i(t) \left(-\gamma_i \bar{x}_i + g_i(\bar{x}_j \omega_{ij}) + v_i \right) \right) = \\ & = -\lambda_0 \sum_{i=1}^n M_i \left(\max(0; \bar{x}_i - A_i) \right)^2 - \sum_{i=1}^n p_i(t) \gamma_i \bar{x}_i + \sum_{i=1}^n \left(\max_{\omega_{ij}} \left(p_i(t) g_i(\bar{x}_j \omega_{ij}) \right) - \lambda_0 K \sum_{j=1}^n \omega_{ij}^2 \right) + \\ & \quad + \sum_{i=1}^n \max_{v_i} \left(p_i(t) v_i - \lambda_0 L \bar{v}_i^2 \right). \end{aligned}$$

Conjugate vector functions are calculated according to the expressions:

$$\dot{p}_k(t) = -\frac{\partial H}{\partial y_k}(t+h);$$

where $y_k = x_k(t-h)$; thus,

$$\dot{p}_k(t) = 2\lambda_0 M_k \max(0; x_k - A_k) + p_k(t) \gamma_k - \lambda \sum_{i=1}^n p_i(t+h) \bar{\omega}_{ik} \frac{e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)} \right)^2};$$

and satisfy the transversality conditions $p_k(t) = -\lambda_0 \frac{\partial \Phi}{\partial x_k} = -2\lambda_0 (x_k(T) - A_k)$.

The maximum principle allows us to reduce the problem of optimal process control to solving a boundary value problem:

$$\dot{x}_i(t) = -\gamma_i x_i(t) + \sum_{j=1}^n \bar{\omega}_{ij}(t) g(x_j(t)) + \bar{u}_i(t); \quad i, j = \overline{1, n};$$

$$\dot{x}_i(t) = -\gamma_i x_i(t) + g_i(z_i(t)) + u_i(t);$$

Wherein $x_i(0) = a_i$; $i, j = \overline{1, n}$; where $\bar{\omega}_{ij}$, $\bar{u}_{ij}$ – optimal controls.

7. Numerical solution of a discrete optimal control task

Let us obtain a numerical solution to the discrete optimal control problem. For this, we will reduce the original problem to a discrete form. We split the segment $[0, T]$ into q segments. Sampling step

$\Delta t = \frac{T}{q}$. Let us approximate the system of differential equations according to the Euler scheme [26]:

$$\dot{x}_i(t^l) \approx \frac{x_i^{l+1} - x_i^l}{\Delta t}.$$

We denote $x_i(t) = x_i^l$; $h = v\Delta t$; $x_i(t-h) = x_i^{l-v}$, after that we approximate the integrals $\int_0^T M_i(\max(0; x_i(t) - A_i))^2 dt$, $\int_0^T u_i^2(t) dt$ and $\int_0^T \omega_{ij}^2(t) dt$ by the method of rectangles:

$$\int_0^T M_i(\max(0; x_i(t) - A_i))^2 dt \approx \sum_{l=0}^{q-1} \sum_{i=1}^n M_i(\max(0; x_i(t) - A_i))^2 \Delta t;$$

$$\int_0^T u_i^2(t) dt = \sum_{l=0}^{q-1} \sum_{i=1}^n (u_i^l)^2 \Delta t;$$

$$\int_0^T \omega_{ij}^2(t) dt = \sum_{l=0}^{q-1} \sum_{i=1}^n \sum_{j=1}^n (\omega_{ij}^l)^2 \Delta t.$$

We get the recurrence expression for calculating the phase trajectory:

$$x_i^{l+1} = x_i^l + (-\gamma_i x_i^l + g_i^l(z_i^l) + u_i^l) \Delta t;$$

where $z_i^l = \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}$; $g_i^l(z_i^l) = \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}}}$; thus, $x_i^{l+1} = x_i^l + \left(-\gamma_i x_i^l + \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}}} + u_i^l \right) \Delta t;$

wherein $(g_i^l)_{w_{ij}} = \lambda \frac{e^{-\lambda \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}} x_j^{l-v}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}} \right)^2}$; $(g_i^l)_{x_k^m} = \frac{e^{-\lambda \sum_{j=1}^n \omega_{ik}^m x_k^{m-v}} \omega_{ik}^{m-v}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ik}^m x_k^{m-v}} \right)^2}.$

Initial values:

$$x_i(0) = a_i; \quad x_i(t) = \varphi_i(t); \quad t = \overline{-h, 0}. \quad (6)$$

The objective function will take the form:

$$I = S \sum_{i=1}^n (x_i^q - A_i)^2 + \sum_{l=0}^{q-1} \sum_{i=1}^n M_i(\max(0; x_i^l - B_i))^2 \Delta t + L \sum_{l=0}^{q-1} \sum_{i=1}^n (u_i^l)^2 \Delta t + K \sum_{l=0}^{q-1} \sum_{i=1}^n \sum_{j=1}^n (\omega_{ij}^l)^2 \Delta t \rightarrow \inf. \quad (7)$$

Using the Lagrange multiplier method, we obtain the necessary optimality condition. Let us compose the Lagrange function [27]:

$$L = \lambda_0 S \sum_{i=1}^n (x_i^q - A_i)^2 + \sum_{l=0}^{q-1} \sum_{i=1}^n M_i(\max(0; x_i^l - B_i))^2 + L \sum_{l=0}^{q-1} \sum_{i=1}^n (u_i^l)^2 \Delta t + K \sum_{l=0}^{q-1} \sum_{i=1}^n \sum_{j=1}^n (\omega_{ij}^l)^2 \Delta t + \\ + \sum_{l=0}^{q-1} \sum_{i=1}^n p_i^{l+1} \left(x_i^{l+1} - x_i^l - \Delta t \left(-\gamma_i x_i^l + \frac{1}{1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}^l x_j^{l-v}}} + u_i^l \right) \right).$$

Let us write down the stationarity conditions:

$$\frac{\partial L}{\partial x_k^m} = p_k^m - p_k^{m+1} + p_k^{m+1} \gamma_k \Delta t - \Delta t \sum_{i=1}^n p_i^{m+1} \frac{\partial g_k^m(z_k^m)}{\partial x_k^m} + 2M_k \lambda_0 \Delta t (\max(0; x_k^m - B_m)) = 0; \\ k = \overline{1..n}; \quad m = \overline{0..q-1}.$$

$$\frac{\partial L}{\partial x_k^m} = p_k^m - p_k^{m+1} + p_k^{m+1} \gamma_k \Delta t - \Delta t \sum_{i=1}^n p_i^{m+v+1} \omega_{ik}^{m+v} \frac{e^{-\lambda \sum_{j=1}^n \omega_{ik}^m x_k^{m-v}}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ik}^m x_k^{m-v}} \right)^2} + 2M_k \lambda_0 \Delta t (\max(0; x_k^m - B_m)) = 0;$$

$$k = \overline{1...n}; \quad m = \overline{0...q-1}; \quad (8)$$

$$\frac{\partial L}{\partial x_k^q} = 2\lambda_0 S(x_k^q - A_k) + x_k^q = 0; \quad k = \overline{1...n}. \quad (9)$$

Let's write down the gradient values:

$$\frac{\partial L}{\partial \omega_{km}^s} = -p_k^{s+1} \frac{\partial g_k^m(z_k^m)}{\partial \omega_{km}^s} \Delta t + 2K \Delta t \omega_{km}^s; \quad k = \overline{1...n}; \quad s = \overline{1...n}; \quad m = \overline{0...q-1};$$

$$\frac{\partial L}{\partial \omega_{km}^s} = -\lambda \Delta t p_k^{s+1} x_m^{s-v} \frac{e^{-\lambda \sum_{m=1}^n \omega_{km}^s x_m^{s-v}}}{\left(1 + e^{-\lambda \sum_{m=1}^n \omega_{km}^s x_m^{s-v}}\right)^2} + 2K \Delta t \omega_{km}^s; \quad (10)$$

$$\frac{\partial L}{\partial u_k^m} = 2L u_k^m \Delta t - p_k^{m+1} \Delta t; \quad k = \overline{1...n}; \quad m = \overline{0...q-1}. \quad (11)$$

Let us express the impulses in terms of these equations p_k^m and p_k^q :

$$p_k^m = p_k^{m+1} - p_k^{m+1} \gamma_k \Delta t + \gamma \Delta t \sum_{i=1}^n p_i^{m+v+1} \omega_{ik}^{m+v} \frac{e^{-\lambda \sum_{k=1}^n \omega_{ik}^m x_k^{m-v}}}{\left(1 + e^{-\lambda \sum_{k=1}^n \omega_{ik}^m x_k^{m-v}}\right)^2} - 2M_k \lambda_0 \Delta t \left(\max(0; x_k^m - B_k)\right); \quad (12)$$

$$p_k^q = -2\lambda_0 S(x_k^q - A_k). \quad (13)$$

After completing the passage to the limit, then we check the correspondence with the boundary value problem, thereby checking the correctness of the solution:

$$\frac{p_k^{m+1} - p_k^m}{\Delta t} = p_k^{m+1} \gamma_k - \lambda \sum_{i=1}^n \omega_{ik}^{m+v} p_i^{m+v+1} \frac{e^{-\lambda \sum_{j=1}^n \omega_{ij}^m x_j^{m-v}}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}^m x_j^{m-v}}\right)^2} + 2M_k \lambda_0 \left(\max(0; x_k^m - B_k)\right);$$

passing to the limit $\lim_{\Delta t \rightarrow 0} \frac{p_k^{m+1} - p_k^m}{\Delta t}$, we get:

$$\dot{p}_k(t) = 2\lambda_0 M_k \max(0; x_k - B_k) + p_k(t) \gamma_k - \lambda \sum_{i=1}^n p_i(t+h) \omega_{ik}(t+h) \frac{e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}}{\left(1 + e^{-\lambda \sum_{j=1}^n \omega_{ij}(t) x_j(t-h)}\right)^2};$$

$$p_k^q = -2\lambda_0 S(x_k^q - A_k);$$

$$p_k(T) = -2\lambda_0 S(x_k(T) - A_k).$$

8. Development of an algorithm for finding a numerical solution

Let's formalize the resulting algorithm:

- 1) We set the parameters of the model.
- 2) We set the parameters of the method: ε – calculation accuracy, α – gradient descent step, q – number of dividing points of the segment and the initial set of controls $[u]^{(0)}$ and $[\omega]^{(0)}$.

3) Using expressions $\mathbf{W}^{(\Sigma)} = \mathbf{W}^1 \cdot \mathbf{W}^2 \cdot \dots \cdot \mathbf{W}^{(p)}$ and $\mathbf{Y} = \mathbf{X} \cdot \mathbf{W}^{(\Sigma)}$, we find a collection $[x]^{(0)}$. We calculate $I^{(0)}$ by the expression $s = \sum_{i=1}^n x_i^2 \cdot w_i$, setting $[M]^{(1)}$.

4) Using expressions (12) and (13), starting from the q -th layer, we calculate $[p]^{(k)}$, $k = 0, 1, \dots$

We calculate $\left[\frac{\partial L}{\partial \omega_{km}^s} : \frac{\partial L}{\partial u_k^s} \right]^{(k)}$ by expressions (10) and (11).

5) Improving control at $k + 1$ steps: $[u]^{(k+1)} = [u]^{(k)} - \alpha \frac{\partial \bar{L}}{\partial u}$, $[w]^{(k+1)} = [w]^{(k)} - \alpha \frac{\partial \bar{L}}{\partial w}$.

6) We calculate $[x]^{(k+1)}$, then $[I]^{(k+1)}$.

7) Compare the values $[I]^{(k)}$ and $[I]^{(k+1)}$. If $[I]^{(k+1)} > [I]^{(k)}$, then we decrease the step of the gradient descent $\alpha = \frac{\alpha}{N}$, where $N \in \mathbb{N}$ we go to step (5) of the same iteration, otherwise go to step (8).

8) We check the condition $\left| [I]^{(k+1)} - [I]^{(k)} \right| < \varepsilon$, if it is satisfied, then we assume that the optimal solution for a given parameter $[M]^{(1)}$ has been found and go to step (9). Otherwise, go to the next iteration $k = k + 1$ and go to step (4).

9) Checking the conditions:

$$\sum_{l=1}^{q-1} \sum_{i=1}^n M_i^l \left(\max(0; x_i^l - A_i) \right)^2 \leq \varepsilon; \sqrt{\sum_{l=0}^q \sum_{i=1}^n \left(x_i^{l(k)} - x_i^{l(k+1)} \right)^2} \leq \varepsilon; \sqrt{\sum_{i=0}^q \sum_{i=1}^n \sum_{j=1}^n \left(\omega_{ij}^{l(k)} + \omega_{ij}^{l(k+1)} \right)^2} \leq \varepsilon.$$

If all conditions are met together, we assume that the solution has been found; otherwise, we determine a new penalty coefficient $[M]^{(1)} = \beta [M]^{(0)}$ and go to step (3).

9. Algorithm for training a formalized neural network

The next step in the methodology for control of helicopters aircraft GTE technical state is training an artificial neural network. It is necessary to choose a training algorithm suitable for a specific task. According to the research results, the best results were obtained using the error backpropagation algorithm. This is primarily due to the fact that the backpropagation algorithm has been created and optimized for the selected neural network architecture and activation function, which makes it possible to maximize its potential [22, 23].

Its main idea is that the change in the weights of synapses takes into account the local gradient of the error function. The difference between the real and correct responses of the neural network, determined on the output layer, propagates in the opposite direction (fig. 4) – towards the flow of signals.

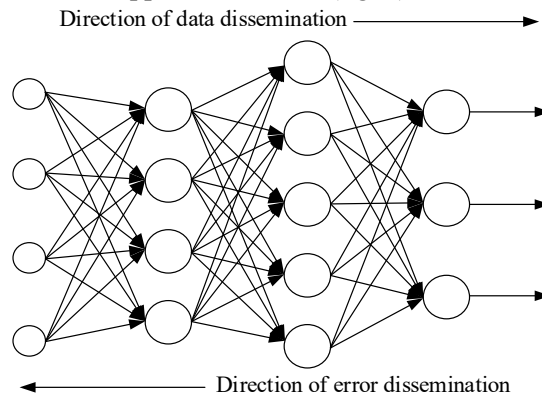


Figure 4: Backpropagation method [23] for multilayer fully connected neural network

As a result, each neuron is able to determine the contribution of each of its weights to the total network error. The simplest learning rule corresponds to the steepest descent method, that is, changes in synaptic weights are proportional to their contribution to the total error [28].

Of course, with such a training of a neural network, there is no certainty that it has trained in the best way, since there is always a possibility of the algorithm falling into a local minimum. For this, special techniques are used to "knock" the found solution out of the local extremum. If, after several such actions, the neural network converges to the same solution, then we can conclude that the solution found is most likely optimal [29].

The algorithm for training a neural network using the backpropagation procedure is constructed as follows [24, 30]:

1. Apply one of the possible images to the network inputs and, in the normal functioning of the neural network, when signals propagate from inputs to outputs, calculate the values of the latter.
2. Calculate the difference between the ideal and the obtained output values for the output layer. Calculate changes in layer weights.
3. Calculate the difference between the ideal and the obtained values of the output and change in weights for all other layers.
4. Adjust all weights to the neural network.
5. If the network error is significant, go to step 1. Otherwise, end. At step 1, all training images are alternately presented to the network in a random order so that the network, figuratively speaking, does not forget some as it memorizes others.

When the output value goes to zero, the training efficiency decreases markedly. With binary input vectors, on average, half of the weight coefficients will not be corrected; therefore, it is desirable to shift the range of possible values of neuron outputs $[0, 1]$ within the limits $[-0.5 \dots 0.5]$, which is achieved by simple modifications of logistic functions. For example, an exponential sigmoid

$$\sigma(x) = \frac{1}{1 + e^{-x}} \text{ is converted to a kind } \sigma(x) = -0.5 + \frac{1}{1 + e^{-x}}.$$

Now let's touch on the issue of the capacity of a neural network, that is, the number of images presented to its inputs that it is able to learn to recognize. For networks with more than two layers, it remains open. For a neural network with two layers, that is, an output and one hidden layer, the deterministic capacity of the network is estimated as follows:

$$\frac{N_w}{N_y} < C_d \frac{N_w}{N_y} < \log \frac{N_w}{N_y};$$

where C_d – deterministic capacity of the network; N_w – number of adjustable weights; N_y – number of neurons in the output layer.

It should be noted that this expression was obtained with some restrictions. First, the number of inputs N_x and neurons in the hidden layer N_h must satisfy the inequality $N_x + N_h > N_y$. Secondly, $\frac{N_w}{N_y} > 1000$.

The adjective “deterministic” appearing in the name of the capacity means that the obtained capacity estimate is suitable for absolutely all possible input patterns that can be represented by N_x inputs. In reality, the distribution of input images, as a rule, has some regularity, which allows the neural network to generalize and thus increase the real capacity. Since the distribution of images, in the general case, is not known in advance, we can talk about such a capacity only conjecturally, but usually it is twice the deterministic capacity.

10. LSTM structure recurrent neural network modification

According to [23], in this work it is advisable to apply an LSTM network based on a dynamic model of a gas turbine engine based on the classical LSTM structure (fig. 5, a) or LSTM structure with variable memory (fig. 5, b) proposed by Georgy Makaryants and Alexander Kuznetsov. Each cell has one output neuron for predicting some parameter (for example, the gas generator r.p.m.). A collection of cells is a network for predicting multiple parameters. The main difference between LSTM networks and other recurrent networks is the memory tensor. A memory tensor is a variable, information into which can be

record or erased during the operation of the network. The operation of the LSTM network (fig. 5) is based on the principle of memory tensor control using memorization and forgetting nodes:

$$c_t = f_t \circ c_{t-1} + i_t \circ cc_t; \quad (14)$$

where c_t – memory tensor representing a vector of weighted inputs in a step t ; f_t – tensor at the exit from the forgetting node in the step t , representing the sum of the weighted inputs and outputs in the previous step $t-1$; c_{t-1} – step memory tensor $t-1$; i_t – tensor at the output of the input node at step t , which is the sum of the weighted inputs at the current step and outputs at the previous step; cc_t – candidate tensor to write to memory tensor [23].

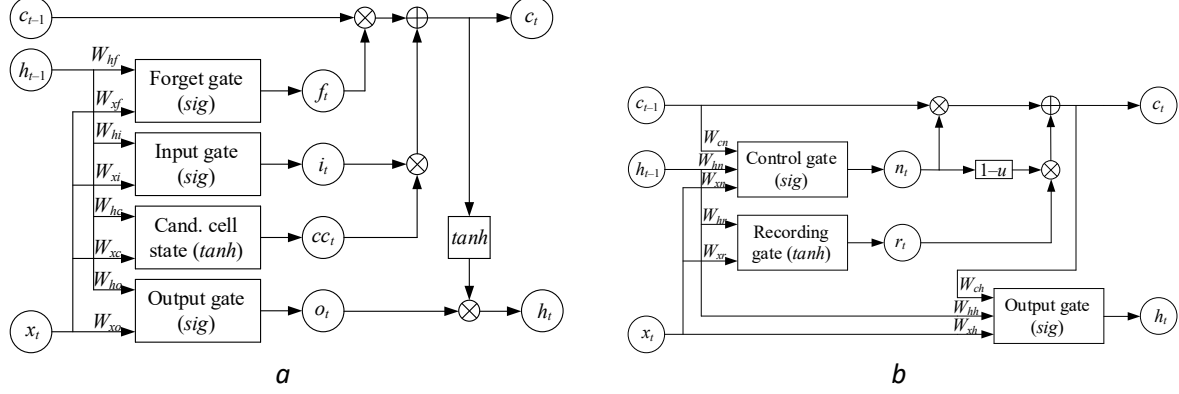


Figure 5: LSTM network diagram [23]: a – classical LSTM structure; b – LSTM structure with variable memory

As mentioned above, to solve the problem of dynamic monitoring of aircraft gas turbine engines, a special architecture of a recurrent neural network with long-short-term memory (LSTM) was developed, presented in [23]. When using LSTM networks, the problem of the vanishing gradient of the LSTM network arises, which the filter mechanism allows to resist. This mechanism makes it possible to regulate the flow of new information into the state vector c_t of the network, as well as the output of the state h_t of the network and updating its state c_t . The network filter vectors are determined according to the expressions [31]:

$$i_t = \sigma \cdot (U_i \cdot x_t + W_i \cdot h_{t-1}); \quad (15)$$

$$f_t = \sigma \cdot (U_f \cdot x_t + W_f \cdot h_{t-1}); \quad (16)$$

$$o_t = \sigma \cdot (U_o \cdot x_t + W_o \cdot h_{t-1}); \quad (17)$$

$$g_t = \sigma \cdot (U_g \cdot x_t + W_g \cdot h_{t-1}); \quad (18)$$

where $\sigma(z) = -0.5 + \frac{1}{1 + e^{-z}}$ – sigmoidal function; $\tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$ – hyperbolic tangent function;

x – input sequence; h – hidden state vector of the network cell; $U_i, U_f, U_o, U_g, W_i, W_f, W_o, W_g$ – network filter weight matrices i, f, o, g ; t – index of the element of the training sequence.

Based on the values of the network filters, its internal state (internal memory) and hidden state vectors are determined [31]:

$$c_t = \tanh(i_t \circ g_t + f_t \circ c_{t-1}); \quad (19)$$

$$h_t = o_t \circ c_t; \quad (20)$$

where t_c – vector of the internal state of the network cell; t_h – vector of the hidden state of the network cell; $\circ$ – elementwise product operation.

Expressions (15) – (20) correspond to the stage of direct signal propagation through the network. This mechanism allows the LSTM network to deal with the vanishing gradient problem when working with long sequences. By training filter parameters ($U_i, U_f, U_o, U_g, W_i, W_f, W_o, W_g$) LSTM network "tunes" its "memory". In this paper, to solve the problem of saturation of the activation function, it is proposed to use an activation function that does not saturate. This approach will allow make network training faster and more accurate. The activation function based on the logarithm, in contrast to [32], avoids saturation when processing large values and is defined by the expression:

$$f(x) = \begin{cases} -0.5 + \ln(x+1), & x \geq 0 \\ -0.5 - \ln(x-1), & x < 0 \end{cases} \quad (21)$$

The advantage of the proposed activation function compared to the hyperbolic tangent function is its "unsaturation" and therefore its application will improve the efficiency of training the LSTM network. The sigmoid activation function is also subject to the saturation problem, but it is not possible to replace it with the proposed one, since it cannot serve as a gateway due to the fact that it does not scale the input values in the range from 0 to 0.5.

11. Formation of training and test subsets

The training set consists of a sufficient number of vectors (10000 vectors are enough for the task at hand), which are input data sets with the correct result. The sets included the main thermodynamic parameters of TV3-117 aircraft engine according to table 1 [33, 34]. Each of these parameters has been assigned a corresponding priority required for the functioning of the neural network. To form the training and test subsets in the work, cross-validation was used [35] to estimate the values of the parameters of TV3-117 aircraft GTE, the results of which are shown in fig. 6.

Table 1

Fragment of the training set

Engine unit	Parameter	1	2	3	4
Input device	ΔP_{in}^*	0.984	0.878	0.812	0.766
	ΔT_{in}^*	0.935	0.864	0.744	0.707
Compressor	ΔP_{comp}^*	0.961	0.902	0.853	0.781
	ΔT_{comp}^*	0.977	0.912	0.834	0.762
Combustion chamber	ΔP_{gas}^*	0.942	0.865	0.776	0.693
	ΔT_{gas}^*	0.954	0.875	0.762	0.689
Compressor turbine	$\Delta P_{comp.turb}^*$	0.988	0.931	0.902	0.894
	$\Delta T_{comp.turb}^*$	0.975	0.917	0.899	0.884
Free turbine	$\Delta P_{freeturb}^*$	0.933	0.898	0.835	0.808
	$\Delta T_{freeturb}^*$	0.946	0.909	0.844	0.801
Output device	ΔP_{out}^*	0.942	0.907	0.891	0.855
	ΔT_{out}^*	0.944	0.908	0.983	0.861

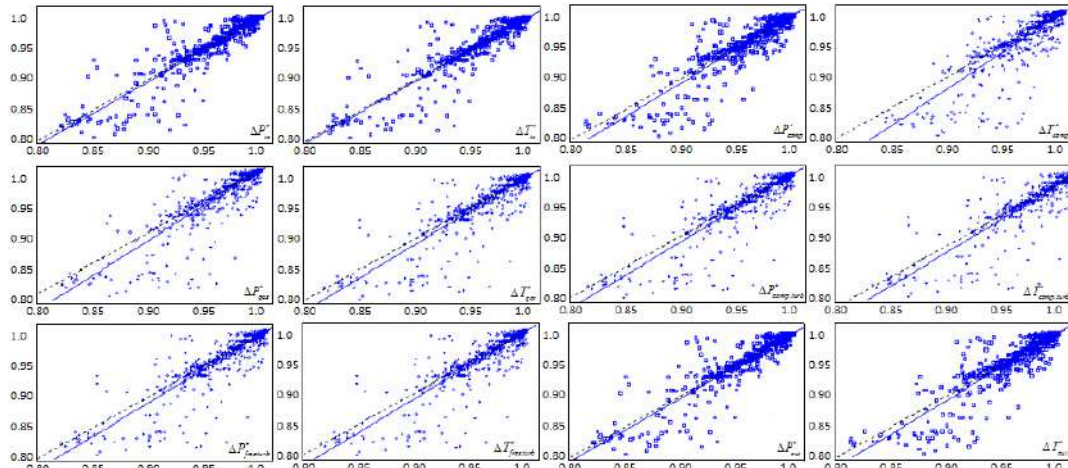


Figure 6: TV3-117 aircraft engine input parameter scatter diagram

12. Quality assessment of a trained neural network

As an example, a test version of a recurrent neural network was developed in the Matlab environment, illustrating the solution to the problem of TV3-117 aircraft GTE technical state control. The neural network was trained according to the following rule [22, 23]. First, all 10000 vectors of the training set were sequentially fed to the network input, with the help of which training was performed using the backpropagation algorithm. Each time, after training, a training test was performed for 2000 sets – the error was checked on 100 arbitrary sets of those already passed (test error control). Then, after completing the entire training, the error was checked for 1000 (10 % of the training set) – the stage of final control. At this stage, the erroneous estimate did not exceed the predetermined threshold of 1.2 % (fig. 7), which is a good result.

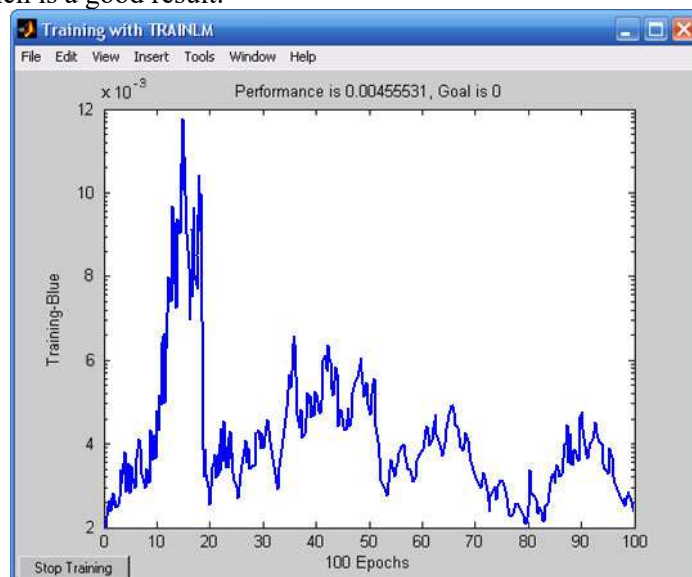


Figure 7: Neural network training results graph

In this case, [36] were used:

- Adam's algorithm as an optimization algorithm [37];
- Accuracy indicator as an objective function, where P – number of objects for which the neural network made the correct decision; N – number of objects in the training set.

Is the binary-crossentropy function that returns the classification error as the logistic loss function $Loss$:

$$Loss = -\frac{1}{N} y_i \log(\hat{y}_i) + (1 - \hat{y}_i) \log(1 - \hat{y}_i); \quad (22)$$

where y_i – true class label; $\hat{y}_i$ – response of the classifier (calculated class label) on the i -th object; N – number of classes.

To assess the quality of neural network training, various quality indicators can be used, in particular, such indicators as *Accuracy*, *Precision*, *Recall*, *F-measure*.

In the case of a binary classification based on the errors matrix of inaccuracies (errors), a 2×2 table can be drawn up (table 2), in which the following designations are used: TP – true-positive solution; TN – true negative decision; FP – false positive decision; FN – false negative decision.

Table 2

Errors matrix

Classification		Class labels in the dataset	
		Positive grade label	Negative grade label
Class labels exposed by the neural network	Positive grade label	TP	TN
	Negative grade label	FP	FN

The *Accuracy* indicator determines the proportion of objects for which the classifier made the right decision:

$$Accuracy = \frac{P}{N} = \frac{TP + TN}{FP + FN + TP + TN}; \quad (23)$$

The *Precision* indicator within a class determines the proportion of objects correctly assigned by the classifier to the class, to the total number of objects assigned by the classifier to the class in the training (test) sample. The higher the *Precision* score, the fewer false positives.

The *Recall* indicator within a class determines the proportion of objects correctly assigned by the classifier to the class to the number of objects of this class in the training (test) sample. The higher the value of the *Recall* indicator, the less false-negative decisions.

In the case of binary classification, *Precision* and *Recall* are defined as:

$$Precision = \frac{TP}{TP + FP}; \quad (24)$$

$$Recall = \frac{TP}{TP + FN}. \quad (25)$$

In the simplest case, the *F-measure* is defined as the harmonic average between *Precision* and *Recall*:

$$F = 2 \frac{Precision \cdot Recall}{Precision + Recall}. \quad (26)$$

The results of training the neural network according to the *Accuracy* and *Loss* indicators are shown in fig. 9 and 10 respectively. As can be seen from fig. 8 and 9, the *Accuracy* indicator approaches one, and *Loss* indicator – tends to zero, which indicates the high accuracy of the model and its minimal error.

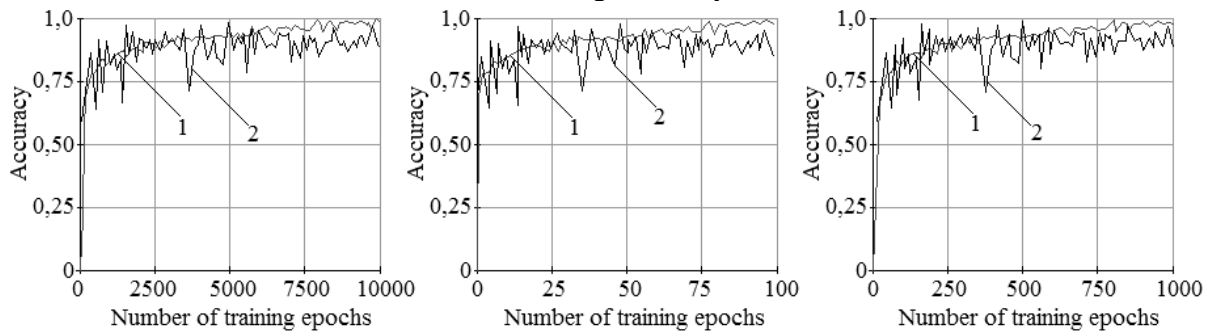


Figure 8: Model training results in terms of the *Accuracy* indicator: 1 – test; 2 – train

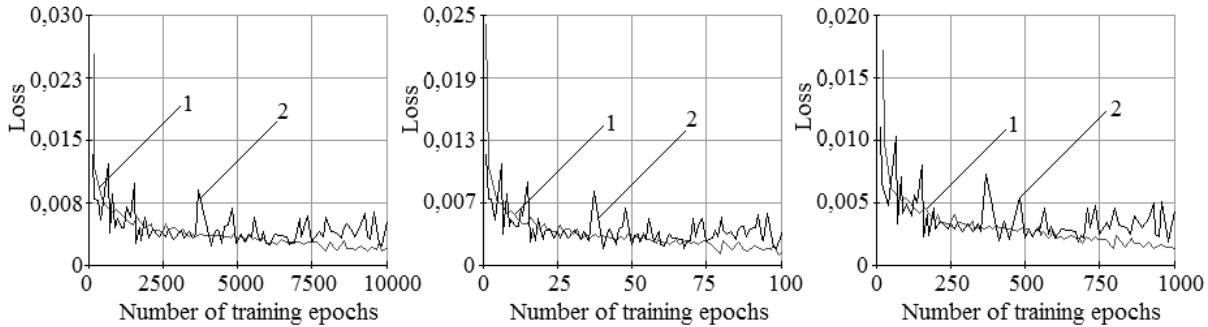


Figure 9: Model training results in terms of the *Loss* indicator: 1 – test; 2 – train

Table 3 shows the averaged values of the model learning outcomes, as well as the mean and variance values for the *Accuracy* indicator.

Table 3

Average values of neural network testing indicators

Accuracy	F-measure	Precision	Recall	Average time, s	Average Accuracy	Dispersion Accuracy
0.98867	0.97234	0.94617	1.0	1204.12	0.99011	0.0000085

13. Results and discussion

In fig. 10 shows a graph of the hypersurface in the space of the controlled parameters of helicopters gas turbine engine in flight modes (for example, the TV3-117 aircraft engine). As a result of its work, the trained neural network divided all the values of the degree of increase in the total pressure in the compressor fed to its inputs into 3 regions (fig. 10), corresponding to the serviceable (blue), faulty (red) and indefinite states where it is difficult to perform separation of parameter values due to their mutual overlap (green). The proposed neural network, integrated into the GTE control system, is capable of real-time correlating the value of the monitored parameter (the degree of increase in the total pressure in the compressor) with one of the areas and, if it enters the area of a faulty state, to give a signal about an incipient malfunction. This information (depending on the degree of danger) can be issued to the crew for timely adoption of the correct decision or to GTE automatic control system the power plant as a whole.

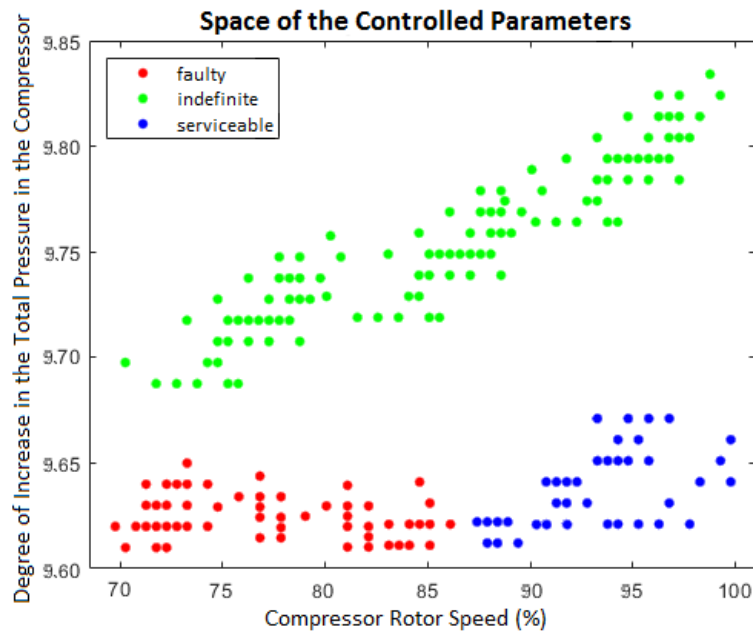


Figure 10: Results of TV3-117 aircraft engine technical state control in flight modes

Analogous research has been carried out using other neural network architectures (multilayer perceptron, convolutional neural network, Hopfield network, Kosko network, Jordan network, Elman network). It should be noted that the use of these neural network architectures did not allow to adequately classify the states of the TV3-117 aircraft engine, namely, the areas of serviceable, faulty and undefined states were not clearly distinguished. The results of a comparative analysis of the accuracy of solving the problem of control of helicopters aircraft GTE technical state in flight modes using the aforementioned neural network architectures are given in table 4.

Table 4

Results of comparative analysis of control accuracy

Neural network architecture	Absolute error (%)		
	Serviceable state	Faulty state	Indefinite state
Modified LSTM network	0.64	0.64	0.65
LSTM network	1.32	1.32	1.35
Multilayer perceptron	2.08	2.12	2.14
Convolutional neural network	16.35	17.48	19.24
Hopfield network	5.66	6.84	6.93
Kosko network	8.32	8.78	8.81
Jordan network	3.76	3.77	3.89
Elman network	3.35	3.35	3.46

Analysis of the table 4 shows that the use of the modified LSTM network makes it possible to solve the problem of control of helicopters aircraft GTE technical state in flight modes with maximum accuracy (the smallest error).

14. Conclusions

In this work, a methodology for control of helicopters aircraft engines technical state in flight modes has been developed, based on the operation of artificial neural networks and the transition to the given indicators of engine thermogasdynamic parameters. The scientific novelty of the result lies in the use of a neural network specially designed and trained on the originally formed a priori sets of engine thermogasdynamic parameters by a neural network, which makes it possible to increase the speed of systems and reduce the load on hardware resources, as well as to topologize the results of evaluating of engines technical state, which makes it possible to more clearly display problem engine nodes (input device, compressor, combustion chamber, compressor turbine, free turbine, output device) and simplify the optimization processes of solutions at a specific node (by reducing the number of engine thermogasdynamic parameters).

The proposed neural network was trained and tested on the practical task of control of TV3-117 aircraft engine technical state, based on the flight data of the Mi-8MTV helicopter. The testing error is calculated, which is no more than 1.2 % of the deviation from the a priori correct result on the vectors of the test set of sets. With an allowable value of 2 %, this allows us to speak about the efficiency of the method.

The use of such on helicopter board system, in contrast to the existing system for control the parameters of a gas turbine engine, will allow control of helicopters gas turbine engine technical state in flight in real time and recognize the failure at an early stage and inform the crew or the engineering staff about it. Becoming, which will be the guarantee of the correctness of the decision on the possibility of using all the potential of the gas turbine engine and will lead to an increase in the level of flight safety.

Thus, as can be seen from the results of experimental studies, recurrent neural networks demonstrate their high efficiency in solving the problem of monitoring the probable class of errors in the operation of equipment in complex dynamic systems (control of helicopters aircraft engines technical state in flight mode).

15. References

- [1] B. Li, Y.-P. Zhap, Y.-B. Chen, Unilateral alignment transfer neural network for fault diagnosis of aircraft engine, *Aerospace Science and Technology*, vol. 118 (2021) 107031. Doi: 10.1016/j.ast.2021.107031
URL: <https://www.sciencedirect.com/science/article/abs/pii/S1270963821005411>
- [2] S. N. Vassilyev, A. Yu. Kelina, Y. I. Kudinov, F. F. Pashchenko, *Intelligent Control Systems, Procedia Computer Science*, vol. 103 (2017) 623–628. 10.1016/j.procs.2017.01.088
- [3] E. L. Ntantis, Diagnostic methods for an aircraft engine performance, *Journal of Engineering Science and Technology*, vol. 8, no. 4 (2015) 64–72. Doi: 10.25103/jestr.084.10
- [4] I. A. Krivosheev, K. E. Rozhkov, N. B. Simonov, Complex Diagnostic Index for Technical Condition Assessment for GTE, *Procedia Engineering*, vol. 206 (2017) 176–181. Doi: 10.1016/j.proeng.2017.10.456
- [5] S. V. Zhernakov, A. T. Gilmanshin, New algorithms for on-board diagnostics of aircraft gas turbine engine based on fuzzy networks, *Bulletin of USATU*, vol. 19, no. 2 (68) (2015) 63–68.
- [6] S. V. Zhernakov, Algorithms for control and diagnostics of an aviation gas turbine engine under conditions of onboard implementation based on neural network technology, *Bulletin of USATU*, vol. 14, no. 3 (38) (2010) 42–56.
- [7] S. Kiakojoori, K. Khorasani, Dynamic neural networks for gas turbine engine degradation prediction, health monitoring and prognosis, *Neural Computing & Applications*, vol. 27, no. 8 (2015) 2151–2192.
- [8] S. V. Enchev, S. S. Tovkach, Diagnostic of the Technical Condition Aviation Engines is Based on Fuzzy logic, *Scientific Bulletin Kherson State Maritime Academy*, no 1 (8) (2013) 216–224.

- [9] V. I. Vasiliev, S. V. Zhernakov, Control and diagnostics of aircraft engine technical state based on expert systems, *Bulletin of USATU*, vol. 9, no 4 (22) (2007) 11–23.
- [10] S. V. Zhernakov, R. F. Ravilov, Control and diagnostics of aircraft engine technical state based on the C-Priz expert system, *Bulletin of USATU*, vol. 16, no 6 (51) (2012) 3–11.
- [11] Yu. M. Shmelov, S. I. Vladov, A. S. Khebda, K. G. Kotliarov, Application of the rules of urban logic in the problem of identification of the technical state of the aviation engine TV3-117, *Scientific notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences*, vol. 30 (69), no 3 (2018) 34–40.
- [12] A. M. Pashayev, D. D. Askerov, C. Ardil, R. A. Sadiqov, P. S. Abdullayev, Condition monitoring system of aircraft gas turbine engine complex, *International Journal of Aerospace and Mechanical Engineering*, vol. 1, no. 11 (2007) 689–695.
- [13] S. Vladov, K. Kotliarov, S. Hrybanova, O. Husarova, I. Derevyanko, S. Gvozdik, Neuro-mechanical methods of control and diagnostics of the technical state of aircraft engine TV3-117 in film regions, *Visnyk of Kherson National Technical University*, no. 1 (72), part 1 (2020) 141–154.
- [14] Y. Shmelov, S. Vladov, A. Kryshan, S. Gvozdik, Application of neural network technologies in the technical state control system of the aircraft engine TV3-117 in flight modes, *Radiotekhnika*, no. 194 (2018) 147–154.
- [15] S. Vladov, Y. Shmelov, R. Yakovliev Control and Diagnostics of TV3-117 Aircraft Engine Technical State in Flight Modes Using the Matrix Method for Calculating Dynamic Recurrent Neural Networks. CMIS-2021: The Fourth International Workshop on Computer Modeling and Intelligent Systems, April 27, 2021, Zaporizhzhia, Ukraine. CEUR Workshop Proceedings (ISSN 1613-0073), vol. 2864 (2021) 97–109.
- [16] S. Vladov, Yu. Shmelov, K. Kotliarov, S. Hrybanova, O. Husarova, I. Derevyanko, L. Chyzhova, Onboard parameter identification method of the TV3-117 aircraft engine of the neural network technologies, *Transactions of Kremenichuk Mykhailo Ostrohradskyi National University*, issue 5/2019 (118) (2019) 90–96. Doi: 10.30929/1995-0519.2019.5.90-96
- [17] V. I. Vasiliev, S. V. Zhernakov, I. I. Muslukhov, Onboard algorithms for monitoring GTE parameters based on neural network technology, *Bulletin of USATU*, vol. 12, no. 1 (30) (2009) 61–74.
- [18] F. A. Del Campo, M. C. Guevara Neri, O. O. V. Villegas, V. G. C. Sanchez, H. J. O. Domínguez, V. G. Jimenez, Auto-adaptive multilayer perceptron for univariate time series classification, *Expert Systems with Applications*, vol. 181 (2021) 115147.
- [19] W. Guo, J. Zhao, J. Zhang, H. Wei, A. Song, K. Zhang, Stability analysis of opposite singularity in multilayer perceptrons, *Neurocomputing*, vol. 282 (2018) 192–201. Doi: 10.1016/j.neucom.2017.12.018
- [20] D. Shi, B. Lam, K. Ooi, X. Shen, W.-S. Gan, Selective fixed-filter active noise control based on convolutional neural network, *Signal Processing*, vol. 190 (2021) 108317. Doi: 10.1016/j.sigpro.2021.108317
URL: <https://www.sciencedirect.com/science/article/pii/S0165168421003546>
- [21] R. Shang, Y. Meng, W. Zhang, F. Shang, L. Jiao, S. Yang, Graph Convolutional Neural Networks with Geometric and Discrimination information, *Engineering Applications of Artificial Intelligence*, vol. 104 (2021) 104364. Doi: 10.1016/j.engappai.2021.104364
- [22] Zgurovsky M., Sineglazov V. and Chumachenko E. (2021) *Artificial Intelligence Systems Based on Hybrid Neural Networks: Theory and Applications*, Publisher: Springer, 734 p.
- [23] A. V. Kuznetsov, G. M. Makaryants, Gas turbine engine dynamic model based on variable-memory LSTM architecture, *Vestnik of Samara University. Aerospace and Mechanical Engineering*, vol. 19, no. 2 (2020) 38–52. Doi: <https://doi.org/10.18287/2541-7533-2020-19-2-38-52>
- [24] Vladov S., Shmelov Yu. And Shmelova T. (2020) Control and diagnostics of TV3-117 aircraft engine technical state in flight modes using neural network technologies, *Kremenichuk Novabook*, 200 p.
- [25] F. Cardin, A. Spiro, Pontryagin maximum principle and Stokes theorem, *Journal of Geometry and Physics*, vol. 142 (2019) 274–286. Doi: 10.48550/arXiv.1812.07875
- [26] Z. Chen, S. Gan, X. Wang, A full-discrete exponential Euler approximation of the invariant measure for parabolic stochastic partial differential equations, *Applied Numerical Mathematics*, vol. 157 (2020) 135–158. Doi: 10.48550/arXiv.1811.01759

- [27] S. Giuffre, Lagrange multipliers and non-constant gradient constrained problem, *Journal of Differential Equations*, vol. 269, issue 1 (2020) 542–562. Doi: 10.1016/j.jde.2019.12.013
- [28] M. Meyer, D. Baldwin, Statistical learning of action: The role of conditional probability, *Learning & Behavior*, vol. 39 (2011) 383–398. Doi: 10.3758/s13420-011-0033-7
- [29] Bodyansky E. V. and Rudenko O. G. (2004) Artificial neural networks: architectures, training, applications. Kharkiv, Teletech, 369 p.
- [30] A. Glushchenko, V. Petrov, K. Lastochkin, Backpropagation method modification using Taylor series to improve accuracy of offline neural network training, *Procedia Computer Science*, vol. 186 (2021) 202–209. Doi: 10.1016/j.procs.2021.04.139
- [31] I. Chernobaev, A. Surkova, A. Pankratova, Modelling of texts using recurrent neural networks, *Proceedings of Nizhny Novgorod State Technical University n.a. R.E. Alekseev*, no 1 (120) (2018) 65–73.
- [32] S. Gluge, R. Bock, G. Palm, A. Wendemuth, Learning long-term dependencies in segmented-memory recurrent neural networks with backpropagation of error, *Neurocomputing*, vol. 141 (2014) 54–64. Doi: 10.1016/j.neucom.2013.11.043
- [33] S. I. Vladov, N. V. Podgornyykh, V. Ya. Teleshun, Mathematical model of TV3-117 aircraft engine compressor for its control and diagnostics of a technical state in the conditions of onboard operation of the aircraft, *The path to success and prospects for development (to the 26th anniversary of the Kharkiv National University of Internal Affairs) : proceedings of the international scientific-practical conference*, November 20, 2020, Kharkiv 112–116.
- [34] R. R. Mukhamedov, GTE mathematical models, *Youth Bulletin of USATU*, no 1 (10) (2014) 35–43.
- [35] J. Wainer, G. Cawley, Nested cross-validation when selecting classifiers is overzealous for most practical applications, *Expert Systems with Applications*, vol. 182 (2021) 115222. doi: 10.1016/j.eswa.2021.115222
- [36] L. Demidova, D. Marchev, Application of recurrent neural networks in the classification problem of failures in the complex technical systems within the framework of proactive maintenance, *Vestnik of RSREU*, no 69 (2019) 135–148. doi: 10.21667/1995-4565-2019-69-135-148
- [37] D. P. Kingma, J. L. Ba, Adam: A Method for Stochastic Optimization, 3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 07–09, 2015. 15 p. URL: <https://arxiv.org/pdf/1412.6980.pdf>

Model of Finding Associative Rules in Inhomogeneous Data of Semantic Networks

Nataliya Boyko, Vladyslav Mykhailyshyn

Lviv Polytechnic National University, Profesorska Street 1, Lviv, 79013, Ukraine

Abstract

The paper examines the issues of hidden connections and potentially useful information from large data sets. Theoretical knowledge about associative rules is substantiated, their influence on connections in multimodal data sets is investigated. The methods of application of associative rules in practice are analyzed. The following are considered in detail: the basic concepts of associative rules and their connection with the idea of logical regularity; ways to determine the "strength" of these connections; basic algorithms for finding patterns; practical implementation of the search for associative rules. The regularities in the "templates" are analyzed: support and confidence value. The correct choice of these values, which directly affect the results of the search for rules, is experimentally determined. Research in this paper aimed to consider the basic concepts and find the Associative Rules both in traditional ways and in heterogeneous data of semantic networks, which creates specific problems when using existing algorithms. The data of semantic networks are analyzed, which in most cases serve a particular field and are highly specialized. The research presents the process of finding associative rules through the work of the classical Apriori algorithm and an alternative algorithm for finding associative rules. Previously, this problem was considered only to a small extent. The results of experiments on accurate SW data showed promising results.

Keywords ¹

Artificial Intelligence Systems, Big Data Mining, Associative Rule, Database, Semantic Web, Data Mining, Health-e-Child.

1. Introduction

The primary purpose of data mining is to "reveal" the hidden connections and potentially useful information from large datasets [3, 7, 8]. Associative rules are one of the ways that help to identify these connections. Currently, there are many areas where the search for associative rules is used, as in IT (search for associations between data in the list of databases transactions, analysis of weblogs) and in the consumer sphere (the problem of the product basket, product placement, demand forecasting), areas of marketing (search for market segments, trends, identification of firms clients groups). This work will be discussed in detail:

- The basic concepts of associative rules and their connection with logical regularity.
- Ways to determine the "strength" of these connections.
- Basic algorithms for finding frequency.
- Practical implementation of searching for associative rules.

For the first time, searching for associative rules arose in the consumer area: it was necessary to identify specific "patterns" of consumer purchases to increase sales of goods due to these data. The Associative rule acquired of the form: "Event X is followed by event Y", as a result of which it is possible to obtain a certain regularity - if the purchase (transaction) has a set of goods (elements) X,

¹The Fifth International Workshop on Computer Modeling and Intelligent Systems (CMIS-2022), May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: nataliya.i.boyko@lpnu.ua (N. Boyko); vladyslavmykhailyshyn@gmail.com (V. Mykhailyshyn)
ORCID: 0000-0002-6962-9363 (N. Boyko); 0000-0003-1889-9053 (V. Mykhailyshyn)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org) Proceedings

then with some probability to assume that the set Y will also appear in it [4]. These rules are characterized by support and confidence values. The correct choice of these values directly affects the search results of the rules. Yes, if the support value is too large, the algorithm's results will be already known and quite noticeable. On the other hand, too small the support value will help identify many different patterns, but there can be some doubts about their reliability. The same with confidence value - the design will be less "valuable" at too low values [1, 6].

The department grouped all products in grocery stores to find what they needed more quickly. It reduces spending time shopping in a store and is also interested in buying something else. Keep in mind that associative rules will not help, in this case, the consumer's personal preferences. Still, with their help, it is possible to find connections between items in each purchase transaction (as opposed to filtering preferences, which considers all purchases of one consumer to recommend to him goods or services in the future). Therefore, the data to search for associative rules are regarded as separate purchases of different consumers by one group.

2. General theoretical information and the concept of associative rule

The purpose of AR [5, 10] is to find all the relationships (also called associations) between datasets, elements of certain regularity between date [13]. The basic concept of AR can be represented as follows: let $S = \{s_1, s_2, \dots, s_m\}$ the list of things, then $T = \{t_1, t_2, \dots, t_n\}$ the list of transactions (purchases), where each transaction is a set of things from the list S (that is $t_i \subseteq S$). Exactly AR is presented in the form $X \rightarrow Y$, where $X \subset S, Y \subset S$ and $X \cap Y = \emptyset$ (X and Y - set of things) [13]. Let our shopping database DB look like this way (Table 1):

Table 1

Example of rule: $\{B, C\} \rightarrow \{A\}$

ID	Items
0001	A, F, B, C
0002	A, C
0003	A, D
0004	D, E, C
0005	F, A, D

Let's say we want to analyze how sales are related to certain goods in the store. In this case, the S list will include all goods in the store, and the transaction (purchase) will be a set of things in the buyer's basket. Let's find AR: $\{A, F\} \rightarrow \{C\}$ (where $\{A, F\} = X, \{C\} = Y$): transaction $t_i \in T$ must contain a list of things X (that is, it is a subset t_i , "covers" the transaction).

2.1. Concept of Support value

In the case of a product basket problem, consider the following example of a rule: Bread $\rightarrow$ Apples (support value = 20%, confidence value = 45%). The results show that 20% of buyers buy bread along with apples at the time like 45% of buyers who buy bread, they also buy apples. Support value and confidence value determine "power" of this rule [2, 5, 9].

The calculation of the support value rule is to represent transactions $t_i \in T$, which are subordinate $X \cup Y$, and in some way is the probability $P(X \cup Y)$, in other words, it is the amount or percentage of transactions that have a set of specific elements. The support value of the $X \rightarrow Y$ rule will be calculated by the Formula 1:

$$\text{support value} = \frac{\text{support count}(X \cup Y)}{n}, \quad (1)$$

where n - the number of transactions in the T list (in our database); $\text{support count}(X \cup Y)$ - the number of transactions in T that include X and Y [2, 9].

Support value is useful in cases where it's too small value indicates that the rule can happen "accidentally". Support value list of things $\{A, F, C\} = 1/5 = 20\%$ in our DB (Table 1).

2.2. Concept of Confidence value

Confidence value rule consist in representing transactions $t_i \in T$ with values from the list Y, which include X, and in some way is a conditional probability $P(Y|X)$. In other words, confidence value determines how often "things" in list Y appear in transactions with things in list X. Confidence value rule $X \rightarrow Y$ will be calculated by the Formula 2 [12, 15, 7]:

$$\text{confidence value} = \frac{\text{support count}(X \cup Y)}{\text{support count}(X)}, \quad (2)$$

where $\text{support count}(x)$ - the number of transactions in T that include X; $\text{support count}(X \cup Y)$ - the number of transactions in T that include X and Y [15, 7].

The confidence value determines the "predictability" of a rule. When its value is too small, there is a problem of reliability of definition or prediction of Y and X. Confidence value rule $\{A, F\} \rightarrow \{C\} = 1/2 = 50\%$ in our DB (Table 1).

2.3. Concept of Lift value

There is a problem: what to do when confidence value, for example, of rule $\{A, F\} \rightarrow \{C\}$, is less than $P(\{C\})$? Lift value acts as an indicator of the "predictive power" of the rule compared to a random event, calculated (Formula 3) [8, 17, 21]:

$$\text{Lift value}(X \rightarrow Y) = \frac{P(Y|X)}{P(Y)}, \quad (3)$$

where (if lift value > 1 , then Y will occur more likely at a given X; lift value < 1 - Y will occur less likely at a given X).

Can be represented as $\frac{\text{confidence value}(X \rightarrow Y)}{\text{support count}(Y)}$ (if X and Y are independent, the value of lift value will be equal to 1; if X and Y occur more often than if they were independent, lift value > 1) [11, 15, 19].

3. Materials and methods

Studies of semantic annotations show that the increase in semantic networks (SW) data is constantly growing. Problems with RDF / (S) and OWL data extraction occur with graphics-oriented software. The solution to this problem may be to use new algorithms that use the axioms of ontologies (Tbox) to obtain the corresponding transactions, which will later use traditional association rule algorithms to find the result. With the help of the analyst, you can control the process represented by query templates. [4, 12, 20].

Today dictates new languages to create alternative associative rules. It combines semantic networks and data mining that allows you to carry out the process faster and without errors. Recently, a new direction has been widely developed, which researchers call Ontology Learning [2, 18, 21].

If you consider a database of unstructured and different data types, you need to use data extraction in semantic networks as the most appropriate for this data set. Its application allows the processing of data sets with different semantics and structure. In this study, we will consider the possibility of combining ontological instances expressed by OWL into search mechanisms by traditional algorithm search algorithms. The task is to reduce the space for searching for the data you are looking for.

Large data sets need to be used to find patterns in traditional networks. Semantic data requires a different approach than the one offered by machine learning [9, 12]. Therefore, in the process of finding patterns, there are several problems, including:

- In the process of processing the algorithm, a priori for homogeneous data should apply transactions. Accordingly, each transaction is a subset of the elements. If we use semantic networks, we will see that ontological axioms describe the subject area. Semantic annotations, in this case, are represented by statements that describe each instance through its property. In semantic networks, as an example, you can find a triplet - this is when the information is represented through the subject, predicate and object. When viewing such a scenario, the very definition of transactions and elements is not trivial. In semantic networks, elements can

correspond to both instances and literals. accordingly, a transaction can be defined by a subset of elements that are semantically related in the data warehouse [10, 11].

- The ontology of the semantic network represents data through formal properties and particular semantics. In this case, the network does not have a clear structure; respectively, the instances belonging to it may belong to a specific class and have a different form [19, 20].

Previous work on SW data extraction has focused mainly on clustering and instance classification. However, the presentation methods required for associative data mining are different from the usual clustering and classification tasks. In the general concept of transactions, the rules are specific observations of the frequency of occurrence of a particular set of elements. For example, in vector-numerical form, the presented data sets are the same as in a traditional format. They can be used to cluster and classify data. For our study, it is essential to process semi-structured and heterogeneous data. Therefore, an ontology should be used to identify data quickly. It is also necessary to consider the ways of creating elements and transactions in semantic networks, as they depend on the level of detail and the very structure of semantic data.

3.1. Definition of SW data

Analyzing the above, it can be argued that to incorporate semantics into current web content, and you need to use appropriate technologies. They can be used to provide presentation elements for the language. For example, you can consider specific presentation formats based on XML. Therefore, to describe semantic metadata, we use the resource description language (RDF), which contains three types of elements. The first is the resource, all web objects that a URI can identify. The second is literals, i.e. numbers, atomic values, dates, strings, etc. Third, properties are binary relationships between resources and literals that a URI can identify. The main components of the RDF are triplets: a binary relationship between two resources or between a resource and a literal. The resulting metadata can be considered as a graph, where nodes are resources and literals, and edges are the properties that connect them. RDFS extends RDF to allow you to define triplets by classes and properties. Thus, we can describe a schema that manages our metadata in the same description frame. An ontology web language (OWL) was later proposed to facilitate the work on the semantic description [3, 14, 17].

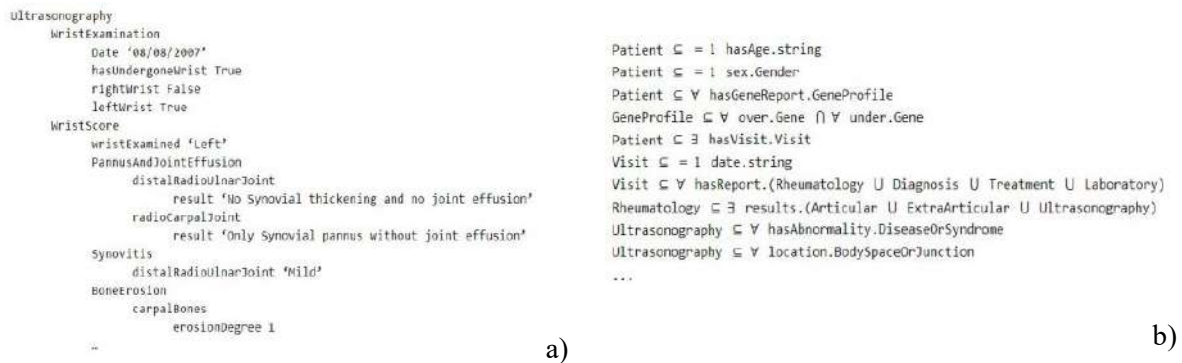


Figure 1: a) Excerpt from a patient's clinical report in the field of rheumatology; b) Axioms of ontologies (Tbox)

One of the areas of application may be medicine - here all the time a huge amount of semantic data is generated. In particular, most semi-structured and very heterogeneous data sources (e.g., laboratory test reports, ultrasound scans, images) are subjected to semantic annotation using UMLS, NCI and Galen ontologies. Suppose we have an excerpt from a clinical report presented in Figure 1a, the semantic annotation of which leads to certain axioms (Figure 1 b) and statements (Figure 2). Axioms in Figure 3 provide the semantics of all information concerning the patient (i.e., medical history, reports, laboratory results) by conceptualizing the domain. Figure 4 provides data on triplets that describe the patient (i.e., semantic annotations (Abox); subjects, predicates, and URI-formatted objects that point to relevant data resources). The generated data also represent complex relationships that are rapidly

evolving with the use of new biomedical research methods. Obviously, traditional analytical tools are not suitable for this type of data [1, 8, 12].

Axioms of ontologies (Tbox) allow to define an area from the point of view of atomic concepts (classes in OWL) and roles (properties in OWL). OWL provides for the union $\cup$, intersection $\cap$ and negation $\neg$, as well as list classes (one Of), existence $\exists$, universality $\forall$ and constraints ($\leq$, $\geq$, $=$) of the atomic concept of R or the inverse $\neg R$.

Subject	Predicate	Object
PTNXZ1	hasAge	'10'
PTNXZ1	sex	'Male'
VISIT1	date	'06/18/2008'
VISIT1	hasReport	RHEX1
RHEX1	damageIndex	'10'
RHEX1	results	ULTRA1
ULTRA1	hasAbnormality	'Malformation'
ULTRA1	hasAbnormality	'Knee'
VISIT1	hasReport	DIAG1
...	...	...

Figure 2: Three patient descriptions

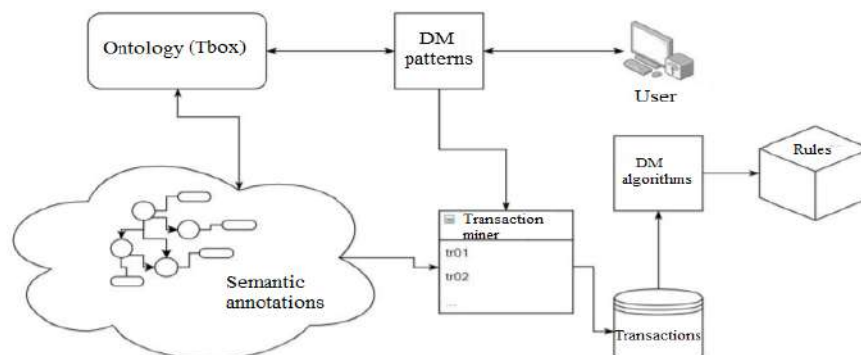


Figure 3: Scheme of the approach of extracting associative rules of semantic annotations

This section presents an overview of the method according to the scheme shown in Figure 3. The user specifies the extraction pattern using the query language syntax. The transaction miner can identify and construct transactions according to a previously defined mining scheme. Finally, the set of received transactions is processed by the traditional pattern mining algorithm, which finds the associative rules according to the minimum values of support value and confidence value defined in the template for the network shown in Figure 4.

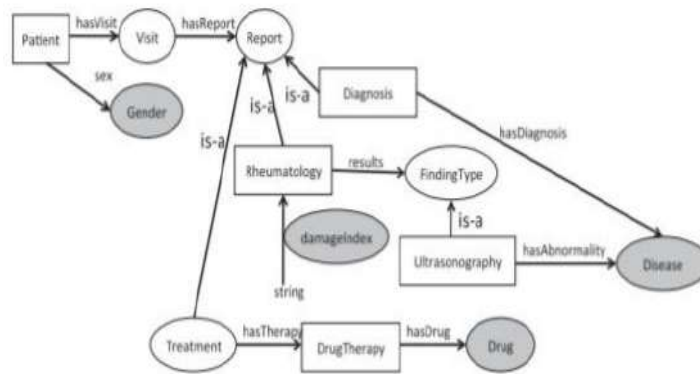


Figure 4: A fragment of the semantic network in the field of rheumatology

Both the ontology and the instances can be represented in the form of subject-predicate-object triplets (Figures 5-6), forming a graph where nodes are resources and literals, and edges are properties that connect. This dynamic graph-based structure contrasts with the well-structured and homogeneous datasets used in conventional associative rule search algorithms. Therefore, to obtain entities and transactions, users must specify the target concept of the analysis and related functions. The features must be relevant to the target concept, i.e., they will be extracted from the subgraph of each instance belonging to the target concept of the analysis.

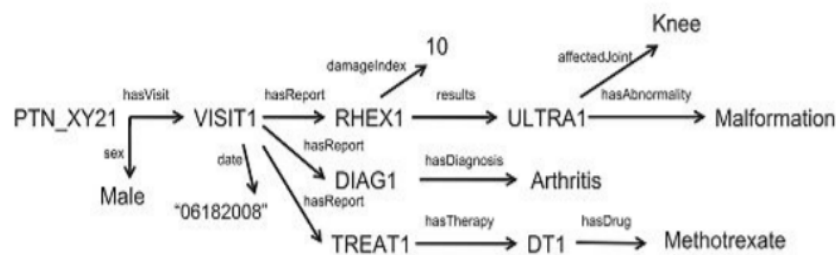


Figure 5: Fragment of the semantic annotations graph

Subject	Predicate	Object	Feature
PTN_XY21	(VISIT1, RHEX1, ULTRA1)	Malformation	Disease
PTN_XY21	(VISIT1, RHEX1)	RHEX1	Report $\cap$ $\exists$ damageIndex
PTN_XY21	(VISIT1, TREAT1, DT1)	Methotrexate	Drug
...	...	...	...

Figure 6: Subjects-subject-object triples

In this method, only important instances (i.e., features) are "extracted" and combined from the entire repository and embedded in regular transactions, capturing implicit data at the schema level in the ontology. You can then apply existing associative rules search algorithms. This type of search rule will become increasingly valuable in future research on both machine learning and SW data. In the future, you can apply generalized query schemes using ontology axioms, as well as automatically detect important instances and search for associative rules. In addition, the method can be used in a variety of

scenarios where mining tasks are transaction-oriented. The problem of the research is the use of data in the ontology to filter and narrow the identified rules, as well as to express the goals of the user. Another important area worth researching is the combination of clustering mining algorithms and associative rules. Previously, this technique has been implemented through hierarchical clustering based on a set of subjects (FIHC). Basically, the FIHC algorithm generates clusters from frequent sets of elements, which in turn constitute cluster descriptors. A new approach could be an algorithm based on finding frequent pairs of objects, which provides more homogeneous clusters and better descriptions than those obtained from FIHC. Also, many studies involve the use of more complex algorithms for data extraction and the formation of transactions from them, the study of their efficiency. No less interesting is the development of new algorithms for data exchange, which are based on semantically enriched elements of the generated transactions.

4. Algorithms for searching Associative Rules

This section will take a closer look at both the well-known AR search algorithms and the new AR search algorithm in semantic networks.

4.1. Traditional search algorithms of AR

AIS algorithm. The first algorithm for finding associative rules, called AIS, was developed by IBM Almaden Agrawal, Imielinski, and Swami in 1993. From this work began an interest in associative rules; in the mid-90s of the last century came the peak of research in this area, and since then every year there are several new algorithms. In the AIS algorithm, candidates for multiple sets are generated and counted "on the fly" while scanning the database [1, 17, 21].

SETM algorithm. The creation of this algorithm was motivated by the desire to use the SQL language to calculate frequent sets of goods. Like the AIS algorithm, SETM also generates candidates "on the fly" based on database transformations. To use the standard SQL join operation to form a candidate, SETM separates the candidate formation from their count.

The inconvenience of AIS and SETM algorithms is the excessive generation and calculation of the Support value of too many candidates, which as a result are not provided often. To improve their performance, the Apriori algorithm was proposed. [7, 13]

Apriori algorithm. The work of this algorithm consists of several stages - the formation of candidates and the counting of candidates. Candidate generation is the stage at which the algorithm, by scanning the database, creates many i -th candidates. At this stage, their Support value is not calculated. Candidate counting is the stage at which the Support value of each i -th candidate is calculated. Candidates whose Support value is less than the minimum value set by the user (min Support value) are also rejected here. The other i -th sets will be the ones that are often found in the database - that is, if the set $\{A, B\}$ is common, then the sets $\{A\}$, $\{B\}$ will also be common. This property is the Support value property (Formula 4): [2, 5, 8]

$$\forall X, Y: (X \subseteq Y) \Rightarrow \text{Support value}(X) \geq \text{Support value}(Y), \quad (4)$$

where X, Y - sets of elements.

Looking at the algorithm of simple search of values, in it there are 2^n variants of sets at the given n elements (Figure 7). Suppose we have a set (AB) with a low value Support value - the Apriori algorithm "cuts off" AB and its derivative sets, thereby accelerating (Figure 8).

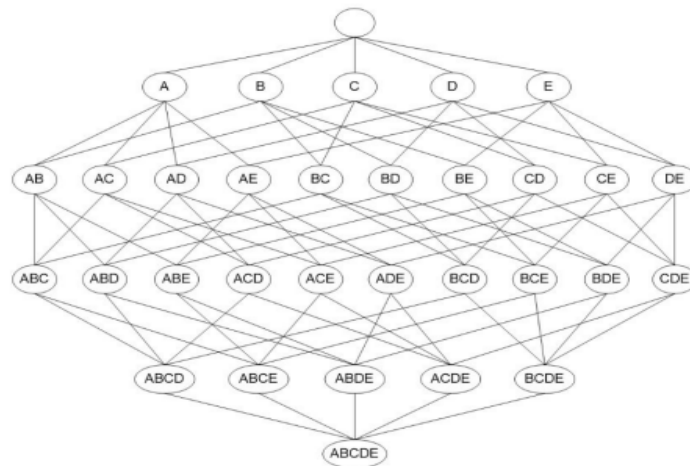


Figure 7: Number of sets (2^n) for these elements

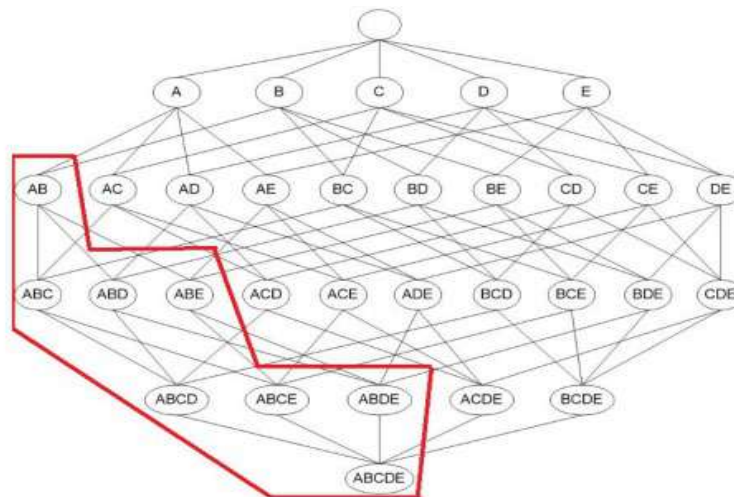


Figure 8: Clipping sets with low Support value

Let's consider the Apriori algorithm on an example, for this purpose we will change a little and we will expand our Table 1 with data (Table 2):

Table 2

Additions to table 1 by associative rules

ID	Items
001	A, B, C
002	B, C
003	B, A, D, C
004	E, D
005	A, B, C, D
006	F

Set min Support value = 3 (Figure 9).

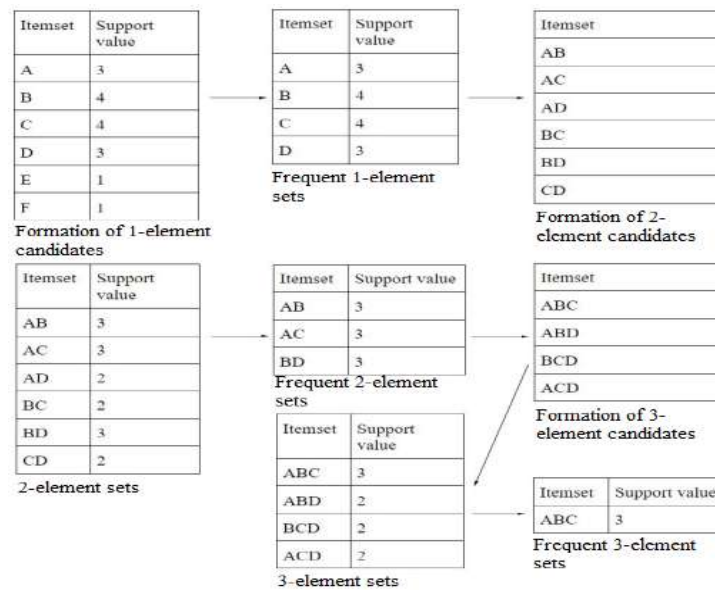


Figure 9: Apriori algorithm

At the first stage (Figure 9), there is a formation of 1-element candidates. Next, the algorithm calculates the Support value of 1-element sets. Sets with a Support value less than the specified (in our case 3) are cut off. In the example, these are sets E and F, which have *Support value* = 1. The remaining sets of elements are considered to be common: A, B, C, D.

Next is the formation of 2-element candidates, counting their Support value and cutting off sets from *Support value* < 3. The remaining 2-element sets AB, AC, BD, participate in the further work of the algorithm.

Continuing the work, the algorithm at the last stage forms 3-element sets of goods: ABC, ABD, BCD, ACD, calculates their Support value and again cuts off sets from *Support value* < 3. The result - a set of ABC products is the most common (Figure 9).

Among the varieties of the Apriori algorithm are the following:

- **AprioriTID.** The peculiarity of this algorithm is that the database of elements is not used to calculate the Support value of the recruitment candidates after the first step. For this purpose, the candidate coding performed in the previous steps is used. In the following steps, the size of the encoded sets can be much smaller than the database itself, thus saving significant resources.
- **AprioriHybrid.** Analysis of the running time of the Apriori and AprioriTID algorithms shows that in earlier steps Apriori achieves better speed than AprioriTID; however, AprioriTID works better than Apriori at later steps. In addition, they form the same procedure for candidate sets. Based on this observation, the AprioriHybrid algorithm is proposed to combine the best properties of the Apriori and AprioriTID algorithms. AprioriHybrid uses the Apriori algorithm in the initial steps and moves to the AprioriTID algorithm when large sets of memory can be used. However, switching from Apriori to AprioriTID requires resources.

Some authors have proposed other algorithms for finding associative rules, which were also improvements to the Apriori algorithm. One of them is the DHP algorithm, also called the hashing algorithm (proposed by J. Park, M. Chen and P. Yu, 1995). Based on its probabilistic calculation of sets of candidates, valid for reducing the count of candidates for the duration of the Apriori algorithm. The reduction is provided by the fact that of the k-element sets of candidates, in addition to the step of reducing the passage of the hashing step. In the algorithm at the k-1 stage during the selection of the candidate, the so-called hash table is created. Each hash table entry is a counter of all reference values of k-element sets that correspond to a row in the hash table. The algorithm uses this information on the k-th to reduce the elements of the candidate sets. After reducing the subset, as in Apriori, the algorithm can delete the candidate set if its value in the hash table is less than the specified Support value.

Also to other advanced algorithms: PARTITION, DIC, the algorithm of "sample analysis". PARTITION algorithm (proposed by A. Savasere, E. Omiecinski and S. Navathe, 1995). This algorithm

of partitioning (division) is contained in the database of scanning operations through the section of its section, each of which can fit in RAM. In the first place in each of the sections using the Apriori algorithm displays sets that are common. The second confirms the importance of supporting each such set. Thus, the stages are available on all other common data sets. To compare the operation of the algorithms, it would also be essential to analyze the DIC algorithm (Dynamic Itemset Counting), which was proposed by scientists S. Brin R. Motwani, J. Ullman and S. Tsur in 1997, divides the database into several blocks, each of which is marked by so-called "start points", and then cyclically scans the database.

4.2. Search algorithm of AR in semantic networks and data

Until now, AR search methods have been applied to traditional data in tabular format or on the basis of graphs. This section explores the problem of finding rules in semantic web data and proposes a new approach to finding APs directly from semantic web data. This approach takes into account the complex nature of semantic web data in contrast to traditional data and, in contrast to existing methods, eliminates the need to convert data and involve end-users in the search process. In trying to apply this search to atypical data, we encounter certain problems and differences compared to traditional ones:

- Heterogeneous: traditional mining algorithms work with homogeneous data sets in which instances are stored in a well-ordered system, and each instance has predefined attributes. But the semantic data are heterogeneous. This means that specific instances of categories / domains (e.g. people, cars, medicines, etc.) based on the same ontology or individual ontologies may have different characteristics.
- There is no clear definition of transactions: in conventional information systems, data is stored in databases using predefined structures, and these structures can be used to recognize transactions and thus extract them from the data set. Then the traditional AR search algorithms process these transactions. For example, in the case of a "buyer basket", transactions are formed from products that are purchased together, and these products will have the same ID as the transaction ID. Conversely, in a semantic network, different attributes for an instance may be formed at different times, and therefore an instance may have an attribute that does not exist in another instance of the same type.
- Multiple relationships between entities: Traditional AR search algorithms to generate large sets of elements take into account only the values of the objects and assume that there is only one type of relationship between the entities (for example, purchased together). But in semantic data, there are many relationships between entities. In fact, predicates are relations between two entities or between one entity and one value.

Because semantic annotations are encoded in RDF / (S) and OWL, you should extend SPARQL with new elements that can specify a search pattern. The syntax is somewhat similar to Microsoft Data Mining Extension (DMX), which is an SQL extension for working with DM models in Microsoft SQL Server. The extended SPARQL grammar is shown in Figure 10 and Figure 11 shows an example of the SPARQL view for AR search schemes (Formula 5):

$$Q = (Patient, Report, \{Disease, Drug, Report \cap \exists damageUndex\}). \quad (5)$$

```

Query ::= Prologue(SelectQuery|ConstructQuery|DescribeQuery|AskQuery|MiningQuery)
MiningQuery ::= CREATE MINING MODEL Source '{'
                Var 'RESOURCE' 'TARGET'
                (Var('RESOURCE'|'LITERAL')
                 'MAXCARD1'? 'PREDICT'? 'CONTEXT'?)+'}'
                DatasetClause* WhereClause UsingClause MeasuresClause
UsingClause ::= 'USING' SourceSelector BracketedExpression
MeasuresClause ::= 'ADD MEASURE' measure +

```

Figure 10: Advanced SPARQL grammar for the CREATE MINING MODEL query

The SPARQL query has been expanded by adding a new character called MiningQuery. The body of the query consists of variables that the user targets when searching for data. Next to each variable, we define its type: RESOURCE for variables that contain RDF, and LITERAL for those that have regular data types. In case we want to find patterns with only one variable, we add the keyword MAX-CARD1 to the variable. By default, found templates can contain more than one occurrence of each variable. In addition, we define the "sequence" of this rule by adding the PREDICT keyword (optional). Finally, the TARGET keyword refers to the analyzed resource, which should be an ontology concept. The purpose of the analysis determines the set of rules obtained. In WhereClause, we specify restrictions on previous variables. The advantage is that the user's knowledge of the ontology structure is not required. Therefore, users only need to specify the type (concept of ontology) to which the variables refer.

```
CREATE MINING MODEL {
    ?patient RESOURCE TARGET
    ?drug RESOURCE
    ?jadi LITERAL
    ?disease RESOURCE PREDICT
    ?report RESOURCE CONTEXT
}
WHERE {
    ?patient rdf:type Patient .
    ?drug rdf:type Drug .
    ?disease rdf:type Disease .
    ?report rdf:type Report .
    ?report damageIndex ?jadi .
}
USING Apriori (SUPPORT = 0.06, CONFIDENCE = 0.7)
ADD MEASURE lift, leverage
```

Figure 11: Extended view of the SPARQL CREATE MINING MODEL query

Let the user choose the patient as the desired "concept" of the analysis. The set of characteristics that will make up the transaction includes diagnosed diseases, prescribed drugs and damage rate. Finally, the transaction will be based on the details of the report, i.e. the transactions will not include the characteristics in all reports in general, but only the characteristics of each doctor's report.

The variable jadi refers to the index of injury to the patient's joint, the user specifies the report and damageIndex as a property of the resource and the type of data from which they can be obtained. UsingClause defines the name and parameters of the algorithm.

Because we do not ask the user to specify the exact relationship, the query model introduces some ambiguity about the elements that perform the transaction. Thus, the same conceptual changes (selected features) can be used under different contexts of ontology. For example, Disease can diagnose the patient's own illness or the illness of a family member. This ambiguity becomes a problem in determining what the intentions of the users really are. In fact, the user can use this ambiguity by specifying in the extended SPARQL query to understand the ontology using the "triplets" WHERE. However, this task can be cumbersome. For the query to be really correct, the user can select the desired context using CONTEXT added to the corresponding concept. In addition, the system will build transactions, taking into account all possible contexts.

Recalling the form of subject-predicate-object triplets (Fig. 2.2.4), they will be useful for the above-mentioned AR search scheme $Q = (Patient, Report, \{Disease, Drug, Report \cap \exists damageIndex\})$. Instances of RHEX1, RHEX1, TREAT1 will belong to the concepts of context Q. The transactions of the elements obtained from these three compositions are shown in Figure 12:

```
Transactions
1. {RheuExam.Disease.Malformation, RheuExam.RheuExam.damageIndex → 10}
2. {Treatment.Drug.Methotrexate}
3. {RheuExam.Disease.Malformation, RheuExam.Disease.BadRotation,
RheuExam.RheuExam.damageIndex → 15}
4. {Treatment.Drug.Methotrexate, Treatment.Drug.Corticosteroids}
```

Figure 12: Transactions of elements of triples of compositions

The algorithm itself will follow the following steps:

- First, compute sets of common elements that reduce their total number (especially when there are a large number of transactions).
- Then the sets of elements are truncated by the method described in subsection 4.1 (Figure 9) with *Support value* < 0.7 to filter out those that combine frequent and rare elements. These transactions are usually false.
- Finally, you can get rules from element sets by specifying min *Confidence value* = 0.8.

5. Results of research and experiments

To ensure that the algorithm is correct and relevant, it will be tested on real-world OWL instances of patient observation. According to the example in Fig. 1b, annotations were formed based on the Health-e-Child (HeC) project. These annotations correspond to the ontology of the project. The structure of semantic annotations is very heterogeneous and contains information about 588 patients classified into three different groups according to their disease: juvenile idiopathic arthritis (JIAPatient), heart disease (CardioPatient) and neurological disease (NeuroPatient). The total number of semantic annotations is 629,000, which is an average of more than 1,000 annotations per patient.

To avoid errors, query schemes were automatically generated for 12 different concept concepts (disease, treatment, medication, ...) and 3 concepts for contexts: patient, visit and report. It is worth noting that the Report concept has 20 sub-concepts that correspond to the various clinical reports of the HeC project. The current implementation of transaction extraction has been developed on the basis of the ontology indexing system, which also provides a simple mechanism for creating ontological indexes. To confirm the relevance and results of the found transactions, there is a range of different AP search algorithms, among which genetic algorithms (GA) for AP search have recently been proposed.

On the Table 3 shows three selected contexts for experiments, as well as the number of generated transactions and their average length.

Table 3

Results of experiments on the number of generated transactions and their average length

Context	Transactions	Medium length
Patient	588	29.57
Visit	1458	12.84
Report	3608	5.24

The number and nature of transactions received in each context are completely different and will therefore affect the rules created. More general contexts tend to generate longer transactions, which in turn increases the likelihood of obtaining more rules. Instead, more specific contexts generate smaller transactions, which narrows the scope for detecting rules. This discrepancy in the nature of transactions necessitates adequate adjustment of the minimum Support value threshold of each set-in order to be able to find the association rules.

In the Table 4 shows the number of created rules together with their average Confidence value, average Lift value and ϕ -coefficient for three sets of transactions.

Table 4

Results of experiments on the number of generated transactions and their average length

Context	Min Support value	Amount of rules	AVG Confidence value	AVG Lift value	ϕ -coefficient	Correlation
Patient (588)	0.187	109	0.993	2.944	0.796	0.678

Visit (1458)	0.047	93	0.976	9.975	0.865	0.480
Report (3608)	0.017	151	0.964	27.69	0.836	0.169

All created rules have a high Confidence value. In addition, the more limited the context, the better the rules are formed. Moreover, the ϕ -coefficient shows a strong correlation in all cases.

In the Table 5 shows the effect of applying certain restrictions (i.e. selecting only specific report types) in the search template.

Table 5

The result of the impact of the application of the imposed restrictions in the search template

Transactions	Amount of rules	AVG Confidence value	AVG Lift value	% of report rules
All reports (588)	655	0.943	4.125	83
Rejected 5 (585)	22	0.963	6.966	56
Rejected 12 (438)	26	0.934	6.652	35

Each line displays received transactions and rules for all reports, canceling the 5 most common reports and discarding the 12 most common reports in the patient context. This table also includes the percentage of rules that contain items from different reports.

Figure 13 analyzes the coverage of the formed rules by different thresholds of support as an indicator of their quality.

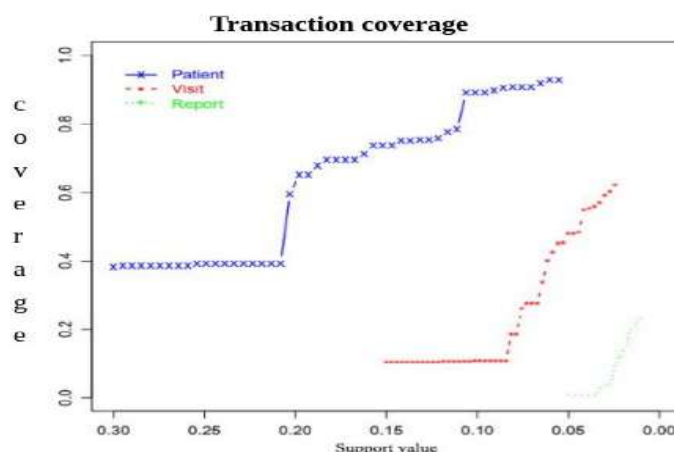


Figure 13: Coverage of transactions achieved by generated rules with different minimum threshold levels Support value

The rules obtained from the Patient set achieve good coverage with relatively high thresholds of Support value. However, other received sets of rules are not able to confirm the high percentage of transactions. The fact is that the length of other sets of transactions is shorter, which usually reduces the number of detected AR. In the case of the Report transaction set, the coverage is even less because the transactions are derived from different types of reports. Therefore, good rules may arise, but with very low Support value thresholds. In these cases, it would be advisable to use more sophisticated AR search algorithms that are not based on the concept of Support value.

In Figure 14 shows the average Confidence value of the generated rules with different threshold Support value.

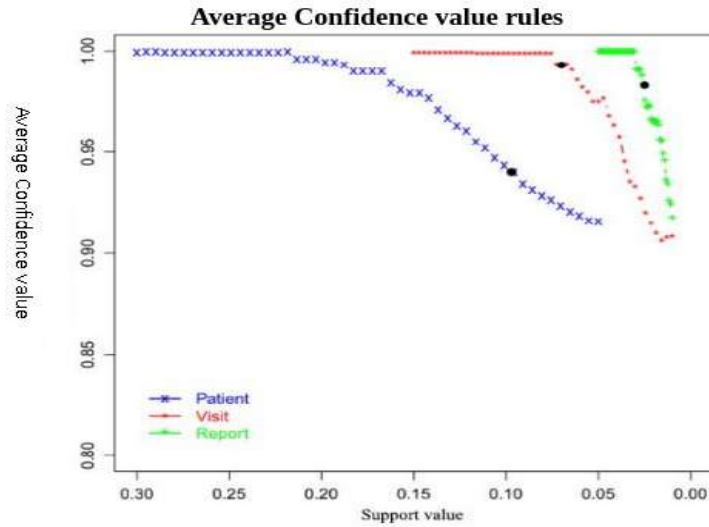


Figure 14: Average Confidence value of generated rules with different thresholds Support value

The support value for the three sets of transactions remains high even for low thresholds, which confirms the quality of the rules.

Based on the two previous measures (coverage and average Confidence value) we can select a minimum Support value threshold for each set of transactions and further analyze the quality of the rules obtained. In Figure 15 shows the coating, multiplied by the average Confidence value. For each set of transactions, the Support value threshold is selected, at which both measures are maximally involved.

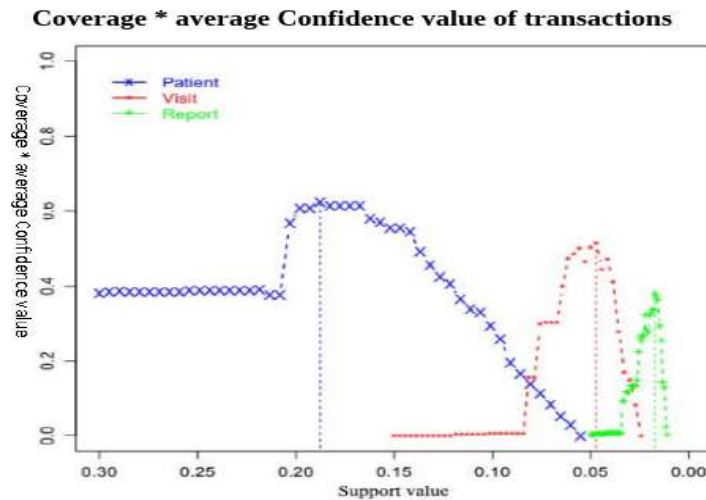


Figure 15: Coverage multiplied by the average Confidence value of the generated rules with different Support value thresholds

Finally, the three Figures 16, 17 and 18 show an example of AR obtained in the context of Patient, Visit and Report.

Rule	Support value	Confidence value	Lift value
{JIAPatient. (Disease.disease)→ oligoarthritis, JIAPatient. (Finding.sacroiliac_ tenderness & Anatomy.sacroiliac_ joint & MedicalProcedure.p resence)} ⇒ {JIAPatient. (Symptom.inflammatory_ pain)}	0.26	1.0	18.833
{CardioPatient. (MedicalProcedure. genetic_research & ExternalActivities.ge netic_molecular), CardioPatient. (Molecular.tbx5)} ⇒ {CardioPatient. (Molecular.5_nkx2)}	0.107	1.0	9.187
{NeuroPatient. (ExternalActivities.e ndocrinology)} ⇒ {NeuroPatient. (MedicalProcedure.r esonance_magnetic_ material_imaging_c ontrast &Chemical.metal)}	0.107	0.969	8.380
{JIAPatient. (Symptom.tenderne s & Quantity.score)→1} ⇒ {JIAPatient. (Finding.pain)→1}	0.144	1.0	6.917

Figure 16: The result of an associative rule obtained in the context of the Patient

When considering the results, the Support value can also be interpreted as a percentage, for example for the first rule it was 0.260 - this means that sets of diseases in the rule occur in 26% of transactions. In many ARs found in the table above, the Confidence value is close to 1. Considering the first rule in the example, this means that in 100% of transactions a patient with oligoarthritic and lumbar pain has active tissue inflammation in this department. Lift value characterizes how good the prediction is and how interdependent the factors are. The greater the value, the greater the dependence of factors, i.e. how much the presence of one factor affects another. With low Lift values, on the contrary, the lower it is, the greater the negative effect one factor has on another.

Rule	Support value	Confidence value	Lift value
{CardioVisit. (MedicalProcedure. auscultation_lung), CardioVisit. (Quantity.compensa tion)} ⇒ {CardioVisit. (Finding.breath_so unds & Quality.equality)}	0.074	1.0	11.950
{JIAVisit. (Chemical.methotrexate)→Weekly} ⇒ {JIAVisit. (Chemical.methotrexate)}	0.08	0.990	7.197

Figure 17: The result of an associative rule obtained in the context of Visits

In this case, for example, the first rule shows good results in Confidence and Lift value. When listening to the patient's lungs with a stethoscope, the doctor can better judge the presence of problems with them and better assesses their general condition.

Rule	Support value	Confidence value	Lift value
{JLADiagnosis. (Finding.fever)} ⇒ {JLADiagnosis. (Finding.erythematous_rash)}	0.010	0.812	57.439
{Surgery. (Finding.surgery_performed), Surgery. (MedicalProcedure.resonance_magnetic_material_imaging_contrast & Chemical.metal)} ⇒ {Surgery. (MedicalProcedure.tumour_removal)}	0.012	0.815	46.137
{JLALaboratoryOPBG. (Immunology.antibody_antinuclear)} ⇒ {JLALaboratoryOPBG. (Immunology.rheumatoid_factor)}	0.0133	0.85	37.951

Figure 18: The result of an associative rule obtained in the context of Reports

This Figure 18 has some really interesting context-based rules Report. At observation at the patient of a fever in most cases it specified on the appearance of erythema, indicating a complex inflammation of the joint or tissues.

Most of the operations before which magnetic resonance imaging is performed with using additional chemical compounds of iron, were just for removal tumors. However, with high Confidence and Lift value, these rules are low Support value, which indicates the small number of occurrences of these sets in transactions.

6. Conclusions

Summarizing all the above, research in this work was directed to consider the basic concepts and search for AR in both traditional ways and in inhomogeneous data of semantic networks, which creates certain problems when using existing algorithms. It is worth noting that one of the most popular areas of application for AR search still remains consumer and marketing. Semantic network data in most cases serve for a specific field and are highly specialized. Probably that's why direction you can do a lot of interesting research, one of which is the search associations among heterogeneous data.

A new method for finding ARs from inhomogeneous ones was also presented data in semantic networks expressed in RDF / (S) and OWL. Previously, this problem considered only to a small extent. Experiments on real SW data show good results. An interesting problem for future work is data mining in the ontology for filtering and cutting off the detected rules. Yet one important area that can be considered in the future concerns combination of clustering and AR search algorithms for generalization of arrays documents. This technology has previously been implemented to some extent hierarchical clustering of sets (FIHC). Basically, the FIHC algorithm generates clusters of sets of elements, which, in turn, make up the cluster descriptors. A new approach based on hierarchy has also recently been proposed element sets, which provides more homogeneous clusters and better descriptions than those obtained from FIHC. Undoubtedly, each algorithm can be improved and improve, apply better ways of embedding data to generated transactions and study their effectiveness. It is no less interesting development of new data exchange algorithms based on SW data and are accelerated by new ways of processing the generated transactions.

7. References

- [1] C. Giannella, H. Jiawei, P. Jian, Mining frequent patterns in data streams at multiple, in : ESMA 2018 IOP Conf. Series: Earth and Environmental Science, 2019, pp. 61-84. doi:10.1088/1755-1315/252/3/032219
- [2] B. Patel, Vishal H. Bhemwala, A. Patel, Analytical, Study of Association Rule Mining Methods in Data Mining. International Journal of Scientific Research in Computer Science Engineering and Information Technology, Vol. 3, 2018, pp. 818-831. DOI:10.32628/CSEIT1833244
- [3] N. Boyko, K. Kmetyk-Podubinska, and I. Andrusiak, Application of Ensemble Methods of Strengthening in Search of Legal Information, in: Lecture Notes on Data Engineering and Communications Technologies, Vol. 77, 2021, pp. 188-200. https://doi.org/10.1007/978-3-030-82014-5_13
- [4] P. Jian, H. Jiawei, B. Mortazavi-Asl, PrefixSpan: mining sequential patterns efficiently by prefix-projected pattern growth, in : Proc of the 17th International Conference on Data Engineering, 2001, pp. 215-224.
- [5] N. Boyko, N. Tkachuk, Processing of Medical Different Types of Data Using Hadoop and Java MapReduce, in: The 3rd International Conference on Informatics & Data-Driven Medicine (IDDM 2020), Växjö, Sweden, November 19-21, 2020 pp. 405-414.
- [6] L. Duanyang, F. Jian, L. Xiaofan, A frequent sequence pattern mining algorithm based on logic. Journal of Computer science, Vol. 42 (5), 2015, pp. 260-264.
- [7] W. Xindong, X. Fei, H. Yongming, A sequence pattern with wildcards and one-off conditions Dig. Software journal, Vol. 24(8), 2013, pp. 1804-1815.
- [8] Zh. Jialu, Y. Jun, H. Jing, Constrained association rule mining based on a set of transaction ids Mining algorithm. Computer engineering and design, Vol. 34(5), 2013, pp. 1663- 1667. doi:10.1088/1755-1315/252/3/032219.
- [9] N. Boyko, A look through methods of intellectual data analysis and their applying in informational systems, in: XIth International Scientific and Technical Conference Computer Sciences and Information Technologies (CSIT), 2016, pp. 183-185.
- [10] H. Hong Juan, Zh. Jian, Ch. Shaohua, Constraint maximum frequent itemset mining based on frequent pattern trees Algorithm. Computer engineering, Vol. 37(9), 2011, pp. 78-80.
- [11] W. Xiuzhi, Association classification algorithm based on intelligent optimization of support and confidence. Computer applications and software, Vol. 30(11), 2013, pp. 184-186.
- [12] Ch. Shutong, X. From Rich, B. Hongwei, Research on efficient privacy protection frequent pattern mining algorithm. Computer science, Vol. 42(4), 2015, pp. 194-198.
- [13] Ch. Aidong, L. Guohua, F. Fan, Uncertain data association rules that satisfy uniform distribution Mining algorithm. Computer research and development, Vol. 50, 2013, pp. 186-195.
- [14] Sh. Yan, D. Min, L. Qiliang, Mining methods for association rules of Marine and continental climate events. Journal to Ball information science, Vol. 16(2), 2014, pp. 182-189.
- [15] R. Idoudi, K. S. Ettabaa, B. Solaiman, K. Hamrouni, Ontology Knowledge Mining Based Association Rules Ranking, in: Procedia Computer Science Published by Elsevier, Vol. 96, 2016, pp. 345-354. DOI: 10.1016/j.procs.2016.08.147
- [16] R. Paul, T. Groza, J. Hunter and A. Zankl, Semantic interestingness measures for discovering association rules in the skeletal dysplasia domain. Journal of Biomedical Semantics Vol. 5(8), 2014, pp. 2-13.
- [17] O. Daramola, A. Ibukun, O. Okuboyejo, Semantic association rule mining in text using domain ontology. International Journal Metadata Semantic Ontology, Vol. 12(1), 2017, pp. 12-28.
- [18] M. Barati, Q. Bai, Q. Liu, Mining semantic association rules from RDF data. Knowledge Based System, Vol. 133, 2017, pp. 183–196.
- [19] L. Galárraga, C. Teflioudi, K. Hose, F.M. Suchanek, Fast rule mining in ontological knowledge bases with AMIE++. The International Journal on Very Large Data Bases, Vol. 24, 2015, pp. 707–730.
- [20] L.A. Galárraga, C. Teflioudi, K. Hose, F. Suchanek, AMIE: Association Rule Mining Under Incomplete Evidence in Ontological Knowledge Bases, in: Proceedings of the 22nd International

Conference on World Wide Web, WWW'13, Rio de Janeiro, Brazil, 13–17 May 2013; ACM: New York, NY, USA, 2013; pp. 413–422.

- [21] K. A. Kale, R.P. Sonar, Review on Mining Association Rule from Semantic Data. International Journal of Computer Science and Information Technologies, Vol. 7 (3), 2016, 1328-1331

The Pseudorandom Key Sequences Generator Based on IV-Sets of Quaternary Bent-Sequences

Elena Bakunina^a and Artem Sokolov^b

^a National University "Odessa Law Academy", Fontanska road, 23, Odessa, 65000, Ukraine

^b Odessa Polytechnic National University, Shevchenko Avenue, 1, Odessa, 65044, Ukraine

Abstract

The development of modern cryptography, steganography, and communication systems with noise-like signals often requires the use of non-binary generators of pseudorandom key sequences. The main requirements for such a cryptographic construct are the compliance of gamma generated by them with the criteria of NIST stochastic tests and the implementation of the concepts of diffusion and confusion with respect to the generated gamma and the cryptographic key. In this paper, we empirically confirm the prospects of combining the advantages of quaternary bent-sequences, which are characterized by the highest nonlinearity value, with the advantages of the linear feedback shift registers, which are characterized by a high stochastic quality of the generated sequences. We present the scheme of a pseudorandom key sequences generator based on binary linear feedback shift registers and the IV-sets of quaternary bent-sequences. It has been established that the presented scheme generates pseudorandom key sequences corresponding to all the NIST stochastic tests, while the maximum values of the nonlinearity of quaternary bent-sequences provide a high level of implementation of the concept of confusion. At the same time, construction of IV-sets of quaternary bent-sequences was found, which provides the best stochastic properties of the generated gamma. The developed generator is characterized by high values of the number of protection levels and the simplicity of the algorithmic implementation. Compared to combining two binary generators of pseudorandom key sequences based on dual sets of binary bent-sequences, 40% fewer linear feedback shift registers are required to ensure the operation of the developed generator. The practical application of the developed generator is justified in systems that require many-valued logic pseudorandom key sequences.

Keywords

Pseudorandom key sequences generator, bent-sequence, many-valued logic.

1. Introduction and statement of the problem

The Pseudorandom Key Sequences Generator (PRKSG) is a very important construct used in modern cryptographic applications as well as in other areas of science and technology: PRKSG is the basis for the functioning of modern stream encryption algorithms, organizing modern efficient modes for block encryption algorithms [1], defining the steganographic path in modern steganographic algorithms, functioning of modern radio communication systems based on Frequency-Hopping Spread Spectrum (FHSS) technology [2], etc.

Today, there are many different layouts for the PRKSG construction: the classical scheme based on Linear Feedback Shift Register (LFSR) with the use of nonlinear elements [3, 4], schemes based on the theory of dynamic chaos [5], cellular automata [6], etc. Regardless of the chosen scheme, the following basic requirements must be applied to PRKSG being developed today: they must be characterized by the high level of stochastic quality (the level of which is measured by the compliance with a generally accepted set of NIST stochastic tests [7]), the ability to generate pseudorandom key

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: elenabakunina72@gmail.com (E. Bakunina); radiosquid@gmail.com (A. Sokolov)

ORCID: 0000-0002-5700-7321 (E. Bakunina); 0000-0003-0283-7229 (A. Sokolov)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org) Proceedings

sequences based on a short cryptographic key (at the same time, the PRKSG must provide a high level of implementation of Shannon's concepts of diffusion and confusion [8] between the key element and the generated gamma), the simplicity of the algorithmic implementation, and the high performance of the PRKSG algorithm.

The classical LFSR-based PRKSG layout with a nonlinear element fully fulfills the mentioned requirements, in particular, its modification based on bent-sequences proposed in [9].

The development of the theory of quantum cryptography, as well as the recently proposed cryptographic algorithms based on many-valued logic functions [10], have led to the need to create PRKSG operating on a non-binary alphabet $\bar{A} = \{0, 1, \dots, q-1\}$, $q \neq 2$. Thus, the scheme of the PRKSG based on LFSR and dual sets of bent-sequences was generalized to operate over the alphabet $\bar{A} = \{0, 1, 2\}$, which corresponds to the problems of quantum cryptography.

However, the vast majority of information today is stored, processed, and transmitted in binary form, which makes it especially relevant to consider the possibility of constructing PRKSG that would operate over the alphabets $\bar{A} = \{0, 1, \dots, 2^k - 1\}$, primarily over the quaternary alphabet $\bar{A} = \{0, 1, 2, 3\}$. The use of the quaternary alphabet fundamentally makes it possible to simplify the PRKSG schemes by generating twice more gamma bits during one cycle, and also increases the possibility of implementing the principles of diffusion and confusion in the PRKSG [8]. The quaternary alphabet also provides a much greater variety of algebraic constructions that can be used to build high-quality PRKSG.

Steganography is another important application of quaternary PRKSG in view of the peculiarities of the choice of the steganographic path: the ability to select one of the three color components or skip a given pixel from the process of information embedding.

These facts substantiate the tasks of further research on the possibilities of improvement of the classical layout for constructing the PRKSG operating over the quaternary alphabet.

The *purpose* of this paper is to develop a quaternary PRKSG based on IV-sets of bent-sequences.

2. The theoretical foundations for the proposed PRKSG

The classic layout of the PRKSG implies the use of LFSR as the main component.

Definition 1 [11]. An LFSR is a shift register consisting of d memory cells, the value of the input element of which is determined by the value of the function constructed in accordance with the primitive irreducible over the Galois field $GF(2)$ polynomial of degree d .

It is known [10] that the use of LFSR makes it possible to achieve good diffusion, and also provides a sufficiently high stochastic quality of the generated pseudorandom key sequences, which will be shown below in Table 2 on the example of LFSR built on the basis of the primitive irreducible polynomial

$$f(x) = x^{83} + x^{46} + x^{45} + x + 1. \quad (1)$$

The disadvantages of LFSR include the fact that they generate pseudorandom key sequences in accordance with a fairly simple rule defined by a primitive irreducible polynomial.

This does not allow them to provide a sufficiently high level of nonlinearity of the relationship between the elements of the short key and the generated pseudorandom key sequence. In other words, these constructions do not provide a sufficient level of confusion. This circumstance leads to the possibility of launching attacks against such a generator, for example, using the Berlekamp-Massey algorithm.

One of the historical attempts to overcome this shortcoming is the Geffe Generator, which provides a significantly higher level of confusion compared to the direct application of the LFSR through the use of several interconnected LFSR units. However, this level of confusion is also insufficient due to the simplicity of the links between the pooled LFSR units.

In modern PRKSG, this drawback is eliminated by using a non-linear element in conjunction with LFSR, which is often represented by the Boolean function with a high level of nonlinearity. Thus, the block diagram of PRKSG in its classical layout is shown in Fig. 1.

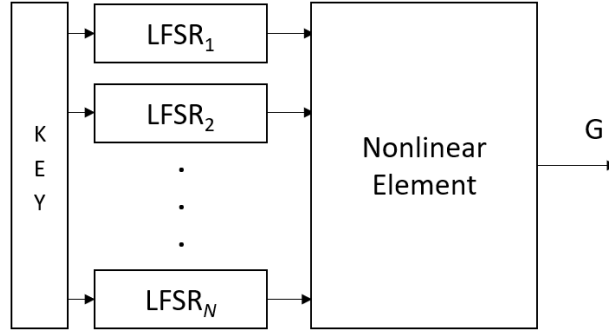


Fig. 1: Block diagram of PRKSG based on LFSR and nonlinear element

The ideal option for nonlinear element is to use such special algebraic constructions as bent-functions (whose truth tables are called as bent-sequences), which have the maximum level of nonlinearity, and, accordingly, provide the highest level of confusion implemented by PRKSG.

In accordance with the definition [9], a binary sequence $B = [b_0, b_1, \dots, b_i, \dots, b_{N-1}]$, where $b_i \in \{\pm 1\}$ are coefficients, of even length $N = 2^{2m}$, $i = 0, 1, \dots, N-1$ is called a bent-sequence if it has uniform absolute values of its Walsh-Hadamard spectrum $W_B(\omega)$, which can be represented in matrix form

$$W_B(\omega) = BA, \quad \omega = 0, 1, \dots, N-1, \quad (2)$$

where A is the Walsh-Hadamard matrix of order N .

The Walsh-Hadamard matrices are constructed in accordance with the Sylvester construction, which is specified using the following recurrent rule

$$A_{2^k} = \begin{bmatrix} A_{2^{k-1}} & A_{2^{k-1}} \\ A_{2^{k-1}} & -A_{2^{k-1}} \end{bmatrix}, \quad A_1 = 1. \quad (3)$$

Despite the successful layout of the block diagram shown in Fig. 1, the practical application of bent-sequences in it is extremely difficult due to bent-sequences imbalance, as a result of which the gamma generated by PRKSG is also unbalanced, and, therefore, does not meet the most basic criteria of stochastic quality.

Another undoubted disadvantage of using binary bent-sequences is their existence only for lengths $N = 2^{2m} = 4, 16, 64, 256, \dots$, while binary bent-functions of an odd number of variables do not exist, which limits the scalability of PRKSG schemes.

In [9], for the binary case, a solution to this problem was proposed based on the use of dual sets of maximally nonlinear bent-functions.

This PRKSG scheme showed excellent characteristics both in terms of performance and in terms of the quality of the generated pseudorandom key sequences, which is confirmed by both a set of stochastic tests [12], as shown in [9], and a set of NIST stochastic tests [7]. The number of protection levels of this scheme reaches the value $Y \approx 2,67 \cdot 10^{49} \approx 2^{165}$, which exceeds the number of protection levels of the AES-128 block symmetric cipher [13].

This scheme was also modified for the ternary case in [14] using LFSR based on primitive irreducible polynomials over the Galois field $GF(3)$, as well as triple sets of bent-sequences. Despite the absence of specific NIST tests for ternary pseudorandom sequences, the sequences generated by PRKSG [13] fully correspond to the set of tests [12], which suggests their high level of quality.

To construct a quaternary PRKSG, it is proposed to use the complete class of quaternary bent-sequences described in [15] and synthesized in [16]. Let us introduce the definition of the Vilenkin-Chrestenson transform that we will need, as well as the definition of the many-valued logic bent-sequence.

Definition 2 [16]. The coefficients of the Vilenkin-Chrestenson transform of a q -valued logic function is the vector obtained by multiplying its truth table T (represented in the exponential form) by the complex conjugate of the Vilenkin-Chrestenson matrix

$$\Omega_A = T \cdot \bar{V}_{16}. \quad (4)$$

while for the case of 4-functions the Vilenkin-Chrestenson matrix is constructed according to the following recurrent rule

$$V_{4^{k+1}} = \begin{bmatrix} V_k & V_k & V_k & V_k \\ V_k & V_k + 1 & V_k + 2 & V_k + 3 \\ V_k & V_k + 2 & V_k & V_k + 2 \\ V_k & V_k + 3 & V_k + 2 & V_k + 1 \end{bmatrix}, \quad (5)$$

where “+” is the operation of addition mod4, the matrices V are represented in symbolic form, i.e. the summation is performed with respect to the indices z_i , and

$$V_4 = \begin{bmatrix} z_0 & z_0 & z_0 & z_0 \\ z_0 & z_1 & z_2 & z_3 \\ z_0 & z_2 & z_0 & z_2 \\ z_0 & z_3 & z_2 & z_1 \end{bmatrix} = \begin{bmatrix} e^{j0} & e^{j0} & e^{j0} & e^{j0} \\ e^{j0} & e^{j\frac{\pi}{2}} & e^{j\pi} & e^{j\frac{3\pi}{2}} \\ e^{j0} & e^{j\pi} & e^{j0} & e^{j\pi} \\ e^{j0} & e^{j\frac{3\pi}{2}} & e^{j\pi} & e^{j\frac{\pi}{2}} \end{bmatrix}. \quad (6)$$

Definition 3 [16]. For the Vilenkin-Chrestenson matrix of order $N = q^k$, a bent-sequence

$H = [h_0, h_1, \dots, h_i, \dots, h_{N-1}]$ is a sequence over the alphabet $h_i \in \left\{ e^{j\frac{2\pi}{q}\nu} \right\}, \nu = 0, 1, \dots, q-1$ if it has a

uniform absolute values of the Vilenkin-Chrestenson spectrum, which can be represented in matrix form

$$|\Omega_H(\omega)| = |H \cdot \bar{V}_N| = \text{const}, \omega = 0, 1, \dots, N-1, \quad (7)$$

where V_N is the Vilenkin-Chrestenson matrix of order N over the alphabet

$$h_i \in \left\{ e^{j\frac{2\pi}{q}\nu} \right\}, \nu = 0, 1, \dots, q-1.$$

Bent-sequences are very important mathematical objects for cryptographic constructions development, however, their main drawback, which limits their use in cryptography, in particular, in the problems of development of cryptographically secure PRKSG, is their imbalance, the inequality of the number of symbols $0, 1, \dots, q-1$ contained in them, i.e. $K^0 \neq K^1 \neq \dots \neq K^{q-1}$. This drawback was solved in [9] and [14] by using dual sets and triple sets of bent-sequences respectively, while the definition of dual sets and triple sets was generalized in [16], resulting in the definition of a q -set of bent-sequences which was presented.

Definition 4 [16]. A set of q q -ary bent-sequences is called a q -set if the concatenation of their truth tables is balanced, i.e. it satisfies the relation $K^0 = K^1 = \dots = K^{q-1}$.

The research of the complete class of quaternary bent-sequences of length $N=16$ makes it possible to distinguish the following weight structures, which are presented in Table 1 in the following form $\{K^0, K^1, K^2, K^3\}$.

Table 1

Weight structures of quaternary bent-sequences of length $N=16$

Weight structure	J	Weight structure	J	Weight structure	J	Weight structure	J
{10,0,6,0}	192	{4,6,0,6}	2112	{8,2,4,2}	6336	{3,3,7,3}	10240
{0,10,0,6}	192	{6,4,6,0}	2112	{2,8,2,4}	6336	{3,3,3,7}	10240
{6,0,10,0}	192	{1,5,5,5}	6144	{4,2,8,2}	6336	{2,4,6,4}	25152
{0,6,0,10}	192	{5,1,5,5}	6144	{2,4,2,8}	6336	{4,2,4,6}	25152
{0,6,4,6}	2112	{5,5,1,5}	6144	{7,3,3,3}	10240	{6,4,2,4}	25152
{6,0,6,4}	2112	{5,5,5,1}	6144	{3,7,3,3}	10240	{4,6,4,2}	25152

Based on the data presented in Table 1, as well as on the Definition 4, it is fundamentally possible to construct 3948 IV-sets, which can be formed on the basis of a complete class of quaternary bent-sequences of length $N=16$. As the experiments show, the best stochastic quality of the generated pseudorandom key sequences is provided with the following choice of their structure

$$[H \quad (H+1) \bmod 4 \quad (H+2) \bmod 4 \quad (H+3) \bmod 4], \quad (8)$$

where H is one of the quaternary bent-sequences of length $N=16$, the cardinality of the complete class of which is $J=200704$.

Note that theoretically, as a LFSR, it is possible to use registers based on primitive irreducible polynomials over the extended field $GF(2^2)$. A method for synthesizing such a primitive irreducible polynomials over extended fields is presented in [17]. Based on this method, for example, we can construct following polynomial

$$f(x) = x^{10} + x^9 + 3x^8 + 3x^7 + 2x^6 + x^4 + 3x^3 + x^2 + 2, \quad (9)$$

on the basis of which the LFSR scheme shown in Fig. 2 can be built.

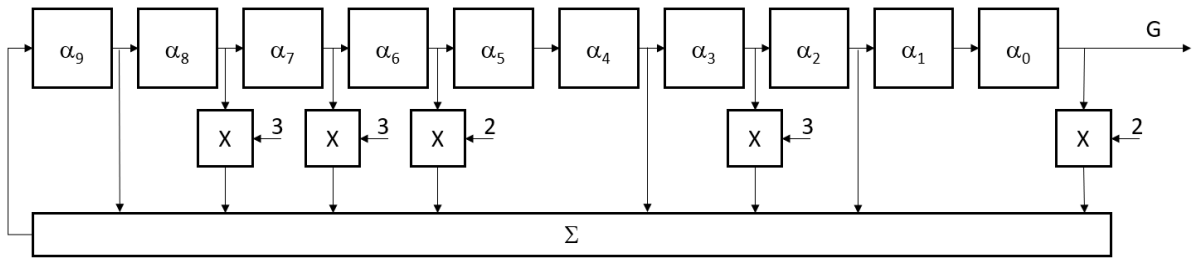


Fig. 2: LFSR scheme based on primitive irreducible polynomial (9)

Note that the operations of multiplication and summation in the LFSR scheme shown in Fig. 2 are performed in the extended Galois field $GF(2^2)$, which arithmetic is determined by a single irreducible polynomial over the Galois field $GF(2)$ of second degree $f(x) = x^2 + x + 1$ in accordance with the following tables

$+$	0	1	2	3	$\cdot$	0	1	2	3
0	0	1	2	3	0	0	0	0	0
1	1	0	3	2	1	0	1	2	3
2	2	3	0	1	2	0	2	3	1
3	3	2	1	0	3	0	3	1	2

(10)

However, the results of the performed research which are presented in Table 2 show that LFSR constructed using extended fields show significantly worse stochastic performance compared to LFSR constructed using prime Galois fields. This circumstance does not allow them to be applied in the developed quaternary PRKSG. In our opinion, a good solution is to use binary LFSR with their subsequent multiplexing to ensure that arguments are supplied to the input of the IV-set of bent-sequences, as well as the choice of a specific quaternary bent-sequence from the IV-set is performed.

3. PRKSG scheme and its characteristics

Based on the above theoretical information, we present a quaternary PRKSG scheme based on a IV-set of quaternary bent-sequences. The following primitive irreducible polynomials are used for constructing LFSR (in accordance with the results of [9], in order to maximize the period of the generated pseudorandom key sequences, the degrees of primitive irreducible polynomials are chosen as coprime)

$$\begin{aligned}
\text{LFSR}_1 : f_1(x) &= x^{23} + x^{17} + x^{11} + x^5 + 1; \\
\text{LFSR}_2 : f_2(x) &= x^{29} + x^{20} + x^{16} + x^{11} + x^8 + x^4 + x^3 + x^2 + 1; \\
\text{LFSR}_3 : f_3(x) &= x^{31} + x^{21} + x^{12} + x^3 + x^2 + x + 1; \\
\text{LFSR}_4 : f_4(x) &= x^{19} + x^9 + x^8 + x^5 + 1; \\
\text{LFSR}_5 : f_5(x) &= x^{17} + x^{16} + x^3 + x + 1; \\
\text{LFSR}_6 : f_6(x) &= x^{37} + x^{28} + x^{18} + x^9 + 1,
\end{aligned} \tag{11}$$

and also, the IV-set of quaternary bent-sequences is chosen in accordance with condition (8)

$$\begin{aligned}
H_1 &= \{0331132333230113\}; \\
H_2 &= \{1002203000301220\}; \\
H_3 &= \{2113310111012331\}; \\
H_4 &= \{3220021222123002\}.
\end{aligned} \tag{12}$$

In Fig. 3, we present the scheme of the developed PRKSG based on the IV-set of quaternary bent-sequences (12) and LFSR constructed in accordance with primitive irreducible polynomials (11).

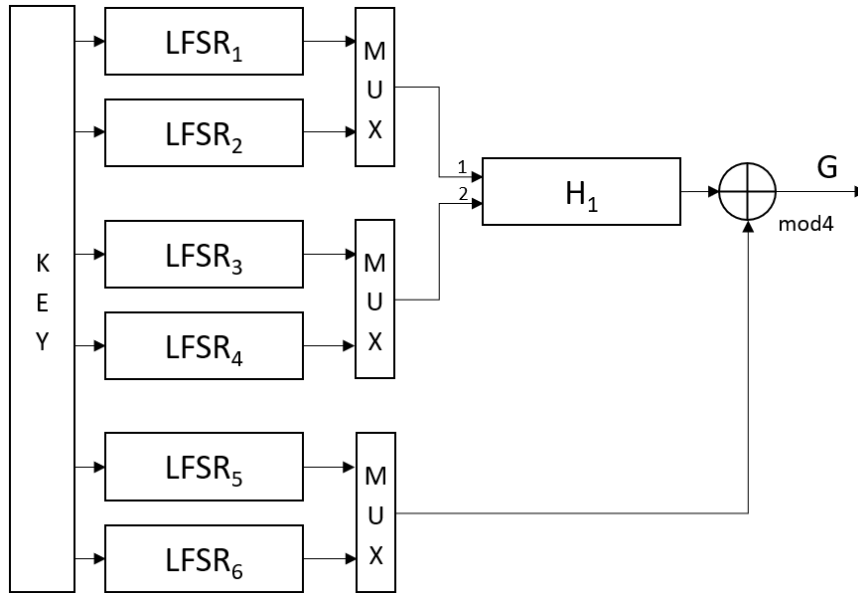


Fig. 3: Scheme of the proposed PRKSG based on IV-sets of bent-sequences (12)

Note that in view of the formation of a IV-set of quaternary bent-sequences in accordance with rule (8), in contrast to the PRKSG schemes presented in [9] and [14], there is no need to store the entire IV-set, as well as implement a selection block, which can be replaced by a summation device, which greatly simplifies the technical implementation of PRKSG and saves required for its implementation memory cells.

Thus, the values of quaternary bent-functions H_2, H_3, H_4 can be formed by adding a constant 1, 2 or 3 to the original bent-function H_1 , respectively. At the same time, the simultaneous operation of LFSR_5 and LFSR_6 allows choosing these values with equal probabilities, which leads to ensuring the equiprobable choice of one of the quaternary bent-sequences from the IV-set at each iteration of the PRKSG operation.

Let us note the features of the operation of the proposed PRKSG scheme based on IV-sets of quaternary bent-sequences. Binary registers LFSR_1 and LFSR_2 generate pseudorandom sequences, which are subsequently multiplexed into quaternary sequences that determine the first input variable of the bent-function (equivalent to corresponding bent-sequence). Similarly, the binary registers

LFSR₂ and LFSR₃ define the second input variable of the quaternary bent-function (equivalent to corresponding bent-sequence).

The registers LFSR₅ and LFSR₆ after the multiplexer generate one of the values from the set {0,1,2,3}, which determines the choice (with help of summation) of quaternary bent-sequence from the IV-set. Next, the corresponding value of the quaternary bent-function (equivalent to corresponding bent-sequence) is calculated, which is the output value of the PRKSG at a current cycle of operation.

The number of protection levels of the developed PRKSG is determined both by the number of possible initial states of the LFSR and by the number of bent-sequences in the complete class. In our case, when using primitive irreducible polynomials (11), as well as the complete class of quaternary bent-sequences [16], the number of protection levels of the developed PRKSG is defined as

$$Y = (2^{23} - 1)(2^{29} - 1)(2^{31} - 1)(2^{19} - 1)(2^{17} - 1)(2^{37} - 1) \cdot 200704 \approx 2^{173,6}, \quad (13)$$

which exceeds the number of protection levels of the AES-128 block symmetric cipher.

At the same time, we note that in order to obtain a quaternary PRKSG based on a combination of two binary PRKSG units based on dual sets of bent-sequences [9], we would have to use 10 LFSR items, which is 40% more than in the proposed scheme. Thus, the use of IV-sets of quaternary bent-sequences makes it possible to simplify the algorithmic implementation of PRKSG in the applications where quaternary pseudorandom key sequences are required.

We note an important property of the proposed PRKSG scheme: it is easily scalable and makes it possible, if necessary, to easily increase the number of protection levels by using quaternary IV-sets of bent-sequences of greater length.

For example, consider the possibility of using quaternary bent-sequences of length $N = 64$. We take as a basis one of the quaternary bent-sequences of given length $N = 64$

$$H_1 = [3333301203213131000001231032020233333012032131312222230132102020], \quad (14)$$

on the basis of which, considering the construction (8), we obtain the IV-set

$$\begin{aligned} H_1 &= \{3333301203213131000001231032020233333012032131312222230132102020\}; \\ H_2 &= \{0000012310320202111112302103131300000123103202023333301203213131\}; \\ H_3 &= \{1111123021031313222223013210202011111230210313130000012310320202\}; \\ H_4 &= \{2222230132102020333330120321313122222301321020201111123021031313\}. \end{aligned} \quad (15)$$

In view of the fact that each of the bent-sequences (15) is the equivalent of the corresponding bent-function of three variables, to construct the corresponding PRKSG we need another two LFSR to ensure the presence of an input signal for each of the variable of equivalent bent-function, as well as two LFSR for ensuring the selection of the bent-sequence from the IV-set.

Thus, in addition to the primitive irreducible polynomials (11), we choose two additional primitive irreducible polynomials

$$\begin{aligned} \text{LFSR}_7 : f_7(x) &= x^{41} + x^3 + 1; \\ \text{LFSR}_8 : f_8(x) &= x^{43} + x^6 + x^4 + x^3 + 1. \end{aligned} \quad (16)$$

On Fig. 4 we present a PRKSG scheme operating using eight LFSR and a IV-set of quaternary bent-sequences (15) of length $N = 64$.

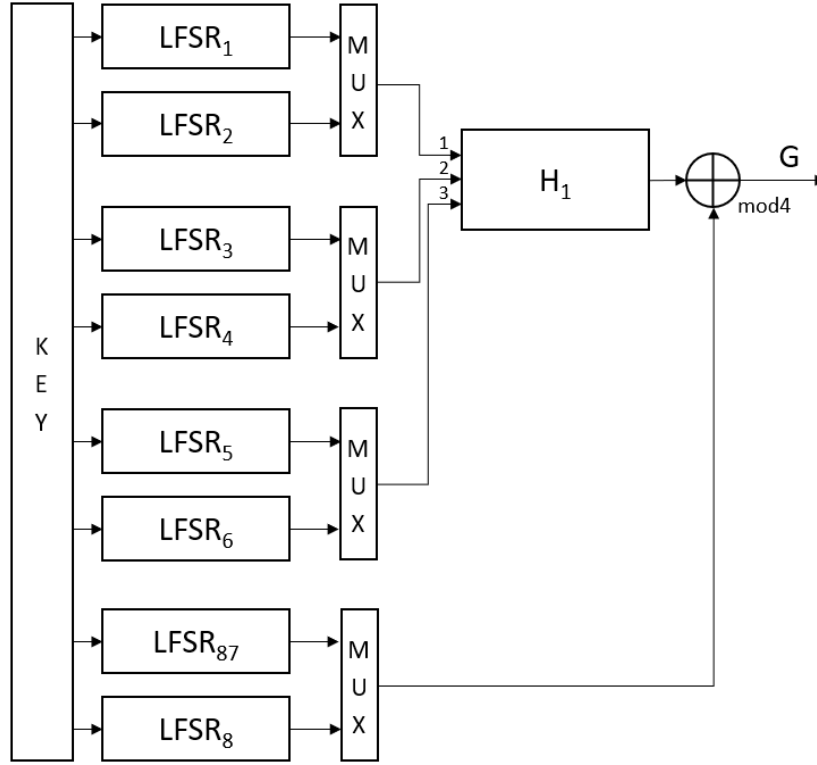


Fig. 4: Scheme of the proposed PRKSG based on IV-sets of bent-sequences (15)

The number of protection levels of the developed PRKSG is determined by the number of possible initial states of the eight LFSR included in it. In the case of using primitive irreducible polynomials (11) and (16), we obtain the following value

$$Y = (2^{23} - 1)(2^{29} - 1)(2^{31} - 1)(2^{19} - 1)(2^{17} - 1)(2^{37} - 1) \times (2^{41} - 1)(2^{43} - 1)(2^{47} - 1) \approx 2^{287}. \quad (17)$$

The obtained value of the number of protection levels exceeds the value of the number of protection levels of the AES-256 cryptographic algorithm.

Note that, to date, there are practically no methods for synthesizing complete classes of quaternary bent-sequences of length $N > 16$ represented in the literature, which leads to the decrease in the total value of protection levels number for developed PRKSG shown in Fig. 4. This circumstance naturally poses the practical problem of developing such a methods in order to increase the number of protection levels for PRKSG based on IV-sets of quaternary bent-sequences.

In Table 2, we present the results of the NIST stochastic tests of the developed PRKSG, its structural elements as well as some known analogues.

Table 2

Results of the NIST stochastic tests

No.	Test	LFSR based on (1)		LFSR based on (9)		PRKSG [10]		Developed PRKSG (Fig. 3)		Developed PRKSG (Fig. 4)	
		P-value	Pass rate	P-value	Pass rate	P-value	Pass rate	P-value	Pass rate	P-value	Pass rate
1	Monobit test	0.51	✓	0.86	✓	0.30	✓	0.09	✓	0.18	✓
2	Frequency within block test	0.74	✓	0.73	✓	0.99	✓	0.47	✓	0.06	✓
3	Runs test	0.72	✓	0.85	✓	0.11	✓	0.40	✓	0.72	✓
4	Longest run ones in a block test	0.14	✓	0.80	✓	0.04	✓	0.75	✓	0.19	✓

5	Binary matrix rank test	0.76	✓	0	✗	0.08	✓	0.87	✓	0.78	✓
6	DFT test	0.84	✓	0	✗	0.79	✓	0.34	✓	0.68	✓
7	Non overlapping template matching test	0.99	✓	0.97	✓	1	✓	1	✓	1	✓
8	Overlapping template matching test	0.01	✓	0.98	✓	0.26	✓	0.89	✓	0.08	✓
9	Maurers universal test	0.8	✓	0.01	✓	0.16	✓	0.40	✓	0.23	✓
10	Linear complexity test	0.16	✓	0	✗	0.83	✓	0.24	✓	0.93	✓
11	Serial test	0.9	✓	0.99	✓	0.18	✓	0.04	✓	0.81	✓
12	Approximate entropy test	0.9	✓	0.99	✓	0.18	✓	0.04	✓	0.81	✓
13	Cumulative sums test	0.49	✓	0.99	✓	0.18	✓	0.06	✓	0.16	✓
14	Random excursion test	0.33	✓	0.04	✓	0.16	✓	0.17	✓	0.04	✓
15	Random excursion variant test	0.42	✓	0	✗	0.29	✓	0.07	✓	0.03	✓

The analysis of the data presented in Table 2 leads to the conclusion that the proposed PRKSG based on IV-sets of quaternary bent-sequences of length $N = 16$, as well as quaternary bent-sequences of length $N = 64$ has a high level of stochastic quality and fully complies with all the NIST stochastic tests.

The NIST test results presented in Table 2, as well as the obvious simplicity of the technical implementation of the proposed PRKSG structure, allow us to recommend to use both developed PRKSG variants (Fig. 3 and Fig. 4) for practical use.

4. Conclusion

We note the main results of the research:

1. It has been established that the construction of quaternary PRKSG is possible on the basis of binary LFSR, as well as IV-sets of quaternary bent-sequences. At the same time, the configuration of the IV-set was found, which provides the best stochastic properties of the gamma generated by the PRKSG.

2. The scheme of the PRKSG is proposed which is based on binary LFSR and IV-set of quaternary bent-sequences. The proposed PRKSG provides a high level of stochastic properties of the generated sequences and the number of protection levels that can be easily scaled. For the case of use of the complete class of quaternary bent-sequences of length $N = 16$, the number of protection levels reaches the value $Y \approx 2^{209,65}$, while in the case of using bent-sequences of length $N = 64$, the number of protection levels reaches the value $Y \approx 2^{287}$, which exceeds the number of protection levels of the AES-256 cryptographic algorithm. Developed PRKSG also requires 40% smaller number of LFSR for generation of quaternary pseudorandom key sequences compared to combining of two PRKSG, based on dual sets of bent-sequences.

3. The presented PRKSG is of interest from the point of view of practical application where quaternary pseudorandom key sequences are needed, for example stream encryption algorithms operating based on many-valued logic principles, communication systems based on the FHSS technology, steganographic algorithms with the possibility of pseudorandom selection of a steganographic path.

Considering the above material, an important direction for further research is the development of methods for the synthesis of complete classes of quaternary bent-sequences, which will expand the variety of possible structures of the developed PRKSG, as well as increase the number of its protection levels value.

Another possible direction for further development of the proposed PRKSG is its modification in order to use larger values of base q .

5. References

1. G. Megala, S. Prabu, P. Swarnalatha, R. Venkatesan, Analysis on Cryptographic Design Techniques of Stream Ciphers and Attacks, in: P. Swarnalatha, S. Prabu (Ed.), *Blockchain Technologies for Sustainable Development in Smart Cities*, IGI Global, US, 2022, pp. 19–33. doi: 10.1007/978-3-540-40974-8_6.
2. M. Hasan, J. M. Thakur, P. Podder, Design and implementation of FHSS and DSSS for secure data transmission, *International Journal of signal processing systems* (2015) 144–149. doi: 10.12720/ijsp.4.2.144-149.
3. I. V. Agafonova, Cryptographic properties of nonlinear Boolean functions, in: *Proceeding of the Seminar on discrete harmonic analysis and geometric modeling*, Saint-Petersburg, 2007, pp. 1–24.
4. B. Schneier, *Applied Cryptography: Protocols, Algorithms and Source Code in C*, Wiley, 2015.
5. O. Datcu, C. Macovei, R. Hobincu, Chaos based cryptographic pseudorandom number generator template with dynamic state change, *Applied sciences* (2020) 451. doi: 10.3390/app10020451.
6. E. Goncu, A. Kocdogan, M. E. Yalcin, A high speed true random number generator with cellular automata with random memory, in: *Proceedings of the IEEE international symposium on circuits and systems (ISCAS)*, Florence, 2018, pp. 1730017. doi: 10.1109/iscas.2018.8351127.
7. A statistical test suite for random and pseudorandom number generators for cryptographic applications. Gaithersburg, MD: U.S. Dept. of Commerce, Technology Administration, National Institute of Standards and Technology, 2000.
8. C. E. Shannon, *A Mathematical Theory of Cryptography*, Bell System Technical Memo, USA, 1945.
9. M. I. Mazurkov, A. V. Sokolov, N. A. Barabanov, The key sequences generator based on bent functions dual couples, *Proceedings of Odessa Polytechnic University* (2013) 150–156.
10. A. Sokolov, N. Kazakova, L. Kuzmenko, M. Mahomedova, Prerequisites for developing a methodology for estimating and increasing cryptographic strength based on many-valued logic functions, in: *Proceedings of Selected Papers of the Workshop on Cybersecurity Providing in Information and Telecommunication Systems (CPITS 2021)*, Kyiv, Ukraine, 2021, pp. 107–116.
11. M. I. Mazurkov, *Broadband radio systems*, Odessa: Science and Technology, 2010.
12. M. A. Ivanov, I. V. Chugunkov, *Theory, application and evaluation of the quality of generators of pseudorandom sequences*, Moscow, KUDITS-OBRAZ, 2003.
13. FIPS 197. Advanced encryption standard, 2001. URL: <http://csrc.nist.gov/publications/>
14. A. V. Sokolov, O. N. Zhdanov, N. A. Barabanov, Pseudorandom key sequence generator based on triple sets of bent-functions, *Problems of Physics, Mathematics and Technics* (2016) 85–91.
15. K. Schmidt, Quaternary Constant-Amplitude Codes for Multicode CDMA, *IEEE transactions on information theory* (2009) 1824–1832. doi: 10.1109/tit.2009.2013041.
16. A. V. Sokolov, Properties of the full class of quaternary bent-functions of two variables, *Journal of discrete mathematical sciences and cryptography* (2021): 1–14. doi: 10.1080/09720529.2021.1884380.
17. M. I. Mazurkov, Ye. A. Konopaka, The families of linear recurrent sequences based on full sets of Galois' isomorphic fields, *Radioelectronics and Communications Systems* (2005) 42–47.

Decision-making Algorithm in Case of Failure of the Electric Motor of a Multi-rotor Unmanned Aerial Vehicle

Vitaliy Larin^a, Heorhii Rozorinov^b, Oleksandr Hres^c, Volodymyr Rusyn^c, Sergey Subbotin^d, Nina Chichikalo^b and Katerina Larina^b

^a National Aviation University, L. Huzara ave, 1, Kyiv, 03058, Ukraine

^b National Technical University of Ukraine "Igor Sikorsky Kyiv Politechnic Institute", Peremohy ave, 37, Kyiv, 03056, Ukraine

^c Yuriy Fedkovych Chernivtsi National University, Kotsybinsky str., 2, Chernivtsi, 58012, Ukraine

^d National University "Zaporizhzhia Politechnic", Zhukovsky str., 64, Zaporizhzhia, 69063, Ukraine

Abstract

The safety of unmanned aerial vehicle (UAV) flights depends on many factors, such as the absence of failures or malfunctions of aviation equipment, the absence of exposure to adverse environmental phenomena, and the absence of errors by the aircraft crew and engineering personnel. In uncontrolled airspace by the internal affairs authorities, when the flight is carried out at altitudes below 500 feet above ground level (AGL), this task is even more complicated, since at the moment there are no monitoring services and procedures for monitoring VTOL-UAV (Vertical Take-Off and Landing unmanned aerial vehicle) performing operations at the specified altitude range. In this case, the only link of control is the remote pilot, who directly monitors his unmanned aerial vehicle (UAV) when flying in Visual Line-of-Sight (VLOS) mode. Some components of an unmanned vehicle are difficult to control from the point of view of preventing the risk of the likelihood of an incident, which is difficult to prevent due to its high potential danger and the speed of its occurrence.

In this paper, we propose an algorithm that is simple for software implementation and does not require mathematical calculations, the implementation of which requires the presence of devices for measuring the speed of rotation of engines, which are proposed to use Hall sensors installed on each engine of a multi-rotor VTOL-UAV.

To prevent the program from crashing when polling sensors, a double redundancy of sensor readings is provided. Also, in case of confirmation of an engine failure, a module has been introduced into the algorithm that provides for a preliminary shutdown of an engine that is symmetrical to the one whose failure is confirmed. After turning off the remaining engines, the parachute compartment is activated for an accurate landing at low speed.

Keywords 1

Algorithm, motor failure, aviation incident, VTOL-UAV

1. Introduction

The safety of unmanned aerial vehicle (UAV) flights depends on many factors, such as the absence of failures or malfunctions of aviation equipment, the absence of exposure to adverse environmental phenomena, and the absence of errors by the aircraft crew and engineering personnel. These factors are classified in the International Civil Aviation Organization (ICAO) risk matrix [1].

In controlled airspace, any manned aerial vehicle is under constant control of the crew and air traffic control units. The controllability of an unmanned aerial vehicle is a much more difficult task, since its crew- the remote pilot's team - are at a considerable distance and can assess the

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: vjlarin@gmail.com (A. 1); grozoryn@gmail.com (A. 2); o.hres@chnu.edu.ua (A. 3); rusyn_v@ukr.net (A. 4); subbotin@zntu.edu.ua (A. 5); Hni312123@iit.kpi.ua (A. 6); l.katerina@gmail.com (A. 7)

ORCID: 0000-0002-5042-2426 (A. 1); 0000-0002-6095-7539 (A. 2); 0000-0002-8465-193X (A. 3); 0000-0001-6219-1031 (A. 4); 0000-0001-5814-8268 (A. 5); 0000-0002-3619-5992 (A. 6); 0000-0002-2315-0891 (A. 7)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org) Proceedings

controllability of their aircraft based on telemetry data that comes along the downstream branch of the C2 channel.

In uncontrolled airspace by the internal affairs authorities, when the flight is carried out at altitudes below 500 feet above ground level (AGL), this task is even more complicated, since at the moment there are no monitoring services and procedures for monitoring VTOL-UAV (Vertical Take-Off and Landing unmanned aerial vehicle) performing operations at the specified altitude range. This altitude range is currently of particular interest to users, since it allows one to implement new or update existing high-tech production and economic processes. In this case, the only link of control is the remote pilot, who directly monitors his unmanned aerial vehicle (UAV) when flying in Visual Line-of-Sight (VLOS) mode, or based on telemetry data, with various Radio Line-of-Sight (RLOS) mode options. Some components of an unmanned vehicle are difficult to control from the point of view of preventing the risk of the likelihood of an incident, which is difficult to prevent due to its high potential danger and the speed of its occurrence. These components include the UAV electric motor. Failure of the UAV motor (sudden or partial) is referred to the group of aviation equipment failures.

2. Related works

In manned aviation, special attention is paid to assessing the risks of aircraft flight deviation from a given flight trajectory. To assess such risks, various methods are often used based on calculating the probability density function, while many publications on the topic of aviation risk assessment suggest a complex use of various options for calculating the probability density function. In particular, in [2], they propose to use a complex model called by the authors Triple Univariate Generalized Error Distribution (TUGED), which is supported by the Maximum Likelihood Method. The authors, by modeling, confirm the effectiveness of the proposed model by checking its compliance with the Chi-square, Bayesian and Akaike criteria.

In [3], it is indicated that the failure rate for various reasons for unmanned aircraft is 1/103 per flight hour, which is two orders of magnitude higher than the value of the same indicator for manned civil (commercial) aviation. Considering the above, the task of developing a decision-making algorithm in the event of a failure of the VTOL-UAV electric motor during flight is an urgent task to reduce the risk of an incident and, accordingly, increase the safety of VTOL-UAV flight.

According to [4], failures due to engine failure account for 411 cases per 1000 failures in general due to failures of various UAV components. It should be noted that this statistics took into account only the failures of internal combustion engines, which are now used on VTOL-UAV, performing long flights, more than one hour. In addition, the data takes into account the type of VTOL-UAV of the traditional aircraft type. In the same source, at the level of the components of the propulsion system, various failures are detailed - failures of the engine itself (63 cases), failures of the ignition system (97 cases); fuel system failures (120 cases), temperature control system failures (48 cases); failures of the management apparatus (83 cases).

As for the VTOL-UAV with an electric motor system, such statistics are either absent or not found. However, according to the experience of the authors, during experimental studies to determine the discharge rate of a LiON battery [5], brushless motors were used as a load, which are widely used as VTOL-UAV engines for vertical take-off and landing of VTOL and small HTOL (abbreviated as VTOL - Vertical Take-Off Landing [6]), weighing less than 25 kg. During the experiments, one of the engines suddenly failed for an unknown reason and was replaced with another. This case occurred under laboratory tests. It is obvious that under conditions of environmental changes during the flight, the probability of motor failure will increase.

In [7], a failure risk assessment was carried out for two VTOL-UAV designs: one VTOL-UAV design for vertical take-off and landing of VTOL (VTOL-UAV of this type is shown in Fig. 1) and another VTOL-UAV design for horizontal take-off and landing of HTOL (VTOL -UAV of this type is shown in Fig. 2, 3).



Figure 1: Multi-engine helicopter for cargo transportation PKM-14 "Saturn" [8]



Figure 2: Twin-engine unmanned aerial vehicle M-7V5 "Sky Patrol" [9]



Figure 3: Unmanned complex of hybrid type Zala Aero Zala Vtol

Zala Aero Zala Vtol is a hybrid unmanned complex, which, according to the developers, reduces the role of the human factor, the number of used and maintained equipment in the flight task, and fully automate flight processes. During operation, the equipment relies on the computing power of the on-board computer ZX1, which uses artificial intelligence technology, which allows to process data in Full HD and transmit HD video and photos via encrypted communication channels to NSU, ensuring effective monitoring before landing. According to the developers, the versatility of the Zala Vtol design makes it fully compatible with all existing Zala target loads, and also allows the installation of additional surveying equipment.

According to preliminary calculations, the crash probability density (CPD) was calculated, which made it possible to determine the frontal impact area, which, in turn, was used to calculate the number of incidents (Expected Level of Safety) over a normalized time range using the Weibel and Hansman model of collision with the ground [10] (an incident in this model is understood as the number of impacts of an emergency board with a person, which leads to his injury or even death). Thus, for the push-type VTOL-UAV of horizontal take-off and landing with one diesel engine, $ELS_{HTOL}=1.47 \cdot 10^{-9}$ incidents/hour was obtained, and for the VTOL of vertical take-off and landing (quadrotor-in-quadrotor (QIQ) octocopter model) $ELS_{VTOL}=4.96 \cdot 10^{-8}$ incidents/hour. Whence it follows that the calculated incident rate for VTOL is almost an order of magnitude higher than for HTOL. The authors of [5] propose the calculation of the safety corridor when planning the flight trajectory, the width of which will depend on the obtained value of the accident probability density. For example, when a VTOL-UAV QIQ model is flying at an altitude of 122 m (400 ft) at a speed of 10 m/s, the width of such a corridor will be 38 m. In this case, the flight trajectory should not run through populated areas or crowded areas.

However, when used within dense urban areas with a high population density, it will be incredibly difficult to plan such a safe route, which will significantly complicate the implementation of the U-space concept.

The concept of constructing a trajectory based on the assessment of multiple risks in the U-space was also proposed in [11]. The authors proposed an extended structure of the UTM Risk Assessment Framework (URAF), which implies its use just for predicting the occurrence of the risk of an incident within the boundaries of a densely populated urban area from a multi-rotor VTOL-UAV. The onboard system, the hardware implementation of which is called the core Flight System, built on the Bayesian trust model, assumes, among other types of assessing the performance of various VTOL-UAV components, including monitoring the VTOL-UAV engine parameters. At the same time, the proposed flight risk assessment system, according to the authors, should be within the competence of

the UAS Traffic Management (UTM) service, which, in fact, has not yet been created. Implementation of the proposed system would be an almost ideal solution to ensure security, but the declared implementation of the system in real time on board a specific VTOL-UAV, in our opinion, seems to be a difficult task, most likely not so much because of the large number of auxiliary hardware components, but also huge volumes of data that will only come from external sources. For example, data on the current population density in the flight area, or data on the strength of the protective coating (strength of the roof material). In addition, the authors did not take into account the very low degree of controllability of the multi-rotor VTOL-UAV in the event of a sudden failure of one of the engines. If such a failure occurs, the remote pilot will not be able to perform a touchdown operation at a new system-defined landing point or return to a departure point.

3. Problem statement

For unmanned aerial vehicles with vertical take-off and landing, there is no algorithm for ensuring the required safety during flight in the event of failure of one (or more) engines.

It is this type of VTOL-UAV that will be used for flights at very low level altitudes (VLL) and within dense urban areas (U-space). Since VTOL-UAV of this type is the most vulnerable to loss of control due to various failures, the development of decision-making algorithms in the event of failure of critical VTOL-UAV components, which include engines, is relevant, since it will help prevent an aviation incident.

4. Results

4.1. Expected risk of an incident due to motor failure

Consider how high the risk of an incident is in the event of a brushless motor failure. For technical devices, the most convenient reliability characteristics are failure rate and mean time between failures.

Failure rate - $\lambda(t)$ is the conditional density of the uptime distribution for time t , provided that no failure of the device has occurred before time t .

$$\lambda(t) = \frac{a(t)}{P(t)}, \quad (1)$$

where $a(t)$ is failure rate, $P(t)$ is probability of failure-free operation. For practical purposes, the following statistical expression for the failure rate can be used

$$\lambda(t) = \frac{n(\Delta t)}{N_{av} \Delta t}, \quad (2)$$

where $n(\Delta t)$ is the number of devices that failed over time, $N_{av} = \frac{N_i + N_{i+1}}{2}$ is the average number of devices that work without failure over time Δt .

From [12] we know the boundary generalized values of the failure rate for small electrical machines, which include brushless motors - they are in the range from 0.01 to $8 \cdot 10^{-4}$ 1 / h.

The mean time between failures is defined as the mathematical expectation of the device operation until the first failure

$$MTBF = \frac{\sum_{i=1}^N t_i}{T}, \quad (3)$$

where t_i is the uptime of the i -th device (until the first failure), h; T is the total operating time of the device.

According to the same source, $MTBF$ for small electrical machines is in the range of 1.000 to 20.000 hours. $MTBF$ performance of high-quality brushless motors is the best among this group.

However, there are two factors, the influence of which, according to [12], significantly reduces the specified reliability characteristics. The first factor is the increased rotational speed expressed in rpm.

Thus, at a nominal rotation speed of 2500 rpm, the guaranteed uptime is 3000 hours, and when the speed rises to 9000 rpm (more than three times), the uptime is reduced to a value from 200 to 600 hours. That is, when the rotational speed is increased by three times, the failure-free operation is reduced by a factor of five or more! Therefore, taking this factor into account is mandatory when assessing the risk of an incident in the event of an engine failure.

The second factor that can significantly increase the likelihood of failure is the ambient temperature. Higher ambient temperatures also reduce guaranteed uptime. The dependences of the decrease in the uptime with a rise in temperature coincide significantly with the dependences under the influence of the factor of increased rotation speed. Thus, when assessing the risk of an incident, the temperature factor must also be taken into account. Moreover, the elimination of the influence of the first negative factor can be achieved by maintaining the nominal speed mode during manual remote piloting, or by introducing algorithmic restrictions when piloting in the autopilot mode. But the influence of the second factor, which is external to the unmanned aviation system, does not depend in any way on the remote pilot's command and thus must be taken into account when creating various algorithms for on-board systems that make it possible to increase the safety of VTOL-UAV flight.

4.2. Additional hardware tool to realization of developed algorithm

To implement the algorithm, it is necessary to have a device for monitoring the speed of rotation of the propeller, which is installed on the electric motor. As a similar device, one can use a speed sensor that operates on the basis of the magnetoelectric Hall effect. Some manufacturers of higher quality brushless motors already integrate Hall sensors into the motor design, which also serve as a positioning support. In this case, the task of providing motor control is simplified. But even if there are no sensors in the design of a brushless electric motor, it is possible to perform local design modifications by connecting an external sensor to each unit.



Figure 4: Hall sensor SS41F from SEC Electronic Inc. in 3-lead SIP package [13]

Such a sensor will record every change in the level of the magnetomotive force, which is induced due to the appearance of an electric current in the stator coils. Thus, information about the speed of each propeller from the Hall sensors will be transmitted to the VTOL-UAV microprocessor-based flight control system. To implement the algorithm, the number of sensors will be required, depending on the number of electric motors on the VTOL-UAV, since it is necessary to control the rotation speed of each electric motor using its own sensor. The main task of the developed algorithm is to monitor the performance of electric motors even in case of failure, which will be identified in the event of either no rotation or insufficient rotation speed of one of the motors. Since such a failure will immediately lead to a loss of control and the emergence of a conflict-hazardous situation, it is necessary to automatically work out a decision on the immediate termination of the VTOL-UAV mission and its landing.

The next device required for the implementation of the algorithm is a software module - a microprocessor or microcontroller. On board VTOL-UAV such a module is mandatory, and very often not alone - one microprocessor provides processing of navigation data from different navigation sensors, the second microprocessor is the core of the flight control system - autopilot. Since modern embedded microprocessors have a high degree of integration and contain several cores on their chip, both of these functions can be assigned to a single, high-performance microprocessor. However, many autopilot designers do separate navigation and control functions. Moreover, the latest trends in the development of on-board systems provide for equipping VTOL-UAV with microprocessors with a battery management system (BMS) [14, 15], which makes it possible to increase the level of VTOL-UAV flight safety. Since the microprocessor of the BMS system interacts with the electric motor, its software will require minor modifications by adding the code of the developed algorithm, which, due to the introduction of hardware controls, will be quite simple and compact in terms of the program code.

4.3. Design of the algorithm

With regard to the choice of the necessary hardware, the block diagram of the decision-making algorithm will have the following form, shown in Fig. 5. At the first step of the software implementation of the algorithm, it is necessary to set the failure counter - x for the i -th engine. The counter of events (motor failures) can take the values of $x[i]=\{0, 1, 2\}$. The event counter is reset to zero during the initialization of the onboard control system before starting the flight. As a failure criterion, we will consider an event whereby the current value of the rotation speed of the i -th engine $V_{rot}[i]$ measured with the i -th sensor (*sensor* $[i]$) will be less than the required value of the speed $V_{reqmnt}[i]$ or even equal to zero, necessary to ensure the controlled flight of VTOL-UAV. As long as the flight is in progress - the algorithm provides for a cyclic poll of the rotation sensors and compares the current measured speed value with the set value generated by the control system. In the case of the first fixation of a failure of one of n electric motors, the counter of the i -th motor failure is incremented for the first time and takes the value $x[i] = 1$. After the initial setting of the failure counter, a re-confirmation check is performed, and if this check gives a negative result, the transition to the next step of the algorithm is performed. In this case, such an electric motor is considered to be conventionally failed, and information is sent to the control system about the repeated sending of a control action to this engine in order to establish the required value of its rotation speed. The control action is usually a PWM signal, the required pulse width of which is set by the processor of the VTOL-UAV control system.

If the condition $V_{rot}[i] < V_{reqmnt}[i]$, is repeatedly satisfied, then the failure counter takes the value $x[i]=2$ and the transition to the algorithm branch responsible for the emergency termination of the flight occurs, since when one electric motor is inoperative, the probability of the origin of an aviation incident increases.

The emergency command block must contain two sequential procedures. In the event of a failure of one of the engines, the multi-rotor type VTOL-UAV immediately loses stability and begins to rotate around a horizontal or vertical axis. If at this moment in time all the engines are stopped at once and then the parachute compartment is opened for such an unstable VTOL-UAV, the parachute will get entangled around the hull and the VTOL-UAV will still make an uncontrolled fall, thus increasing the risk of an incident many times over.

In case of failure of one of the engines of the propeller-driven group, to stabilize the multi-rotor aircraft, it is necessary to immediately shut down the symmetrical engine, which is located diagonally opposite the suddenly failed engine according to the diagram shown in Fig. 6 where symmetric engines for a six-screw VTOL-UAV (hexacopter) are connected by a blue, red or green dashed line, respectively. After such a shutdown, the motors remaining in working condition will stabilize and prevent an uncontrolled fall for some time.

Thus, in the branch of the emergency subroutine, the emergency engine is first assigned a zero index - $j = 0$, then $x[0]$. Following that, the symmetrical (diagonal) motor is turned off. The index of such an engine for a quadcopter will be $x[2]$, for a hexacopter $x[3]$, for an octacopter $x[4]$. To unify the algorithm in the block diagram of the algorithm, the symmetric motor index is denoted as *symmetr*.

After this operation, a short time delay of no more than 3 s is required to stabilize the aircraft. The block-diagram of the algorithm shows this delay as the *_delay* statement.

Further execution of the program involves sending a command to the flight control system to simultaneously turn off all other electric motors, which in a formalized form can be represented as a command $V_{rot}[i] = 0$, executed in a cycle. Next, the sending of the current coordinates is initiated, at which the flight is forcibly interrupted, via the downstream C2 subchannel to the remote pilot to notify him of the location of the VTOL-UAV, and a command is sent to the control system that initiates the opening of the parachute compartment and the release of the parachute. Then VTOL-UAV makes landing.

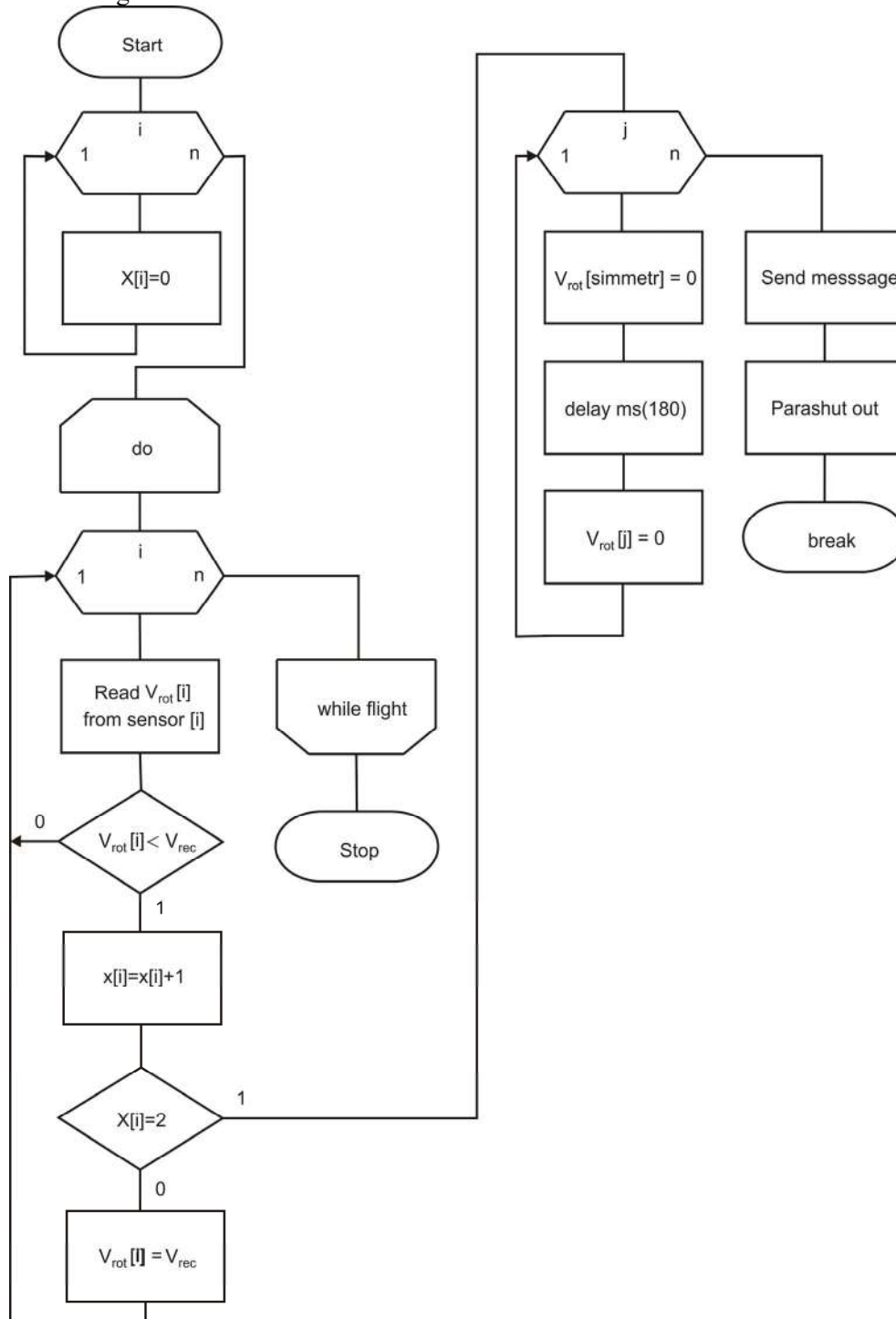


Figure 5: Block diagram of the decision-making algorithm in case of failure of the VTOL-UAV electric motor

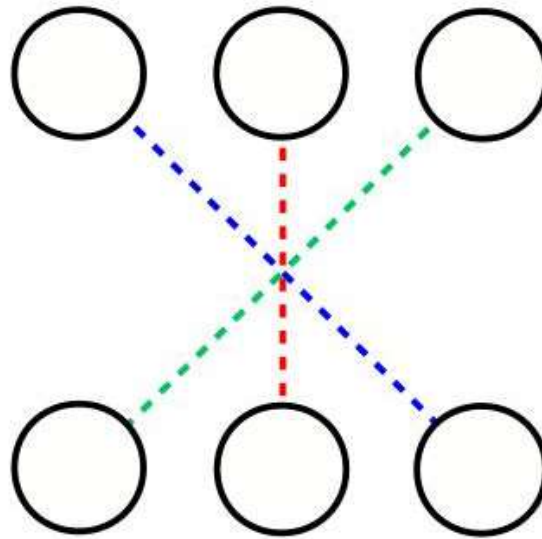


Figure 6: Diagram showing pairs of symmetrical motors for a hexacopter

5. Conclusions

Based on the analysis of sources, two negative factors have been identified that significantly reduce the standard indicators of the uptime of brushless motors and thus increase the risk of an incident: increased values of the speed of rotation of the rotor of the electric motor, which will be present in case of need for high-speed flight mode of VTOL-UAV and increased ambient temperature.

Taking into account the limited time for making a decision and preventing the risk of an incident, an algorithm that is simple for software implementation has been developed that does not require mathematical calculations, for the implementation of which it is necessary to have devices for measuring the speed of rotation of engines. For this purpose it is proposed to use Hall sensors installed on each engine of a multi-rotor VTOL-UAV.

To prevent the program from failure when polling sensors, double redundancy of sensor readings is provided. Also, in case of confirmation of engine failure, a module has been introduced into the algorithm providing for preliminary shutdown of the engine symmetrical to the one who has confirmed the failure. After turning off the rest of the engines, the parachute compartment is activated for an accurate landing at low speed.

6. References

- [1] ICAO Annex 19 'Safety Management', first edition. 2013. 44 p.
- [2] I. Ostroumov, K. Marais, N. Kuzmenko, N. Fala, Triple Probability Density Distribution Model in the Task Of Aviation Risk Assesment. AVIATION 24 (2020) 57-65. doi:/10.3846/aviation.2020.12544.
- [3] E. Petritoli, F. Leccese, L. Ciani, Reliability and Maintenance Analysis of Unmanned Aerial Vehicles. Sensors 18 (2018) 1-16. doi:10.3390/s18093171.
- [4] Austin, R. Unmanned Aircraft Systems - Wiley: Hoboken, NJ, USA, 2010, 323 p.
- [5] Larin V.J., Shcherban, A.P., Ryzhykh V.M., V.P. Maslov, N.V. Kachur Use of infrared thermography method to develop discharging rules for lithium-polymer batteries. Semiconductor Physics, Quantum Electronics and Optoelectronics 22 (2019) 252-256. doi:10.15407/spqeo22.02.252
- [6] V.G. Ambrosia A.C. Watts and E.A. Hinkley. Unmanned Aircraft Systems in Remote Sensing and Scientific Research: Classification and Considerations of Use. Remote Sensing, (2012) 1671-1692.

- [7] Chin E. Lin, Pei-Chi Shao. Failure Analysis for an Unmanned Aerial Vehicle Using Safe Path Planning. *Journal of Aerospace Information Systems* 17 (2020) 358-369.
- [8] Multicopter helicopter to the cargo carrying PKM-14 “Saturnia”. URL: - http://uav.nau.edu.ua/pkm-14_ukr.html.
- [9] Twin-engined unmanned aerial vehicle M-7V5 “Sky patrol”. URL: - http://uav.nau.edu.ua/m-7v5_ukr.html.
- [10] Weibel, R. E., Hansman, R. J., Jr., Safety Considerations for Operations of Different Classes of UAVs in the NAS, AIAA 3rd “Unmanned Unlimited” Technical Conference, Workshop and Exhibit, AIAA Paper 2004-6244, 2004.
- [11] Ersin Ancel, Francisco M. Capristan, John V. Foster. In-Time Non-Participant Casualty Risk Assessment to Support Onboard Decision Making for Autonomous Unmanned Aircraft. AIAA AVIATION Forum, 17-21 June 2019, Dallas, Texas, doi:10.2514/6.2019-3053.
- [12] J. F. Gieras, Permanent Magnet Motor Technology: Design and Applications. 3rd Edition, Taylor & Francis CRC Press Group. 2010, 603 p. ISBN: 978-1-4200-6440-7.
- [13] SS41F Datasheet/Electronic Components Datasheet Search — URL: https://www.alldatasheet.com/view.jsp?Searchword=Ss41f&gclid=EAIaIQobChMI4K-vkvaC8gIVB6p3Ch0oWwb7EAMYASAAEgJOQPD_BwE — available 29.06.2021.
- [14] A. Shcherban System for monitoring the state of the battery of an unmanned aerial vehicle. PhD thesis. Institute of Electrodynamics of NAS of Ukraine, Kyiv, 2020.
- [15] A. Shcherban, V. Larin, V. Maslov, N. Kachur, T. Turu. Intelligent System for Temperature Control of Li-Pol Battery [on-line]. *International Journal of Automation, Control and Intelligent Systems* 4 (2018) 24-28.

Formation of a Method for Estimating the Error of Determining the Coordinates of the Source of a Sound Anomaly

Aleksandr Trunov¹, Zhan O. Byelozyorov¹, Serhii I. Maltsev¹, Maksym Skoroid¹

¹ Petro Mohyla Black Sea National University, Mykolaiv, Ukraine, 68 Desantnikov, 10, Mykolaiv 54000, Ukraine.

Abstract

The methods of determining the coordinates of the source of a sound anomaly are considered. Improving the quality of the equipment of the computerized systems of microphones for registration of intensity of noise was done by the method of Monte-Carlo simulation. The computing for different components and parameters of the system and evaluation of error for comparison of the efficiency of four algorithms was done. As a result, consideration set a task conduct simulation of random processes that form a random component of error into computer systems of microphones and movable unmanned aerial or ground vehicles as a source of a dynamic error. To experimental study the propagation of the front of SA wave, the nature of its attenuation, and physical modeling of the random position of the point of the source SA, it was proposed to use a rotary platform on thrust bearings driven by a stepper motor. The model of rotation control of platform in Matlab by using especially a hybrid two-phase stepper motor model is considered. The comparative analysis of estimates of the maximum possible error of algorithms for symmetric and asymmetric echograms was provided. On the bases of analyzing the algorithms for recognizing and determining the momentum of input time of complex echograms was done. The recommendation for conditions and advantages for algorithm application was formulated.

Keywords

Sound Anomaly, coordinates, random simulation, estimating error, static, dynamic

1. Introduction

The strategy of offensive and defensive actions demonstrates the success of the use of unmanned aerial vehicles and anti-tank systems. These results show that the paradigm of success is dominated by the number of personnel and tanks with artillery successfully changing the paradigm of automation and implementation of computer-integrated technologies, which aims to maximize the preservation of personnel [1-3]. However, its further development and practical application require the development of small-sized means of determining the coordinates of the sound anomaly (SA) [4-5] and ensuring the assessment of the maximum possible error [6-9]. This task is becoming especially important and relevant for devices installed on mobile devices such as unmanned aerial vehicles (UAVs) or ground unmanned vehicles (GUVs) or ground unmanned combat vehicles (GUCVs) [1-3]. In addition to these, the possible use of small-scale means of determining the coordinates for further improvement of monitoring systems for civilian use [4-5] requires the development of methods for estimating the maximum possible error and the formation of measures to reduce it.

2. Analysis of recent publications

The work [1] reveals the advantages of using UAVs for compatible video and audio

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL: trunovalexandr@gmail.com (A. Trunov); zhan.byelozyorov@gmail.com (Z. Byelozyorov); msereg92@gmail.com (S. Maltsev); hahido71@gmail.com (M. Skoroid)
ORCID: 0000-0002-8524-7840 (A. Trunov); 0000-0002-3216-8153 (Z. Byelozyorov); 0000-0001-9094-5259 (S. Maltsev); 0000-0003-1641-9124 (M. Skoroid)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

reconnaissance, which complement each other in information and allows you to enter this information in the electronic fire card of the department. However, despite the advantages of such implementation of computer-integrated technologies, the complication of effective use of metrological measurement scheme used in this article is the lack of binding of coordinates of the carrier and device for determining the coordinates of the source SA to one basic coordinate system. In addition, it is unreasonably assumed that the error of the coordinates of the center of gravity of the UAV, as determined by GPS, is zero, and its relative orientation does not contribute to the magnitude on the error. In work [2] the necessity is highlighted and the reality of NBR application is demonstrated. However, during the operation of the GUAV there will also be a problem of determining the coordinates and its spatial orientation. In [3] there is also the problem of estimating the error and its impact on the result of reliable data entry into the electronic card of the fire department. Thus, the development of acoustic means of detecting a shot from a small weapon, their classification and formation of design requirements are presented in [4]. However, these issues of choosing a metrological scheme and determining the error were not raised and resolved. The development of devices for wireless monitoring and recording the coordinates of the shot as one of the examples of SA [5] opens up new opportunities for further improvement of equipment. Thus, [6] forms a method for estimating the error of measuring the angle directly to the target by a distributed system of sound artillery reconnaissance. However, the error of determining the coordinates of SA is ignored. The methodical error of direction-finding target by sound artillery reconnaissance system [7] is not taking into account the motion of computer system of microphones.

Peculiarities of the application of the sound metric complex in the deployment at short distances with simulation as a combat application, which determines the priority tasks of data reliability and accuracy of the coordinates of the source SA [8]. The proposed method of the evaluation value is applied for Polish aviation in the SBAS APV landing procedure [8]. The experience and results of its implementation into Position Dilution of Precision are demonstrated and compared for the two air tests in Dęblin and Chełm [8].

Analysis of the reasons and factors for the formation of error in the processes of sound reconnaissance or other types of SA shows that further expansion of the scope of application with the use of mobile vehicles will only increase the maximum possible error [9]. The article shows preliminary results of EGNOS Safety of Life service performance in Dęblin in comparison to the results obtained in Olsztyn [9]. The main parameters characterizing a navigational system i.e. accuracy, integrity, continuity, and availability were analyzed in detail [9]. The experience can have to consider as a basis to assess the possibility of implementing the EGNOS APV approach and landing procedures in Dęblin and Olsztyn [9].

The useful contribution to the formation of causes and quantification of error is presented in [10] for a multilevel learning approach based on B-splines for the localization of sound sources. In [11] the algorithm and the process of automatic determination of microphone coordinates are presented. The quality of the created hardware, which shows the tendency to miniaturize and integrate single-chip single-board computers, leads to the spread of examples of use in the rehabilitation of the military, working in high-noise areas [12]. The second success is the use in flexible robotic systems [13]. A special role was played by successes in the development of elements of the CS in the robotic systems of navigation units and positioning of production elements of flexible systems [14].

Enhancement of methods and modernization of means of study opportunities to improve the accuracy and reliability especially in-flight due to geometric calibration of the spacecraft imaging complex using known and unknown landmarks are demonstrated in works [15]. The constructing dynamic scenarios in the work of navigation geographic information systems based on the principles of sound location is also improved based on the development of CS elements [15] and requires comprehensive accounting and modeling of all components that form the error of the device as a whole.

Thus, the search for the causes of methodological and random errors in mobile monitoring systems by UAV, GUV and GUCV has emerged as an unsolved problem, due to growing needs for accuracy assessment and application requirements.

Thus, the availability of innovative hardware [5], which has been created recently, will contribute to the formulation and solution of the problem of developing a perfect metrological scheme and method of correctly estimating the maximum possible error.

3. Purpose and objectives of the study

The aim of the work is to improve the quality of the equipment of the CS of microphones of registration of SA, by simulation, accounting for components that form the error of its evaluation and the introduction of algorithms that reduce it.

To achieve this goal, the following tasks were set:

- Set a task and conduct simulation of random processes that form a random component of error;
- Carry out a comparative analysis of estimates of the maximum possible error of algorithms for symmetric and asymmetric echograms;
- To analyze the algorithms for recognizing and determining the time of complex echograms.

4. Statement of the problem of modeling and research of the influence of the algorithm on the value of the random component of the maximum possible error.

The echogram recordings that formed sound simulators were used as a source of data material [5]. The methods of error estimation theory, probability theory, normal distribution, Monte - Carlo, recurrent approximation and simulation in Matlab methods were used to solve the research problems.

4.1 Statement of the problem of simulation modeling of random processes that form a random component of error.

The CS consisting from four microphones was considered, each of which is located at the vertices of an equilateral pyramid with an equilateral triangle at the base [5]. A connected Cartesian coordinate system is chosen. Its starting point O is taken so that it coincides with the point of intersection of the three bisectors of the triangle base AKD. In some cases, it was assumed that the fourth microphone is turned off, then the pyramid degenerates into a triangle. The scheme of arrangement of microphones as point and stationary receivers is presented in fig. 1.

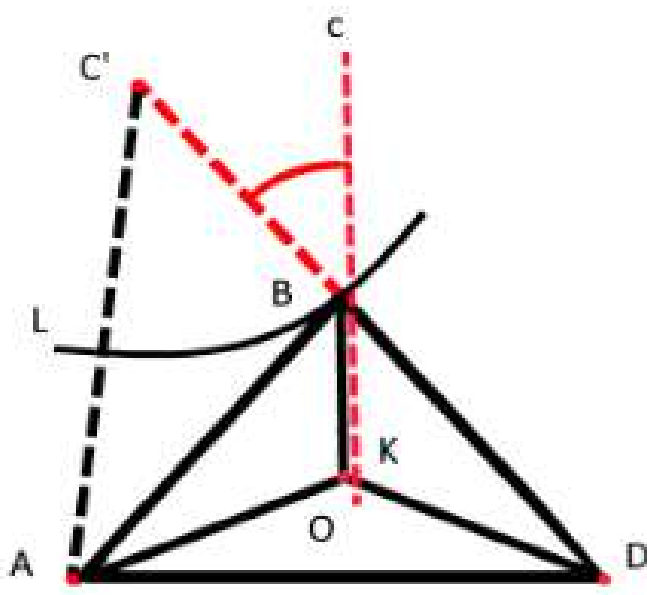


Figure 1: Scheme of four microphones location as point receivers

Next, we assume that SA are point and stationary sources. Let us denote by the letter C the point where the source SA is located. When the source is shifted to the point C', the rays CB and C'B form

an angle $\angle CBC'$, which changes randomly, which is described by the normal law of Gaussian density. Under these conditions, as indicated, the range of change of the angle $\angle CBC'$ will be a limited interval $[-5 / 6\pi, 5 / 6\pi]$. In the simulation, the intermediate value of the angle $\angle CBC'$ was calculated from this range by the expression:

$$\angle CBC' = \alpha - \Delta\alpha + 2\Delta\alpha k_i, \quad (1)$$

where k_i - is a random number evenly distributed in the interval $[0, 1]$.

The angle $\angle ABC'$ is calculated by eq. (1) as a random variable and forms the elements of the sample:

$$\angle ABC' = 180 - \angle CBC' - 30 = 150 - \angle CBC'. \quad (2)$$

The distance from the wave front to the microphone will also be a random variable and will be equal to:

$$AL = \sqrt{r^2 - 2ra \cos \angle ABC' + a^2} - r. \quad (3)$$

In [5-9] it was assumed that the front of the acoustic wave during propagation did not change the amplitude, and hence the intensity of the wave. This assumption led to the formation of an error in determining the coordinate SA. It was mainly supposed in recent works that the amplitude of the intensity of the wave is inversely proportional to the square of the distance between the sound source and receiver. As a result of this assumption, the intensity of the wave recorded by the first microphone will change its values when fixed by subsequent microphones depending on the distance to the source of SA. As a result of such attenuation, the amplitude of intensity of the wave will change by:

$$\Delta i = I_{\max} \left[\frac{1}{(r + AL)^2} - \frac{1}{r^2} \right]. \quad (4)$$

For scientifically based choice of equipment for experimentally studying the propagation of the SA wave, the nature of its attenuation and modeling the random position of the point of origin of the source SA, it was proposed to use a rotary platform on thrust bearings driven by a stepper motor to position determine by eq. (1)-(4). The Matlab's model of stepper motor control is implemented using specially a hybrid two-phase model.

The input data as a specified parameter was changed in the dialog box in accordance experiment program as selected from the specific data. Each of the motor phases can be powered by two PWM converters on H-bridge MOS transistors. The auxiliary DC source is generated by a DC voltage drop in range from 12 to 28 V. Two not connected controllers independently measure and control currents by comparators. The MOSFET drives signal also are referenced by these controllers. The current pulsation is controlled by the hysteresis band of the comparators. In the formation of frequency-modulated output signals, their variable is selected depending on the parameters of the engine and generated by the simulation algorithm.

The single-phase excitation scheme, the formation of a rectangular voltage drop pulse with the specified parameters of amplitude and time length according to the values specified in the dialog box makes such a scheme convenient and easy to conduct experimental simulations based on Monte Carlo simulation. The motion and rotation position of the stepper motor rotor is controlled by STEP and DIR input modulated DC voltage from the output Signal Builder unit. In accordance with values from the range, 0 to 2 A and 0 to 500 steps/s in correspondence value are chosen in dialog mask will be selected current amplitude and step speed. In the fig. 2 are shown an example STEP and DIR input and on fig. 3 is shown the oscillogram of shaft rotation. This oscillogram demonstrates the possibility to modulate the area of growth, of the constancy and of the fall.

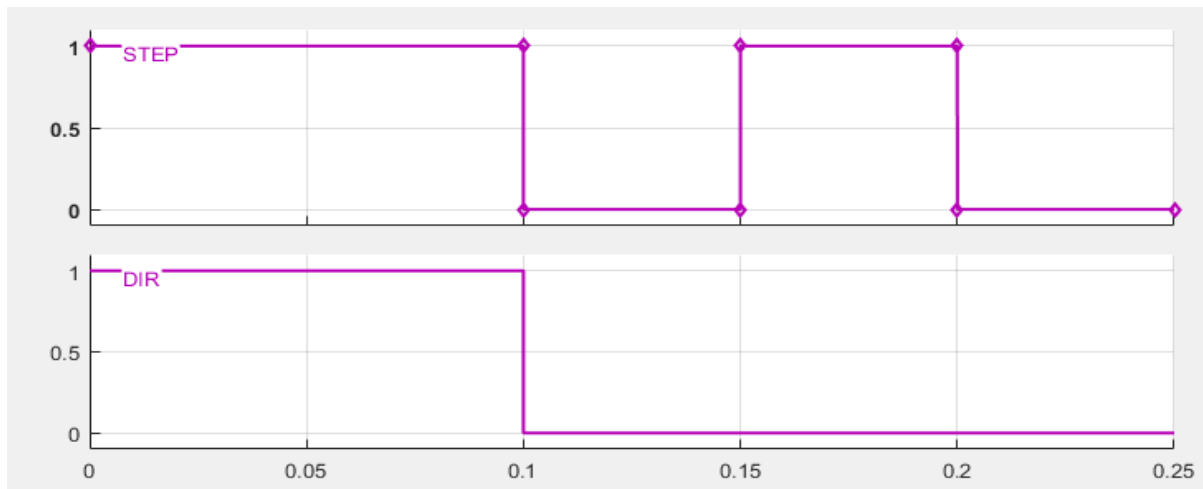


Figure 2: Appearance of signals for stepper motor control

The cover of the oscillogram fig.2 as time function is shown how a set of values 0 and 1 in STEP and DIR output from the Signal Builder effected on the character of the output signal: 1.0 a positive value initiates the rotation of the motor; a zero value stops the rotation. The DIR generally controls the direction of rotation. A positive value (1.0) imposes a positive direction, and a zero value initiates the opposite rotation. The operation of the stepper motor drive is illustrated by the main characteristics of the signals: voltage drop, current, torque, angular position, speed, and time decrement. Each of these characteristics a function of time displayed on the Scope unit.

For simulation was performed fixed decrement of time 1 μ s. As was determined by the analysis of the results of the experiments this value is quite good enough because provides acceptable accuracy for PWM. To distinguish higher PWM accuracy is required, a decreased decrement of time, but for this case, the simulation will be so slower. The one type of realized signal applied to the engine are shown in Fig. 2. The process of rotation of the rotor by a stepper motor in Matlab is displayed as an oscillogram: the angle defined in degrees as a function of time. An example showing the reaction to random numbers of the angle of rotation of the platform is presented in Fig. 3. The engine from the initial position begins to rotate the rotor clockwise to a time of 0.105 ms, then fixed its position for 0.05 ms, and 0.155 ms changes the direction of rotation counterclockwise to 0.205 ms, and then fixed again..

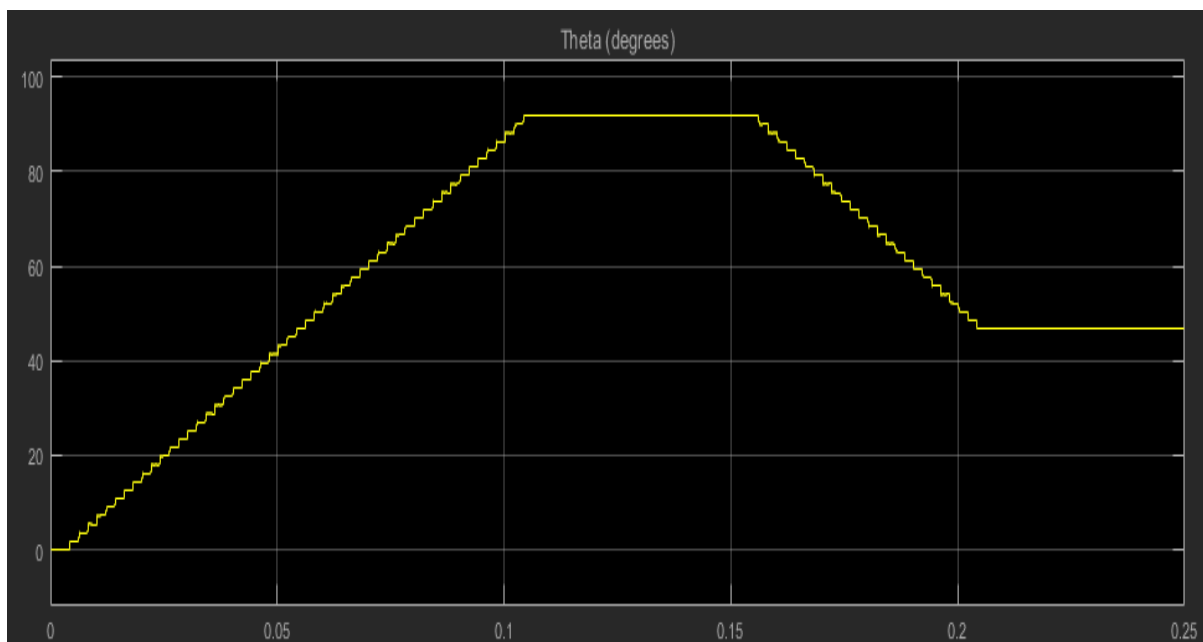


Figure 3: Oscillogram of shaft rotation

Such fragments of the presented movements correspond to a random set of numbers generated by simulation. Thus, the simulation of the rotational movements of the platform confirms the suitability of the system, which is proposed for physical modeling. Thus, the error resulting from the hypothesis of intensity constancy, causes an error in determining the time according to the echogram:

$$\Delta t = \frac{\Delta i}{\nabla i} \Big|_{i=a} \quad (5)$$

However, if the CS of the microphones is located on the UAV, the coordinates of which are determined by the GPS with an error, then due to the movement there are additional components that should be considered together with the static error.

Consider the case of monitoring and control UAV when the physical field wave intensity is scalar and inhomogeneous. In this case, the error of its measurement is determined in the first approximation by field fluctuations and sensor properties:

$$\Delta i = \left| \overline{grad i} \right| \left(\frac{\Delta r}{2} + \frac{\Delta r_g}{2} + v_i t_{np} \right) + (\Delta F_{CT}) = \sqrt{\frac{3}{2}} \left| \frac{\partial i}{\partial l} \right|_{\max} \left[|\Delta l|_{\max} + \Delta l_g + \bar{v}_i t_{np} \right] + |\Delta i_{CT}|, \quad (6)$$

where r – generalized coordinate, l – coordinate in the direction of greatest intensity increase, Δl_g – maximum size of the measuring part of the microphone in this direction, v_i – UAV speed, t_{np} – signal conversion time, Δi_{CT} – static sensor intensity error.

The last expression shows that the error of determining the intensity field using sensors located on the UAV depends not only on the static error $|\Delta i_{CT}|$, the value sensor, but also on the time constant, the error of the speed sensors, and hence the acceleration. The latest shows that inflated requirements for static error of sensors are not always justified. In addition, it is possible to reduce this component of the error by reducing the size of the sensor as an averaging zone. Despite the increase in error, the advantage of this system is that due to its installation on an inexpensive UAV, it is relatively inexpensive and maneuverable and durable. The latter allows you to read data at different heights. This fact allows us to obtain information bypassing the natural terrain and other obstacles, such as trees, buildings, etc., while the UAVs themselves can be at a safe distance and height and even in the rear, both outside the affected area and over enemy positions.

4.2. Comparative analysis of estimates of the maximum possible error of algorithms for symmetric and asymmetric echograms

The formation of the algorithm for determining the time of the phases of echograms according to the sound time series, which are registered by different microphones, is the main source of errors [5-9] as a static component. In this regard, it was determined to develop a key element of the general algorithm of the module program [5]. The task of the module is to establish conditions that ensure the uniqueness and uniformity of fixation of the identical moment of time in different echograms of one SA. Consideration of SA, in the form of a long-term function of intensity change, leads to the hypothesis of phase invariance of any of its points during propagation. Based on the above, it was assumed that such an echogram point was found, and there is an unambiguous relationship to determine it. Based on one of the conclusions of [5] on the calibration procedure, as a tool to expand the functions and capabilities of the monitoring CS, a four-point schedule using the method of recurrent approximation (MPA) was used. According to the expression of the analytical series of MPA [5], the approximate value of the time of entry of the acoustic wave is calculated. In the case where the discrepancy between two consecutive moments of time does not satisfy the required accuracy, the approximation is repeated until the required convergence is reached.

Thus, the introduction of the intensity of the SA as a continuous function that is integrated with the square, the norm of which allows several ways to determine the characteristic phase according to the calibration of multipliers λ_{oi} , as one that uniquely determines the time of registration. Analysis of the results of comparing the amplitude of the wave and the norm, the center of gravity or the relative location of the point of length as an unambiguous criterion on which to base the algorithm for fixing SA is set as one of the most important tasks to minimize error.

To establish practical recommendations for the choice of the algorithm for fixing the time of entry of the same phase of the wave into the microphone, simulation modeling was used, in particular the Monte-Carlo method. When placing the microphones at the vertices of an equilateral triangle at a distance of 0.6 m and a distance of 100 m from the event, the magnitude of the error was analyzed. A simple SA, represented by the triangular law of intensity distribution, with symmetric and asymmetric law of change of intensity in time and their composition, was considered. The distributions and intensity parameters of several examples are given for symmetric dependence in table 1 for asymmetric in table 2, and for the composition in table 3.

Table 1

Dependences of intensity distribution on time for symmetric SA

№	Relative time, calculated from the beginning of SA, t, s	Relative intensity			
		Example 1	Example 2	Example 3	Example 4
1	0	0	0	0	0
2	0,1	0,2	0,2	0,4	0,5
3	0,2	0,4	0,6	0,8	0,7
4	0,3	0,6	0,8	0,9	0,8
5	0,4	0,8	0,9	1	0,9
6	0,5	1	1	1	1
7	0,6	0,8	0,9	1	0,9
8	0,7	0,6	0,8	0,9	0,8
9	0,8	0,4	0,6	0,8	0,7
10	0,9	0,2	0,2	0,4	0,5
11	1	0	0	0	0

Table 2

Dependences of intensity distribution on time for asymmetric SA

№	Relative time, calculated from the beginning of SA, t, s	Relative intensity			
		Example 1	Example 2	Example 3	Example 4
1	0	0	0	0	0
2	0,1	0,3	0,6	0,8	0,35
3	0,2	0,6	0,9	0,9	0,7
4	0,3	0,8	1	1	1
5	0,4	1	0,9	0,8	0,9
6	0,5	0,9	1	0,9	0,8
7	0,6	0,6	0,9	0,7	0,5
8	0,7	0,4	0,8	0,5	0,4
9	0,8	0,3	0,6	0,4	0,3
10	0,9	0,2	0,2	0,2	0,1
11	1	0	0	0	0

Table 3

Dependences of intensity distribution on time for the composition SA

№	Relative time, calculated from the beginning of SA, t, s	Relative intensity			
		Example 1	Example 2	Example 3	Example 4
1	0	0	0	0	0
2	0,1	0,3	0,6	0,8	0,2
3	0,2	0,6	0,9	0,9	0,7
4	0,3	0,8	1	1	1
5	0,4	1	0,7	0,8	0,7
6	0,5	0,8	1	0,9	0,8
7	0,6	0,6	0,7	0,7	0,5
8	0,7	0,7	0,6	0,5	0,7
9	0,8	0,4	0,5	0,4	0,5
10	0,9	0,2	0,3	0,2	0,1
11	1	0	0	0	0

The results of calculations showing the effect of influences of distance from microphone to the source of the SA on the error are presented in table 4.

Table 4

Analysis of the influence of intensities and parameters of SA on the magnitude of the error

Example SA	Distance to SA, r, m	Maximum intensity error	Absolute time error Δt , s	Relative coordinate error, %
1	25	7,41E-05	1,74121E-05	0,0689
1	50	9,43E-06	1,0943E-05	0,0538
1	75	2,81E-06	1,02811E-05	0,0532
1	100	1,19E-06	1,01189E-05	0,0531
2	25	7,41E-05	1,37061E-05	0,0789
2	50	9,43E-06	1,04715E-05	0,054
2	75	2,81E-06	1,01405E-05	0,0532
2	100	1,19E-06	1,00595E-05	0,0531
3	25	7,41E-05	1,37061E-05	0,1719
3	50	9,43E-06	1,04715E-05	0,0343
3	75	2,81E-06	1,01405E-05	0,0273
3	100	1,19E-06	1,00595E-05	0,0267
4	25	7,41E-05	2,48242E-05	0,0327
4	50	9,43E-06	1,1886E-05	0,0215
4	75	2,81E-06	1,05621E-05	0,0213
4	100	1,19E-06	1,02379E-05	0,0212

As evidenced by the analysis of the calculation results (see table 4) for symmetric SA (examples of table 1), the distance to SA has a significant effect on the magnitude of the error. Thus, when recording the intensity, it is assumed that the maximum value of the intensity of the wave coming to each of the microphones was assumed to be constant. The latter assumption also makes a significant contribution to the magnitude of the overall methodological error, as the intensity decreases as the wave propagates from the first microphone to the next. It is known that almost all [5] algorithms do not take this fact into account. In addition, the time used by the system is synchronized with the time of reading information from the microphones, and therefore, as a second factor, introduces a time

error equal to half the sample time of reading. Thus, in the simulation, simulating the changes in intensity, the value of the maximum possible error of the maximum intensity, the absolute error of the time and the relative error of the coordinate SA are estimated. The results of the analysis show that for the selected four examples of symmetric distributions of intensity SA, the same regular changes in the dependence of the maximum possible error are observed. Increasing the distance between the microphones increases the amount of error, but it should be noted that this impairs the mobile quality of the system. Data demonstrating the grounds for such conclusions are presented in table 5 for example 4 of the symmetric variant of echograms. Thus, with increasing distance between the microphones, significantly increases the maximum possible error and, as a consequence, the relative error of determining the coordinates.

Table 5

Analysis of the influence of the distance between the microphones on the error? algorithm 5

№	The distance between mic- rophones a , m	Distance to SA, r , m	Maximum intensity error	Absolute time error Δt , s	Relative coordinate error, %
1	0,2	25	2,53E-05	1,10118E-05	0,0145362
2	0,6	25	1,29648E-05	7,41E-05	0,0206952
3	1,0	25	0,000121	1,48284E-05	0,0287785
4	1,4	25	0,000165	1,66079E-05	0,0372271

To select the algorithm for determining the time, the simulation and comparison of processes by the criterion of the magnitude of the relative error, the number of operations and the time required to fully calculate the time of entry of the wave into the microphone. In [5-9] the authors established that the determining factor of the three criteria is the accuracy of coordinate determination. Its value is determined by the error of fixing the time of entry of the wave to each of the microphones. In this regard, four algorithms were considered.

Algorithm 1.

The calculated time is the arithmetic mean time between the moment of time at which the value of 0.5 of the maximum intensity in the jump phase is reached and the time at which the intensity decreases to 0.5 of the maximum intensity.

Algorithm 2.

The calculated time is the time of the center of gravity of the echogram at ten points, and its countdown is from the point in time at which the value of 0, 5 maximum intensity is reached.

Algorithm 3.

The calculated time is the time of the arithmetic mean between the time of the jump phase from 0, 2 values of maximum intensity and level to 0.2 of its decrease in maximum intensity.

Algorithm 4.

The estimated time is according to the recurrent algorithm according to MPA [5] for which the estimated time is as the time of equivalent areas, according to the recurrent relation.

First of all, their work was studied and the influence of factors on the error for each of the four algorithms was studied. The properties of these algorithms were studied. The general conclusion of this analysis is that the generalized first and third algorithms are the simplest and most effective for symmetric intensity distributions. The use of a given value of the constant from 0.2 to 0.5 in the formal essence of the algorithm does not change. However, it leads to unified values of zero and one. In addition, the choice of a constant value of 0.2 significantly reduces the total calculation time, which is critical for specific applications. Algorithm 4 is efficient in terms of action and accuracy, but requires the organization of complex calculations. The accuracy of algorithm 4 is justified for complex non-symmetric intensity distributions, and for symmetric simple SA it is advisable to use algorithm 1.

4.3. Analysis of algorithms for recognizing and determining the time of complex echograms

On the basis of the conducted analysis the generalized algorithm of definition of time of occurrence of a sound wave of difficult structure ON as follows was formulated.

Algorithm 5. The calculated time is the arithmetic mean time between the time of the jump time from 0.2 to the value of the maximum intensity and the first subsequent maximum. To analyze the advantages and disadvantages of Algorithm 4, we analyze the effect of the distance between the microphones on the amount of error when using Algorithm 4.

Table 6

Analysis of the influence of distance between the microphones on the error value using algorithm 4

№	The distance between microphones a, m	Distance to SA, r, m	Maximum error intensity, Δi	Absolute error of time $\Delta t, c$	Relative error of the coordinates of its decline, %	Total calculation time $\Delta t_{min}, s$
1	0,2	25	2,53E-05	1,253E-05	0,000143128	0,20002265
2	0,6	25	7,41E-05	1,741E-05	0,000182726	0,20004716
3	1,0	25	0,000121	2,207E-05	0,023603	0,20007036
4	1,4	25	0,000165	2,652E-05	0,029351	0,20009259

To confirm this conclusion, we present the calculations by algorithm 4 (see Table 6). Thus, the calculations of the variant of table 5 show a significant reduction in error due to the choice of algorithm 5. However, the calculation of time costs shows that in the case when the origin is unknown and to be determined, the simple algorithm 5 prevails because it has a millisecond time.

Thus, when the accuracy of determining the value of the coordinate by algorithm 5 is sufficient, and the time of determining the coordinate is critical, the advantage over algorithm 4 is algorithm 5. In addition, it should be noted that for MES, military, with an error of centimeters, time is decisive, and for other purposes there is an accuracy of coordinates of a source SA.

5. Conclusions

1. The task of simulation modeling and estimation of the magnitude of the components of the maximum possible error that occurs when determining the coordinates of SA is set and carried out. Comparative analysis of the simulation results shows that the distance to the source SA and the relative position and distance between the microphones of the CS significantly affects the amount of static error. Thus, when increasing the relative distance from 41.66 to 82.33, the error decreases from 0.0689% to 0.0538%, and then almost unchanged for the case of symmetric echograms.

2. The choice of the algorithm does not significantly affect the estimation of the static maximum possible error in determining the coordinates SA for symmetric echograms, however for asymmetric echograms the choice of the algorithm significantly recognizes its value.

3. The algorithm for determining the moment of time of the identical phase of the occurrence of the wave front is determined by the required accuracy or the required calculation time. For critical accuracy, more preferable is the MPA algorithm 4, and at critical time, acquires algorithm 5.

6. References

- [1] Trunov, A., Byelozyorov, Z. (2020). Forming a method for determining the coordinates of sound anomalies based on data from a computerized microphone system. *Eastern-European Journal of Enterprise Technologies*, 2 (4 (104)), 38–50. doi: <http://doi.org/10.15587/1729-4061.2020.201103>
- [2] The Pentagon received at its disposal the first robot mule, 2022. URL: <https://phys.org/news/2012-12-legged-squad-ls3-darpa-four-legged.html>
- [3] A. Trunov, P. Kazan, V. Aliksieiev, O. Korolova, O. Sliusarenko, I. Dronyuk, Functioning Model of The Ground Robotic Complex, in: *Proceedings of The International Scientific and Technical Conference on Computer Sciences and Information Technologies (CSIT'21)*, IEEE, 2, p. 128–131. 2021, doi: 10.1109/CSIT52700.2021.9648595.
- [4] Axel Michelsen and Ole Næsbye Larsen. Pressure difference receiving ears. Topical review. IOP publishing. bioinspiration & biomimetics Bioinsp. Biomim. 3 (2008) 011001 (18pp) doi:10.1088/1748-3182/3/1/011001
- [5] Zh.Byelozyorov, A. Trunov, Increasing quality of the wireless module for monitoring and supervision of sound series of the expanded purpose, *Eastern-European Journal of Enterprise Technologies*, vol. 6/5 (114) (2021), 28–40. DOI: <https://doi.org/10.15587/1724061.2021.247658>
- [6] R. Kochan, O. Kochan, B. Trembach, P. Falat, K. Warwas. Theoretical Error of Bearing Method in Artillery Sound Ranging *Proceedings of the 2019 10th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, IDAACS 2019*, 2019, 2, стр. 615–619, 8924450
- [7] R. Kochan, B. Trembach, O. Kochan. Methodical error of targets bearing by sound artillery intelligence system ISTCMTM. *Scientific journal "Measuring Equipment and Metrology. Lviv 2019; Volume 80(3) : pp.10-14.* <https://doi.org/10.23939/istcmtm2019.03.010>
- [8] Kamil Krasuski, Damian Wierzbicki. Monitoring Aircraft Position Using EGNOS Data for the SBAS APV Approach to the Landing Procedure. *Journal Sensors* 20(7):1945. March 2020 DOI:10.3390/s20071945
- [9] Ciećko, A.; Grunwald, G. The comparison of EGNOS performance at the airports located in eastern Poland, *Technical Sciences* 2017, 20(2), pp. 181–198
- [10] L. Nguyen, J. Valls, X. Qiu: Multilevel B-Splines-Based Learning Approach for Sound Source Localization, *IEEE Sensors Journal* (2019). doi:10.1109/ JSEN.2019.2895854
- [11] D. Döbler, G. Heilmann, M. Ohm, Automatic detection of microphone coordinates, 3rd Berlin Beamforming Conference, 2010. URL: <http://www.bebec.eu/Downloads/BeBeC2010/Papers/BeBeC-2010-15.pdf>
- [12] J. Suern Yong, D.Y. Wang, Impact of noise on hearing in the military, *PMCID*, PMC4455974. PMID: 26045968 (2015). URL: <https://mmrjournal.biomedcentral.com/articles/10.1186/s40779-015-0034-5>.
- [13] C. Rascon, I. Meza, Localization of sound sources in robotics: A review, Elsevier Masson SAS, 2017. URL: <https://doi.org/10.1016/j.robot.2017.07.011>
- [14] X. Yang, H. Xing, X. Ji: Sound Source Omnidirectional Positioning Calibration Method Based on Microphone Observation Angle, *Complexity JOUR*, vol. 04. 2018. October. DOI: 10.1155/2018/2317853
- [15] A. I. Tkachenko "Fuzzy" Estimating of the In-Flight Geometric Calibration Parameters. *Том 51, Выпуск 3*, 2019, pp. 48-55. DOI: 10.1615/JAutomatInfScien

Method for Generating Pseudorandom Sequence of Permutations Based on Linear Congruential Generator

Emil Faure^{1,2}, Eugene Fedorov¹, Iryna Myronets¹ and Svitlana Sysoienko¹

¹ Cherkasy State Technological University, Shevchenko Blvd., 460, Cherkasy, 18006, Ukraine

² The State Scientific and Research Institute of Cybersecurity Technologies and Information Protection of the State Service for Special Communications and Information Protection of Ukraine, M. Zaliznyaka Str., 3, bl. 6, Kyiv, 03142, Ukraine

Abstract

The results of the study of the graph of states of a linear congruential generator (LCG) are considered and theoretically substantiated. A model of a generalized graph of LCG states has been developed. It represents each connected component of the graph in the form of cycles equipped with tree products, allows classifying the types of connectivity components of the graph of LCG states and investigating the influence of parameters on its topology. A method for generating a pseudorandom sequence (PRS) of numbers based on the linear congruential method is presented. This method allows generating uniformly distributed numbers regardless of the topology of the graph of LCG states and, consequently, minimizing the time spent on choosing its parameters, and increasing the size of the space of their allowable values to achieve the maximum period. Computer implementation of the algorithm for generating PRS of permutations based on LCG with any type of graph of its states has allowed increasing the speed of the generator compared to the permutation generator using the modern Fisher-Yates algorithm.

Keywords

Pseudorandom sequence, permutation, shuffle, random interleaver, linear congruential generator, graph of states, monad, topology, monad graph

1. Introduction

Pseudorandom sequences (PRSs) [1-2] are widely used to solve a wide class of problems. In particular, they are used to protect information from unauthorized access [3-5], to control the integrity of information [6], to form signals that provide covert data transmission and implement the modeling of complex systems and objects [7-9], to form permutations for factorial coding of information [10-12].

The linear congruential generator, the linear-feedback shift register (LFSR), and the additive generator are the simplest ones in design, high performance and ones of the most commonly used generators. Their design is often the basis of high-quality pseudorandom number generators (PRNGs) with a long repetition period.

Linear congruential generator (LCG) is proposed by D. H. Lehmer in 1949 [13]. It implements a recurrent relation ($f: S \rightarrow S$ transition function) of the form:

$$s_i = |K \cdot s_{i-1} + C|_M, \quad (1)$$

where K is the multiplier;

C is the growth;

M is the module, $K, C, s_0 \in \mathbf{Z}_M$.

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: e.faure@chdtu.edu.ua (E. Faure); fedorovee75@ukr.net (E. Fedorov); i.myronets@chdtu.edu.ua (I. Myronets); s.sysoienko@chdtu.edu.ua (S. Sysoienko)

ORCID: 0000-0002-2046-481X (E. Faure); 0000-0003-3841-7373 (E. Fedorov); 0000-0003-2007-9943 (I. Myronets); 0000-0002-0009-337X (S. Sysoienko)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

Obviously, $s_i \in [0; M - 1]$, and the power of the space of LCG states $|S| = M$.

Congruential sequence always forms repeating cycles [1].

LCGs are very sensitive to changes in parameters. Numerous works on the theory and application of congruential generators are aimed at choosing their parameters and assessing the quality of the obtained PRSs. Among them are both classical [1, 14-17] and modern works [18-22].

At the same time, despite the large number of studies on the choice of LCG parameters and experimental evaluation of its properties, most of them are aimed at improving the "randomness" of the formed sequence and do not take into account the structure of space of LCG states S .

Among the works devoted to the analysis of the structure of the space of LCG states, one of the main ones is the thorough scientific work of G. Marsaglia [23]. In addition, the work [24] is devoted to the development of the theory of PRS construction based on the LCG and LFSR.

The analysis of the simplest PRNGs, LCG and LFSR, shows their limitations due to the need to select parameters in order to ensure necessary PRS statistical properties. In particular, the allowable values of the generator parameters to ensure the maximum PRS period are shown in Table 1.

Table 1
PRNG parameters to achieve the maximum PRS period

Method	Period	Valid values of PRNG parameters	Size of space of permissible values of PRNG parameters
Linear congruential method [13]	$T = M$	1) $Gcd(C, M) = 1$; 2) $ K - 1 _p = 0$ for $\forall$ simple $p : M _p = 0$; 3) if $ M _4 = 0 \Rightarrow K - 1 _4 = 0$	$\varphi(M) \cdot P$, where P is the number of K values that satisfy conditions 2 and 3
Method based on LFSR [25-27]	$T = 2^n - 1$	Generator polynomial $G_n(x)$ is primitive one	Corresponds to the number of primitive polynomials of n degree
Method for PRS generating based on concatenation of LCG cycles [24]	$T = M$	$Gcd(K, M) = 1$	$\varphi(M) \cdot M$
Method for PRS generating based on concatenation of LFSR cycles [24]	$T = 2^n$	Generator polynomial $G_n(x)$ generates a cyclic structure of the graph of LFSR states	Corresponds to the number of polynomials of n degree that generate the cyclic structure of the graph of LFSR states

The purpose of the work is to generate PRS of permutations with high performance and necessary statistical properties without the need to select LCG parameters.

For the further study and analysis of PRNG construction based on sequential traversal of all vertices of the graph of LCG states, it is necessary to perform an in-depth study of the structure of the graph of LCG states, extended analysis of the influence of LCG parameters on the structure of its graph, and, consequently, to develop the method for generating PRS based on linear congruential method by sequentially traversing the contour of the graph of states of the generator.

2. Review of the literature

To study and generalize the structure of the graph of LCG states, we shall explore the main approaches to forming graphs of states of PRS generating devices.

2.1. Cycle graphs

A cycle graph [28], also known as a simply n -cycle [29], is a graph containing n nodes and consisting of a single cycle that passes through all its nodes. The cycle graph is denoted by C_n . The number of vertices in C_n is equal to the number of edges, each vertex has a power of 2, any vertex is incidental to two edges.

Cycle graphs are used, for example, to illustrate the structure of multiplicative groups M_n (Figure 1). Such graphs are formed by creating numbered nodes, one for each α element of the surplus class, and constructing cycles obtained by calculating α^i for $i=1,2,\dots$. Each edge of such a graph has a bidirectional character [28].

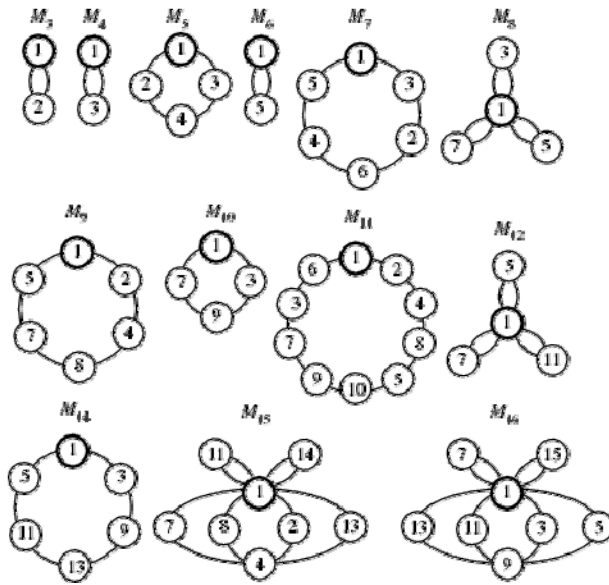


Figure 1: Graphs for some small-order multiplicative groups (from [30])

In all graphs of Figure 1 the node with $\alpha=1$ is highlighted because it is a zero cycle: $\beta_j = \left| \alpha^j \right|_M = \alpha = 1$ for any j . Next, we shall use the same notation to represent zero cycles in the graph of LCG states.

Note that in Figure 1 not all represented graphs are cycle graphs. This follows from the fact that multiplicative group by M module can be isomorphic to the product of several cyclic groups (for example, $M_8 \leftrightarrow C_2 \times C_2$, and $M_{15} \leftrightarrow C_2 \times C_4$). In this case, the graph is a combination of several cycle graphs.

To visually represent the structure of the graph of LCG states, it is necessary to use an oriented cycle graph, an oriented version of the cycle graph, in which all arcs are directed in the same direction.

2.2. Algebra of monads and topology of monad graphs

This paragraph, as well as its name, is based on the V. I. Arnold works [31-32].

According to [31], a monad is a representation of a finite set in itself. The monad graph has all elements of this finite set as vertices, and oriented edges connect each element with its image.

In other words, a monad graph is an arbitrary finite oriented graph, from each vertex of which exactly one edge emerges. Iterations of the monad lead any vertex to a cycle attractor.

According to [31], each connected component of the monad graph is a forest of root-oriented root trees, the roots of which are connected by an oriented cycle (topologically circular) from the edges connecting the tree roots.

In other words, connected components of any monad are cycle attractors, which are equipped with root trees attached by their roots to each vertex of the cycle attractor [32]. The number of vertices of the cycle can be equal to 1. In this case, the whole component is one root tree.

Each connectivity component of the graph of any representation of a finite set contains one and only one cycle.

Let S be a finite group, and $f: S \rightarrow S$ representation transforms each of its elements $s \in S$ according to the expression: $f(s) = |K \cdot s + C|_M$, where K is the multiplier; C is the growth; M is the module, $K, C, s_0 \in Z_M$.

In this paper, $f: S \rightarrow S$ representation will be called the monad of S group.

According to [31], the symbols O_n , A_n , T_m , E_n will denote the following oriented graphs:

- O_n = oriented cycle of n vertices;
- A_n = connected graph of $2n$ vertices, which is a cycle of n length, equipped with n single-edge trees, which are included one in each of n vertices;
- T_{2^n} = root tree with $2n$ vertices and n floors except the root, which branches binarily on $1, \dots, n-1$ floors; the root is considered to be the zero floor, and it also includes two edges: one is from itself and one is from a single vertex of the first floor;
- E_n = root tree with n vertices, from each of which the edge leads directly to the root (so that $E_2 = A_1 = T_2$);
- $D_n = 4n$ -vertex graph, consisting of O_n cycle of n length, equipped in each of its vertices with three input edges (form together with this corresponding to the cycle vertex the root tree $D_1 = E_4$).

For example, graphs of monads for additive cyclic groups in the field Z_n have the form shown in Figure 2.

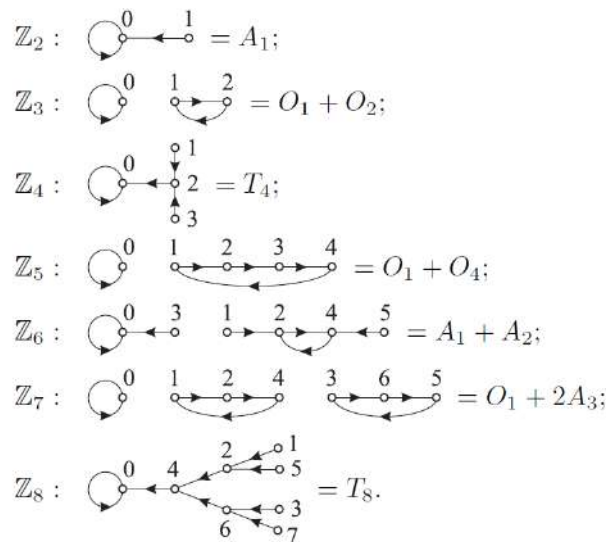


Figure 2: Graphs of monads for some simple cyclic additive groups (from [31])

According to [31], a monad that acts on the direct product $X * Y$ component by component: $(A * B)(x, y) = (Ax * By)$ is called the $A * B$ product of A and B monads that act on X and Y , respectively. The number of elements of a monad product is equal to the product of the number of elements of monad coefficients.

In [31] it is also shown that the graph of the monad product is the product of graph coefficients:
 $[graph(A * B)] = [graph(A)] * [graph(B)] / A_n = A_1 * O_n, D_n = D_1 * O_n$.

Multiplying any root tree T by O_n equips n -cycle O_n with root trees of T type with roots at all points in the cycle.

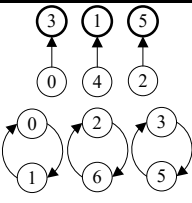
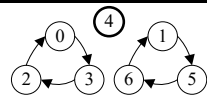
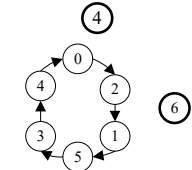
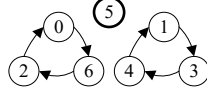
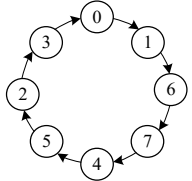
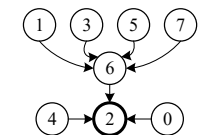
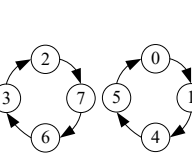
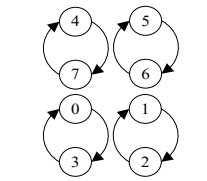
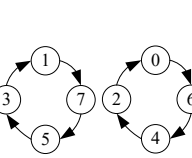
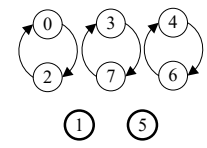
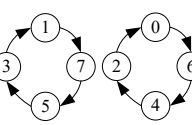
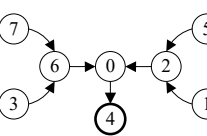
3. Materials and methods

In this section, we shall investigate the LCG topology and generalize the graph of its states to develop a method for generating permutation sequences.

3.1. Graphs of linear congruential generator

Examples of graphs of monads of S group for some LCG parameters are shown in Table 2.

Table 2
Oriented graphs of LCG states for some of its parameters

LCG parameters			Graph of states	LCG parameters			Graph of states
M	K	C		M	K	C	
6	4	3		7	2	3	
7	6	1		7	4	6	
7	3	2		8	4	2	
8	5	1		8	7	3	
8	3	1		8	7	2	
8	1	6		8	6	4	

Let us analyze the structures of graphs of monads of LCG S group. We shall summarize the analyzed structures and present in Table 3 some graphs typical to LCG.

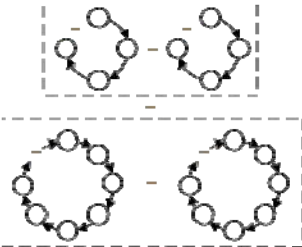
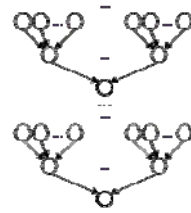
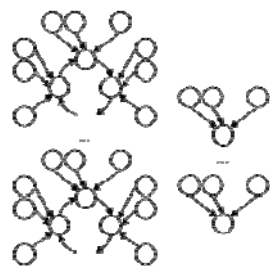
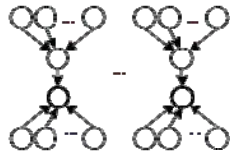
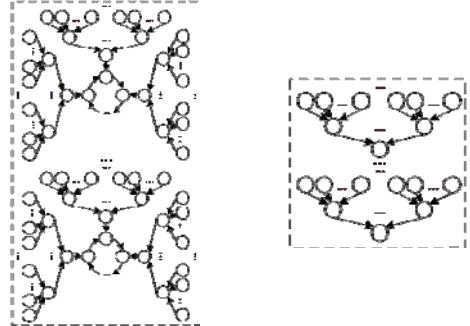
Extending the regularities given in [24], we give some graphs of LCG states and parameters for which these graphs are typical:

1. O_M – for:

- a) $M \geq 2, K = 1, C = 1;$
- b) $M = 2^p, K = 4l + 1, C = 2m + 1, l, m \in \mathbb{Z} \geq 0;$
- c) M is simple, $K = 1, C \geq 1;$

Table 3

Generalized graphs of LCG states

1. Generalized graph of cycles		
Groups of cycles of t_i length $\left(i \geq 1, t_i \geq 1, \sum_i d_i t_i = M \right)$		$\sum_i d_i O_{t_i}$
2. Generalized graph of trees		
Group of a^n - vertex trees with n floors $(a \geq 2, da^n = M)$		dT_{a^n}
3. Combinations of cycles and trees		
A group of cycles of t length, equipped in each of its vertices with input edges $(n-1)$, and a group of root trees with n vertices, from each of which the edge leads directly to the root $(n(dt + k) = M)$		$d(E_n * O_t) + kE_n$ $(d(T_{n^1} * O_t) + kT_{n^1})$
A group of zero cycles with single-edge trees included in them, equipped in each of their vertices with root trees with n vertices, from each of which the edge leads directly to the root $(dn = M/2)$		$d(E_n * A_1)$ $(d(T_{n^1} * T_{2^1}))$
A group of cycles of t length, equipped with t a^n -vertex trees with n floors, and a^n group of a^n - vertex trees with n floors $(n \geq 2, a^n(dt + k) = M)$		$d(T_{a^n} * O_t) + kT_{a^n}$

2. dO_l – for $M = 2^p$, $K = 4l - 1$, $l \in \mathbb{Z} \geq 1$, and $C = 2m + 1$, $m \in \mathbb{Z} \geq 0$. Under these conditions:
 - a) for $l = 2^{k-2}$ ($K = 2^k - 1$), $k \geq 2$, $t = 2M/(K+1)$, and $d = (K+1)/2$ (or $t = 2^{p-l+1}$, $d = 2^{l-1}$) take place;
 - b) for $l = 2k - 1$ ($K = 8k - 5$), $k \geq 1$, $t = M/2$, and $d = 2$ take place;
 - c) for $l = 2k$, $k \geq 1$, $K \neq 2^r - 1$, $r \geq 3$ (that is $k \neq 2^{r-3} : k = 2^{i-4}(2j+1)$, and $l = 2^{i-3}(2j+1)$ ($K = 2^{i-1}(2j+1) - 1$) for $j \in [1; 2^{p-i} - 1]$, $i \in [4; p-1]$), $t = M/2^{i-2}$, $d = 2^{i-2}$ take place;
3. $dO_l + O_l$ – for simple M , $K \geq 2$, $C \in \mathbb{Z}$;
4. a tree with a root, zero cycle, for $M = 2^p$, $K \in \{2l : l \in \mathbb{N}, 0 < l < 2^{p-1}\}$, $C \in \{0, 1, \dots, 2^p - 1\}$, $p \in \mathbb{N}$

3.2. Generalized graph of states of linear congruential generator

Analysis of possible graphs of LCG states shows that they can all be reduced to a single configuration containing a set of d cycles of the same or different length, including zero cycles, and a set of d' precycles (trees) leading to cycles. Under these conditions

$$M = \sum_{i=1}^d t_i + \sum_{j=1}^{d'} t'_j, \quad (2)$$

where t_i is the length of the i^{th} cycle;

t'_j is the number of vertices in the j^{th} tree except for the root.

The generalized LCG graph can be represented as follows:

$$G_{LCG} = (V_{LCG}, A_{LCG}),$$

where V_{LCG} is the set of vertices of the graph;

A_{LCG} is the set of arcs of the graph.

In its turn,

$$V_{LCG} = \{v_k\} \cup \{v'_l\},$$

where $\{v_k\}$ is the set of vertices belonging to the cycles, $k = 1, 2, \dots, \sum_{i=1}^d t_i$;

$\{v'_l\}$ is the set of vertices belonging to the trees, except for their roots, $l = 1, 2, \dots, \sum_{j=1}^{d'} t'_j$;

$$A_{LCG} = \{a_k\} \cup \{a'_l\},$$

where $\{a_k\}$ is the set of arcs belonging to the cycles, $k = 1, 2, \dots, \sum_{i=1}^d t_i$;

$\{a'_l\}$ is the set of arcs belonging to the trees, $l = 1, 2, \dots, \sum_{j=1}^{d'} t'_j$.

According to [33], for a simple M , the expressions $d' = 0$, $t_i = t$ for $\forall i \in [1, d-1]$ and $t_d = 1$ are fair, and expression (2) takes the form: $M = \sum_{i=1}^{d-1} t + 1$.

We present a generalized graph of LCG states in terms of the theory of algebra of monads and the topology of monad graphs.

The analysis of typical oriented graphs of LCG states shows that no more complex patterns, except for the products of trees and cycles, are found in the graphs of monads of $f : S \rightarrow S$ representation. LCG graph is an inconnected combination of cycles equipped with tree products. Since $E_n = T_{n^1} * O_1$, $A_t = T_{2^1} * O_t$, $D_t = T_{4^1} * O_t$, each connected component of LCG graph can be represented as

$T_{a^n} * (T_{b^m} * O_t)$. For example, $O_t = T_{a^0} * (T_{b^0} * O_t)$, and $A_t = T_{a^0} * (T_{2^1} * O_t) = T_{2^1} * (T_{b^0} * O_t)$, $E_n = T_{a^0} * (T_{n^1} * O_t) = T_{n^1} * (T_{b^0} * O_t)$.

Then the generalized graph of LCG states has the form:

$$G_{LCG} = \sum_{i=1}^d d_i \left(T_{a_i^{n_i}} * (T_{b_i^{m_i}} * O_{t_i}) \right), \quad (3)$$

where d is the number of different types of connectivity components of the graph of LCG states;

d_i is the number of connectivity components of the graph of LCG states of the B^{th} type;

a_i, n_i, b_i, m_i, t_i are parameters of connectivity components of the graph of LCG states of the i^{th} type.

In this case $\sum_{i=1}^d d_i a_i^{n_i} b_i^{m_i} t_i = M$.

Determining the rules for calculating the number of graph components $\sum_{i=1}^d d_i$, their types d and values of numbers a_i, n_i, b_i, m_i, t_i through LCG parameters is beyond the scope of this work and requires further research.

3.3. Method for generating sequences of permutations based on linear congruential method

The method for generating LCG-based PRS is as follows.

1. If necessary (for example, to increase the speed of PRS generating or meet the requirements for the spatial complexity of the algorithm that implements the proposed method), the type of graph of LCG states, as well as the conditions to be met by K , C and M LCG parameters to obtain a given type of structure are determined. Determining the type of the graph of LCG states can be performed in accordance with typical graphs presented in Table 3. M parameter determines the area of determination of pseudorandom variable. If the choice of the type of the graph of states is not made, LCG parameters are determined arbitrarily, taking into account the restrictions imposed on them.
2. If necessary (for example, for non-consecutive cycles of the graph of LCG states without precycles (trees) to increase the speed of PRS generating), the representatives of each cycle of the generator (boot vectors (BVs)) are determined and stored in memory.
3. The current LCG BV is determined by (random or deterministic) choosing from the set of stored BVs, if this set is specified, or from the set of integers in the range $[0, M - 1]$.
4. A PRS is formed by the LCG with given parameters until the generator forms unique numbers. In the case of reappearance of any element (not necessarily equal to BV (for a graph containing continuous cycles without precycles (trees), equal to BV)) the generation of the current segment of the sequence is stopped.
5. A new current BV of LCG is determined by its (random or deterministic) choosing:
 - from the set of still unused BVs, if this set is specified;
 - from the set of integers of the range $[0, M - 1]$ except for the numbers present in the formed part of the PRS.
6. PRS is formed for a given BV until the reappearance of the element in the formed sequence (for a graph containing continuous cycles without precycles (trees), equal to the current BV).
7. Transition to item five until the shuffle of all BVs.
8. Transition to item three until the shuffle of all combinations consistently used in items three and five of BV (if they are set and stored in memory).
9. Transition to item two until the shuffle of all BV combinations (if they are set and stored in memory).

Thus, the proposed method allows performing concatenation not only of separate and disjoint cycles in LCG graph, but also of precycles (trees), if they are contained therein.

In addition, the proposed approaches can be used to form PRS based on LFSR with an arbitrary generator polynomial. This is because the use of a reducible polynomial as a generator one for LFSR leads to an increase in the number and the change in the structure of connected components in the generator graph of states.

4. Experiments and results

The proposed approaches to constructing devices for PRS generating based on LCG are used to create software implementations of generators.

The size of the space of allowable values of LCG parameters for the above method is equal to M^2 . The comparison with the corresponding indicators of the analogues in table 1 shows that this method allows increasing the size of the space of allowable values of LCG parameters to achieve the PRS period $T = M$ in $M/\varphi(M)$ times.

Let us study the speed of software implementation of the permutation generator based on the developed method and compare it with the speed of the generator that implements the modern Fisher-Yates algorithm [34]. For the sake of objectivity, the generators were implemented on one platform and tested on one computer with fixed performance indicators. The results are presented in Figure 3.

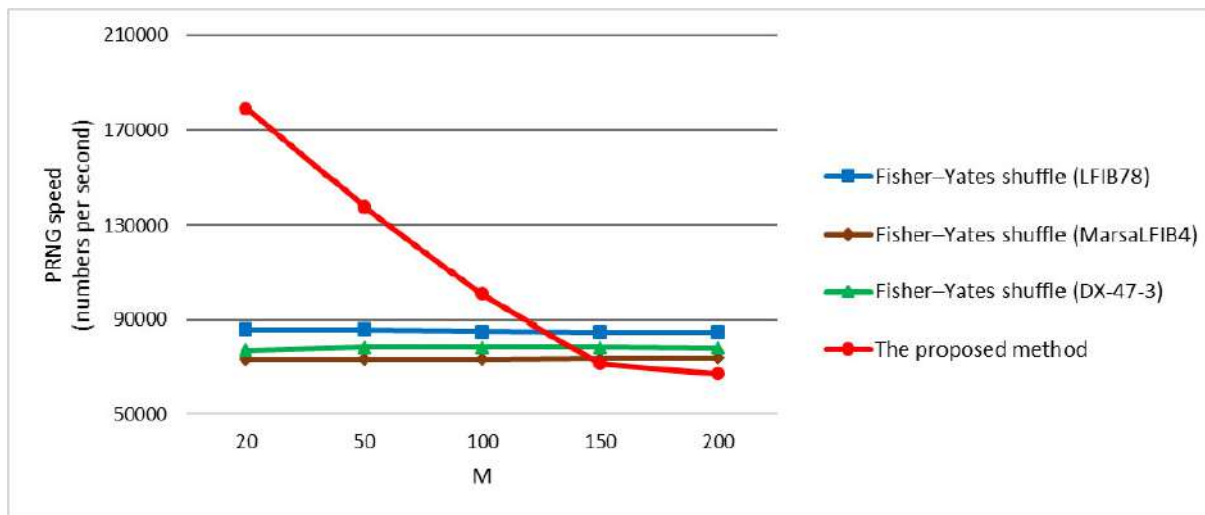


Figure 3: Graphs of dependence of speed of permutation generators on M value

The speed of the developed generator exceeds the speed of the permutation generator using the Fisher-Yates algorithm for $M \leq 125$, which expands the results obtained in [35].

It should be noted that PRS formed according to the proposed method is not cryptographically stable and can not be used in "pure" form in cryptographic transformations, for example, as a gamma for stream ciphers. However, the proposed approaches to PRS generating can be used to implement a multi-stage encryption procedure.

5. Conclusions

The scientific novelty of the study is as follows. It is developed the model for generalized graph of states of linear congruential generator, which allows to carry out the classification of types of connectivity components of the graph of its states and to investigate the influence of parameters on its topology. The developed model has allowed to improve the method for PRS generating based on linear congruential method, which allows to form PRS of uniformly distributed numbers regardless of the topology of the graph of states of linear congruential generator and, as a result, to minimize the

time spent on choosing its parameters and increase the size of the space of their allowable values to achieve the PRS period $T = M$ in $M/\varphi(M)$ times.

Implementation of the algorithm for generating PRS of permutations based on LCG with any type of the graph of its states has allowed to increase the speed of the generator compared to the permutation generator based on PRNG LFIB78 using the Fisher-Yates algorithm for the permutation order $M \leq 125$: in particular, for $M = 20$ – in 2.1 times; $M = 50$ – 1.6 times; $M = 100$ – 1.2 times.

6. Acknowledgements

This research was funded by the Ministry of Education and Science of Ukraine, grant number 0120U102607.

7. References

- [1] D. E. Knuth, The Art of Computer Programming, Vol. 2: Seminumerical Algorithms, 3rd ed., Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA, 1997.
- [2] F. James, L. Moneta, Review of high-quality random number generators, *Computing and Software for Big Science* 4 (1) (2020). doi: 10.1007/s41781-019-0034-3.
- [3] L. Crocetti, S. Di Matteo, P. Nannipieri, L. Fanucci, S. Saponara, Design and test of an integrated random number generator with all-digital entropy source, *Entropy* 24 (2) (2022). doi: 10.3390/e24020139.
- [4] X. Tong, X. Chen, S. Xu, Advances in superlattice cryptography research, [超晶格密码的研究进展] *Kexue Tongbao/Chinese Science Bulletin* 65 (2-3) (2020) 108-116. doi 10.1360/TB-2019-0291.
- [5] S. Sysoienko, I. Myronets, V. Babenko, Practical implementation effectiveness of the speed increasing method of group matrix cryptographic transformation, in: *CEUR Workshop Proceedings*, volume 2353, 2019, pp. 402-412.
- [6] S. Gnatyuk, V. Kinzeryavyy, K. Kyrychenko, Kh. Yubuzova, M. Aleksander, R. Odarchenko, Secure hash function constructing for future communication systems and networks, *Advances in Intelligent Systems and Computing* 902 (2020) 561-569.
- [7] M. Alawad, M. Lin, Survey of stochastic-based computation paradigms, *IEEE Transactions on Emerging Topics in Computing* 7 (1) (2019) 98–114. doi 10.1109/TETC.2016.2598726.
- [8] S. Bandyopadhyay, R. Bhattacharya, *Discrete and Continuous Simulation: Theory and Practice*, CRC Press, 2014. doi: 10.1201/b17127.
- [9] P. A. A. Resende, A. C. Drummond, A survey of random forest based methods for intrusion detection systems, *ACM Computing Surveys* 51 (3) (2018). doi: 10.1145/3178582.
- [10] J. S. Al-Azzeh, B. Ayyoub, E. Faure, V. Shvydkyi, O. Kharin, A. Lavdanskyi, Telecommunication systems with multiple access based on data factorial coding, *International Journal on Communications Antenna and Propagation* 10 (2) (2020) 102-113. doi: 10.15866/irecap.v10i2.17216.
- [11] E. Faure, A. Shcherba, Y. Vasiliu, A. Fesenko, Cryptographic key exchange method for data factorial coding, in: *CEUR Workshop Proceedings*, volume 2654, 2020, pp. 643–653.
- [12] E. Faure, A. Shcherba, B. Stupka, Permutation-based frame synchronisation method for short packet communication systems, in: *Proceedings of the 11th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, IDAACS 2021*, volume 2, 2021, pp. 1073–1077. doi: 10.1109/IDAACS53288.2021.9660996.
- [13] D. H. Lehmer, Mathematical methods in large-scale computing units, in: *Proceedings of a Second Symposium on Large-Scale Digital Calculating Machinery*, Cambridge, Mass., 1949, pp. 141–146.
- [14] W. Freiberger, U. Grenander, *A Short Course in Computational Probability and Statistics*, Springer, New York, 1971.
- [15] M. Greenberger, An a priori determination of serial correlation in computer generated random numbers, *Mathematics of Computation* 15 (76) (1961) 383.

- [16] T. E. Hull, A. R. Dobell, Random number generators, *SIAM Review* 4 (3) (1962) 230-254.
- [17] C. F. Gauss, J. Brinkhuis, C. Greiter, *Disquisitiones Arithmeticae*, reissue ed., Springer, New York, 1986.
- [18] E. V. Faure, A. I. Shcherba, V. M. Rudnytskyi, The method and criterion for quality assessment of random number sequences, *Cybernetics and Systems Analysis* 52 (2) (2016) 277-284. doi: 10.1007/s10559-016-9824-3.
- [19] K. Tsuchiya, Y. Nogami, Long period sequences generated by the logistic map over finite fields with control parameter four, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* E100.A (9) (2017) 1816-1824. doi: 10.1587/transfun.E100.A.
- [20] E. Faure, I. Myronets, A. Lavdanskyi, Autocorrelation criterion for quality assessment of random number sequences, in: *CEUR Workshop Proceedings*, volume 2608, 2020, pp. 675–689.
- [21] P. L'Ecuyer, P. Wambergue, E. Bourceret, Spectral analysis of the MIXMAX random number generators, *INFORMS Journal on Computing* 32 (1) (2020) 135-144. doi: 10.1287/IJOC.2018.0878.
- [22] G. Konsam, M. Loukrakpam, Triple linear congruential generator-based hardware-efficient pseudorandom bit generation, in: T. R. Lenka, D. Misra, A. Biswas (Eds.), *Micro and Nanoelectronics Devices, Circuits and Systems*, volume 781 of *Lecture Notes in Electrical Engineering*, Springer, Singapore, 2022. doi: 10.1007/978-981-16-3767-4_22.
- [23] G. Marsaglia, The structure of linear congruential sequences, in: S. K. Zaremba (Ed.), *Applications of Number Theory to Numerical Analysis*, Academic Press, 1972, pp. 249-285.
- [24] A. Lavdanskyi, *Methods and Tools of Forming Pseudorandom Sequences for Computer Cryptography*, Ph.D. thesis, Cherkasy State Technological University, Cherkasy, Ukraine, 2017.
- [25] R. C. Tausworthe, Random numbers generated by linear recurrence modulo two, *Mathematics of Computation*, 19 (90) (1965).
- [26] S. W. Golomb, *Shift Register Sequences: Secure and Limited-Access Code Generators, Efficiency Code Generators, Prescribed Property Generators, Mathematical Models*, 3rd revised ed., Aegean Park Press, Laguna Hills, CA, USA, 2017.
- [27] H. T. Nguyen, G. N. Pham, A. N. Bui, B. A. Nguyen, N. T. Le, H. T. Pham, Linear feedback shift register and its applications in digital system design, *International Journal of Emerging Technology and Advanced Engineering* 11 (11) (2021) 204-208. doi: 10.46338/IJETAE1121_24.
- [28] D. Shanks, *Solved and Unsolved Problems in Number Theory*, 4rd. ed., AMS Chelsea Pub., New York, 2002.
- [29] S. Pemmaraju, P. S. Skiena, *Computational Discrete Mathematics, Combinatorics and Graph Theory with Mathematica*, 1st ed., Cambridge University Press, Cambridge, U.K., New York, 2003.
- [30] Modulo Multiplication Group, (25.02.2022).
URL: <http://mathworld.wolfram.com/ModuloMultiplicationGroup.html>.
- [31] V. I. Arnold, Topology of algebra: Combinatory of squaring, *Functional Analysis and Its Applications* 37 (3) (2003) 177-190.
- [32] V. I. Arnold, Topology and statistics of formulae of arithmetics, *Uspekhi Matematicheskikh Nauk* 58:4 (352) (2003) 3-28.
- [33] T. V. Mityankina, V. V. Shvydkiy, A. I. Shcherba, Randomization of a sequence of congruential numbers, *Bulletin of the Engineering Academy of Ukraine* 2 (2008) 107-111.
- [34] Paul E. Black, Fisher-Yates shuffle, in *Dictionary of Algorithms and Data Structures* (2019).
URL: <https://www.nist.gov/dads/HTML/fisherYatesShuffle.html>.
- [35] F. Panca Juniawan, H. Arie Pradana, Laurentinus, D. Yuny Sylfania, Performance comparison of linear congruent method and Fisher-Yates shuffle for data randomization, *Journal of Physics: Conference Series*, 1196 (1) (2019). doi:10.1088/1742-6596/1196/1/012035.

Convergence of adaptive operator extrapolation method for operator inclusions in Banach Spaces

Serhii Denysov, Vladimir Semenov

Taras Shevchenko National University of Kyiv, 64/13 Volodymyrska Street, Kyiv, 01161, Ukraine

Abstract

A novel iterative splitting algorithm for solving operator inclusions problem is considered in a real Banach space setting. The operator is a sum of the multivalued maximal monotone operator and the monotone Lipschitz continuous operator. The proposed algorithm is an adaptive modification of the “forward-reflected-backward algorithm” [14]. Step size update rule not require Lipschitz constant knowledge of the operator. Advantage of the proposed algorithm is a single calculation of the maximal monotone operator resolvent and value of the monotone Lipschitz continuous operator on each iterative step. Weak convergence of the method is proved for operator inclusions in 2-uniformly convex and uniformly smooth real Banach space.

Keywords

Maximal monotone operator, operator inclusion, variational inequality, splitting algorithm, convergence, adaptability, 2-uniformly convex Banach space, uniformly smooth Banach space

1. Introduction

Consider real Banach space E . Denote it's dual space as E^* . We study monotone operator inclusion:

$$\text{find } x \in E: 0 \in (A + B)x, \quad (1)$$

where $A: E \rightarrow 2^{E^*}$ is multivalued maximal monotone operator, $B: E \rightarrow E^*$ is monotone, single-valued, and Lipschitz continuous operator.

Many actual problems can be written in the form of (1). Among them are variational inequalities and optimization problems from various fields of optimal control, inverse problem theory, machine learning, image processing, operations research, and mathematical physics [1–6]. Prominent example is the saddle problem that plays an important role in mathematical economics:

$$\min_{p \in P} \max_{q \in Q} F(p, q),$$

where $F: P \times Q \rightarrow R$ is a smooth convex-concave function, $P \subseteq R^n$, $Q \subseteq R^m$ – closed convex sets, which can be formulated as

$$\text{find } x \text{ such that } 0 \in (A + B)x,$$

where $x = (p, q) \in R^{n+m}$ and

$$Ax = \begin{pmatrix} N_P p \\ N_Q q \end{pmatrix}, \quad Bx = \begin{pmatrix} \nabla_1 F(p, q) \\ -\nabla_2 F(p, q) \end{pmatrix},$$

N_P (N_Q) is a normal cone operator for closed convex set P (Q) [1].

Research and development of algorithms for operator inclusions (1) and related problems is a rapidly growing field of applied modern nonlinear analysis [1–4, 7–25].

The most well-known and popular iterative method for solving monotone operator inclusions (1) in Hilbert space is the “forward-backward algorithm” (FBA) [1, 11, 12]

$$x_{n+1} = J_{\lambda}^A(x_n - \lambda Bx_n),$$

where $J_{\lambda}^A = (I + \lambda A)^{-1}$ is the operator resolvent, $A: H \rightarrow 2^H$, $\lambda > 0$.

Note that the FBA scheme includes well-known gradient method and proximal method as special cases [1]. For inverse strongly monotone (cocoercive) operators $B: H \rightarrow H$ FBA method is weakly converging [1]. However, FBA may diverge for Lipschitz continuous monotone operators B .

The condition of the inverse strong monotonicity of the operator B is a rather strong assumption. To weaken it, Tseng [13] proposed the following modification of the FBA:

$$\begin{cases} y_n = J_{\lambda}^A(x_n - \lambda Bx_n), \\ x_{n+1} = y_n - \lambda(By_n - Bx_n), \end{cases}$$

where $B: H \rightarrow H$ – monotone and Lipschitz continuous operator with constant $L > 0$ and $\lambda \in (0, L^{-1})$.

A main limitation of Tseng method is their two calls of B per iteration. Evolution of this idea resulted in the so-called “forward-reflected-backward algorithm” [14]:

$$x_{n+1} = J_{\lambda}^A(x_n - \lambda Bx_n - \lambda(Bx_n - Bx_{n-1})), \quad \lambda \in \left(0, \frac{1}{2L}\right),$$

and related method [15]:

$$x_{n+1} = J_{\lambda}^A(x_n - \lambda Bx_n) - \lambda(Bx_n - Bx_{n-1}), \quad \lambda \in \left(0, \frac{1}{3L}\right).$$

A special case of this algorithm is the algorithm “optimistic gradient descent ascent”, which is popular among specialists in the field of Machine Learning [14, 15].

These schemes are naturally called operator extrapolation schemes. Note that in the three above algorithms there is a requirement to know operator’s Lipschitz constant – which is often unknown or difficult to estimate. To overcome these difficulties, adaptive rules for updating parameter $\lambda > 0$ on each step have been proposed [16].

Some progress has been achieved recently in the study of splitting algorithms for operator inclusions in Banach spaces [2, 17–25]. This progress is mostly related to careful use of theoretical results and concepts from Banach spaces geometry [2, 26–29]. Book [2] contains an extensive material on this topic.

In the article [17] for the solution of inclusions (1) in the 2-uniformly convex and uniformly smooth Banach space the next algorithm is proposed:

$$x_{n+1} = J_{\lambda}^A \circ J^{-1}(Jx_n - \lambda Bx_n), \quad (2)$$

where $J_{\lambda}^A = (J + \lambda A)^{-1} J$ is the resolvent of operator $A: E \rightarrow 2^{E^*}$, $\lambda > 0$, J is the normalized duality mapping from E to E^* . For inverse strongly monotone (cocoercive) operators $B: E \rightarrow E^*$ method (2) weakly converges [17]. Recently, Shehu [18] extended Tseng’s result to 2-uniformly convex and uniformly smooth Banach spaces. He proposed the following weakly converging process for approximating the solution of inclusion (1):

$$\begin{cases} y_n = J_{\lambda}^A \circ J^{-1}(x_n - \lambda_n Bx_n), \\ x_{n+1} = J^{-1}(Jy_n - \lambda_n(By_n - Bx_n)), \end{cases} \quad (3)$$

where $\lambda_n > 0$ can be set using knowledge of Lipschitz constant of the operator B or calculated with linear search. Note that two values of operator B need to be calculated at the iteration step.

In this article we study a new splitting algorithm for solving operator inclusion (1) in Banach space. The algorithm is an adaptive modification of the well-known Malitsky–Tam “forward-reflected-backward algorithm” [14], where the step size update rule not require Lipschitz continuous constant

knowledge for operator B . It's advantage is a single computation of maximal monotone operator A resolvent and B operator value on each iterative step. The method weak convergence theorem is proved for operator inclusions in 2-uniformly convex and uniformly smooth Banach space.

The article structure is the following. Section 2 provides the necessary information from the areas of geometry of Banach spaces and theory of monotonous operators. The proposed adaptive operator extrapolation algorithm is described in section 3. Proof of weak convergence is presented in section 4. Theoretical applications to nonlinear operator equations, convex minimization problems and variational inequalities are given in sections 5 and 6. Some results of numerical experiments are also presented in section 7.

2. Preliminaries

To formulate and prove our results about convergence of algorithms we need some concepts and important facts about the geometry of real Banach space [26–33].

Consider real Banach space E with norm $\|\cdot\|$. Space E^* is the dual space for E . Let $\langle x^*, x \rangle$ is a value of linear bounded mapping $x^* \in E^*$ on $x \in E$ (i.e. $\langle x^*, x \rangle = x^*(x)$). Let's $\|\cdot\|_*$ be dual norm in dual space E^* .

Let $S_E = \{x \in E : \|x\| = 1\}$. Space E is strictly convex, if $\forall x, y \in S_E, x \neq y$ we have $\left\| \frac{x+y}{2} \right\| < 1$ [26]. The modulus of convexity of Banach space E is defined as [27]

$$\delta_E(\varepsilon) = \inf \left\{ 1 - \left\| \frac{x+y}{2} \right\| : x, y \in S_E, \|x-y\| = \varepsilon \right\} \quad \forall \varepsilon \in (0, 2].$$

Space E is uniformly convex, if $\delta_E(\varepsilon) > 0 \quad \forall \varepsilon \in (0, 2]$ [27]. Space E is 2-uniformly convex, if there exists such $c > 0$ that $\delta_E(\varepsilon) \geq c\varepsilon^2 \quad \forall \varepsilon \in (0, 2]$ [27]. It is known that 2-uniform convexity implies uniform convexity. Also the uniform convexity of real Banach space implies its reflexivity [26].

A space E is smooth if the limit

$$\lim_{t \rightarrow 0} \frac{\|x+ty\| - \|x\|}{t} \quad (4)$$

exists for all $x, y \in S_E$ [26]. A space E is uniformly smooth if the limit (4) exists uniformly over $x, y \in S_E$ [26].

Convexity type and smoothness type of spaces E and E^* are in duality relationship [26, 27]: dual space E^* is strictly convex $\Rightarrow$ space E is smooth [26]; dual space E^* is smooth $\Rightarrow$ space E is strictly convex [26]; space E is uniformly convex $\Rightarrow$ dual space E^* is uniformly smooth [26]; space E is uniformly smooth $\Rightarrow$ dual space E^* is uniformly convex [27]. We can reverse the first two implications if the space E is reflexive [26].

Widely used functional spaces L_p ($1 < p \leq 2$) and real Hilbert spaces are uniformly smooth (spaces L_p are uniformly smooth for $p \in (1, \infty)$) and 2-uniformly convex [26–29].

Also recall [1, 28, 31, 32] that a multivalued operator $A: E \rightarrow 2^{E^*}$ is called monotone if $\forall x, y \in E$

$$\langle u-v, x-y \rangle \geq 0 \quad \forall u \in Ax, v \in Ay.$$

A monotone operator $A: E \rightarrow 2^{E^*}$ is called maximal monotone if for any monotone operator $B: E \rightarrow 2^{E^*}$ we have that $\Gamma(A) \subseteq \Gamma(B)$ implies $\Gamma(A) = \Gamma(B)$, where

$$\Gamma(A) = \{(x, u) \in E \times E^* : u \in Ax\}$$

is a graph of A [1, 28, 31].

Lemma 1 ([1, 28]). Let $A: E \rightarrow 2^{E^*}$ be a maximal monotone operator, $x \in E, u \in E^*$. Then

$$\langle u - v, x - y \rangle \geq 0 \quad \forall (y, v) \in \Gamma(A) \quad \Leftrightarrow \quad (x, u) \in \Gamma(A).$$

It is known that if $A: E \rightarrow 2^{E^*}$ is maximal monotone operator, $B: E \rightarrow E^*$ is Lipschitz continuous monotone operator, then $A + B$ is maximal monotone operator [28, 31].

Let us also recall [1] that operator $A: E \rightarrow E^*$ is called inverse strongly monotone (cocoercive) if there exists such a number $\alpha > 0$ (the constant of inverse strong monotonicity) that

$$\langle Ax - Ay, x - y \rangle \geq \alpha \|Ax - Ay\|^2.$$

Inverse strongly monotone operator is Lipschitz continuous, but not every Lipschitz continuous operator is inverse strongly monotone.

Multivalued mapping $J: E \rightarrow 2^{E^*}$, which acts as

$$Jx = \left\{ x^* \in E^* : \langle x^*, x \rangle = \|x\|^2 = \|x^*\|_*^2 \right\},$$

is called normalized duality mapping [30]. We also use the next facts [28, 30, 32, 33]: if Banach space E is smooth then operator J is single-valued [30]; if real Banach space E is strongly convex then operator J is one-to-one and strongly monotone [28]; if Banach space E is reflexive then operator J is onto [28]; if Banach space E is uniformly smooth then on all bounded subsets of E operator J is uniformly continuous [30].

Remark 1. For a Hilbert space $J = I$. Explicit form of operator J for Banach spaces ℓ_p , L_p , and W_p^m ($p \in (1, \infty)$) is provided in [28, 32].

Consider reflexive, strictly convex and smooth space E [26]. The maximal monotonicity of operator $A: E \rightarrow 2^{E^*}$ is equivalent to equality

$$R(J + \lambda A) = E^*$$

for all $\lambda > 0$ [31]. For maximal monotone operator $A: E \rightarrow 2^{E^*}$ and $\lambda > 0$ resolvent $J_\lambda^A: E \rightarrow E$ is defined as follows [31]

$$J_\lambda^A x = (J + \lambda A)^{-1} Jx, \quad x \in E,$$

where J is normalized duality mapping from E to E^* . It is known that

$$A^{-1}0 = F(J_\lambda^A) = \{x \in E : J_\lambda^A x = x\} \quad \forall \lambda > 0.$$

It is also known that the set $A^{-1}0$ is closed and convex [1, 31].

Consider smooth Banach space E [26]. Yakov Alber introduced the convenient real-valued functional [32]

$$D(x, y) = \|x\|^2 - 2\langle Jy, x \rangle + \|y\|^2 \quad \forall x, y \in E.$$

Definition of D implies a useful 3-point identity:

$$D(x, y) - D(x, z) - D(z, y) = 2\langle Jz - Jy, x - z \rangle \quad \forall x, y, z \in E.$$

For strictly convex Banach space E and $x, y \in E$ [32]:

$$D(x, y) = 0 \Leftrightarrow x = y.$$

Lemma 2 ([31]). Let E be a uniformly convex and uniformly smooth Banach space $(x_n), (y_n)$ – bounded sequences of elements from Banach space E . Then

$$\|x_n - y_n\| \rightarrow 0 \Leftrightarrow \|Jx_n - Jy_n\|_* \rightarrow 0 \Leftrightarrow D(x_n, y_n) \rightarrow 0.$$

Lemma 3 ([24, 25]). Let E be a smooth Banach space and 2-uniformly convex. Then for some $\mu \geq 1$ we have:

$$D(x, y) \geq \frac{1}{\mu} \|x - y\|^2 \quad \forall x, y \in E.$$

Remark 2. For Banach spaces ℓ_p , L_p and W_p^m ($1 < p \leq 2$) we have $\mu = \frac{1}{p-1}$ [29]. And for a Hilbert space inequality for Lemma 3 becomes identity.

3. Algorithm

Let real Banach space E be a uniformly smooth and 2-uniformly convex [27]. Let A be a multivalued operator acting from E into 2^{E^*} , and B an operator acting from E into E^* .

Consider the next operator inclusion problem:

$$\text{find } x \in E: \quad 0 \in (A+B)x, \quad (5)$$

Let us denote $(A+B)^{-1}0$ set of solutions of this operator inclusion.

Suppose that the next assumptions hold [20]:

- $A: E \rightarrow 2^{E^*}$ is a maximal monotone operator;
- $B: E \rightarrow E^*$ is a monotone and Lipschitz continuous operator with Lipschitz constant $L > 0$;
- Set $(A+B)^{-1}0$ is nonempty.

Operator inclusion (5) can be formulated as the problem of finding a fixed point:

$$\text{find } x \in E: \quad x = J_{\lambda}^A \circ J^{-1}(Jx - \lambda Bx), \quad (6)$$

where $\lambda > 0$. Formulation (6) is useful, as it leads to well-known simple algorithmic idea. Calculation scheme

$$x_{n+1} = J_{\lambda}^A \circ J^{-1}(Jx_n - \lambda Bx_n)$$

was studied in [17] for inverse strongly monotone operators $B: E \rightarrow E^*$. However, the scheme generally does not converge for Lipschitz continuous monotone operators. Let's use the idea of work [14] and consider modified scheme

$$x_{n+1} = J_{\lambda}^A \circ J^{-1}(Jx_n - \lambda Bx_n - \lambda(Bx_n - Bx_{n-1}))$$

with extrapolation term

$$-\lambda(Bx_n - Bx_{n-1}).$$

And let's use update rule for $\lambda > 0$ similar to one from [16] to exclude explicit use of Lipschitz constant of operator B .

We will assume that we know constant $\mu \geq 1$ from Lemma 3.

Algorithm 1.

Choose some $x_0 \in E$, $x_1 \in E$, $\tau \in (0, \frac{1}{2\mu})$ and $\lambda_0, \lambda_1 > 0$. Set $n \leftarrow 1$.

1. Compute

$$x_{n+1} = J_{\lambda_n}^A \circ J^{-1}(Jx_n - \lambda_n Bx_n - \lambda_{n-1}(Bx_n - Bx_{n-1})).$$

2. If $x_{n-1} = x_n = x_{n+1}$, then $x_n \in (A+B)^{-1}0$, else return to 3.

3. Compute

$$\lambda_{n+1} = \begin{cases} \min \left\{ \lambda_n, \tau \frac{\|x_{n+1} - x_n\|}{\|Bx_{n+1} - Bx_n\|_*} \right\}, & \text{if } Bx_{n+1} \neq Bx_n, \\ \lambda_n, & \text{otherwise.} \end{cases}$$

Set $n \leftarrow n+1$ and return to 1.

Step size sequence (λ_n) which is created by rule on step 3 is non-increasing and bounded from below by

$$\min\{\lambda_1, \tau L^{-1}\}.$$

So, we have $\lim_{n \rightarrow \infty} \lambda_n > 0$.

Let us prove convergence result for proposed Algorithm 1.

4. Proof of convergence

The following lemma contains inequality which is crucial to proof weak convergence of adaptive operator extrapolation method (Algorithm 1).

Lemma 4. The next inequality holds for the sequence (x_n) , generated by adaptive operator extrapolation method (Algorithm 1):

$$\begin{aligned} D(z, x_{n+1}) + 2\lambda_n \langle Bx_n - Bx_{n+1}, x_{n+1} - z \rangle + \tau\mu \frac{\lambda_n}{\lambda_{n+1}} D(x_{n+1}, x_n) &\leq \\ &\leq D(z, x_n) + 2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) - \\ &\quad - \left(1 - \tau\mu \frac{\lambda_{n-1}}{\lambda_n} - \tau\mu \frac{\lambda_n}{\lambda_{n+1}}\right) D(x_{n+1}, x_n), \end{aligned}$$

where $z \in (A+B)^{-1}0$.

Proof. Let $z \in (A+B)^{-1}0$. Then

$$-Bz \in Az.$$

Equality $x_{n+1} = J_{\lambda_n}^A \circ J^{-1}(Jx_n - \lambda_n Bx_n - \lambda_{n-1}(Bx_n - Bx_{n-1}))$ can be formulated as

$$Jx_n - \lambda_n Bx_n - \lambda_{n-1}(Bx_n - Bx_{n-1}) - Jx_{n+1} \in \lambda_n Ax_{n+1}.$$

From monotonicity of A

$$\langle Jx_{n+1} + \lambda_n(Bx_n - Bz) + \lambda_{n-1}(Bx_n - Bx_{n-1}) - Jx_n, z - x_{n+1} \rangle \geq 0. \quad (7)$$

Let's write 3-point identity

$$D(z, x_n) - D(z, x_{n+1}) - D(x_{n+1}, x_n) = 2\langle Jx_{n+1} - Jx_n, z - x_{n+1} \rangle. \quad (8)$$

Using (8) withing (7) we get

$$D(z, x_{n+1}) \leq D(z, x_n) - D(x_{n+1}, x_n) + 2\langle \lambda_n(Bx_n - Bz) + \lambda_{n-1}(Bx_n - Bx_{n-1}), z - x_{n+1} \rangle. \quad (9)$$

Using monotonicity of operator B . We have

$$\begin{aligned} \langle \lambda_n(Bx_n - Bz) + \lambda_{n-1}(Bx_n - Bx_{n-1}), z - x_{n+1} \rangle &= \lambda_n \langle Bx_n - Bx_{n+1}, z - x_{n+1} \rangle + \\ &+ \lambda_{n-1} \langle Bx_n - Bx_{n-1}, z - x_{n+1} \rangle + \underbrace{\lambda_n \langle Bx_{n+1} - Bz, z - x_{n+1} \rangle}_{\leq 0} \leq \\ &\leq \lambda_n \langle Bx_n - Bx_{n+1}, z - x_{n+1} \rangle + \lambda_{n-1} \langle Bx_n - Bx_{n-1}, z - x_n \rangle + \\ &\quad + \lambda_{n-1} \langle Bx_n - Bx_{n-1}, x_n - x_{n+1} \rangle. \end{aligned} \quad (10)$$

Using (10) withing (9) we get

$$\begin{aligned} D(z, x_{n+1}) &\leq D(z, x_n) - D(x_{n+1}, x_n) + 2\lambda_n \langle Bx_n - Bx_{n+1}, z - x_{n+1} \rangle + \\ &\quad + 2\lambda_{n-1} \langle Bx_n - Bx_{n-1}, z - x_n \rangle + 2\lambda_{n-1} \langle Bx_n - Bx_{n-1}, x_n - x_{n+1} \rangle. \end{aligned} \quad (11)$$

Using the rule for λ_n recalculation, we can estimate term $2\lambda_{n-1} \langle Bx_n - Bx_{n-1}, x_n - x_{n+1} \rangle$ in (11) from above. We get

$$\begin{aligned}
2\lambda_{n-1} \langle Bx_n - Bx_{n-1}, x_n - x_{n+1} \rangle &\leq 2\lambda_{n-1} \|Bx_n - Bx_{n-1}\|_* \|x_n - x_{n+1}\| \leq 2\tau \frac{\lambda_{n-1}}{\lambda_n} \|x_n - x_{n-1}\| \|x_{n+1} - x_n\| \leq \\
&\leq \tau \frac{\lambda_{n-1}}{\lambda_n} \|x_n - x_{n-1}\|^2 + \tau \frac{\lambda_{n-1}}{\lambda_n} \|x_n - x_{n+1}\|^2 \leq \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_{n+1}, x_n).
\end{aligned}$$

So, we came to inequality

$$\begin{aligned}
D(z, x_{n+1}) + 2\lambda_n \langle Bx_n - Bx_{n+1}, x_{n+1} - z \rangle + \tau\mu \frac{\lambda_n}{\lambda_{n+1}} D(x_{n+1}, x_n) &\leq \\
\leq D(z, x_n) + 2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) - \left(1 - \tau\mu \frac{\lambda_{n-1}}{\lambda_n} - \tau\mu \frac{\lambda_n}{\lambda_{n+1}}\right) D(x_{n+1}, x_n),
\end{aligned}$$

which was required to prove.

The next theorem states our main result.

Theorem 1. Let Banach space E be a uniformly smooth and 2-uniformly convex, $A: E \rightarrow 2^{E^*}$ be a multivalued maximal monotone operator, $B: E \rightarrow E^*$ be a monotone and Lipschitz continuous operator. Suppose that J (normalized duality map) is sequentially weakly continuous and $(A+B)^{-1}0 \neq \emptyset$. Then we have weak convergence of sequence (x_n) generated by adaptive operator extrapolation method (Algorithm 1) to a point $z \in (A+B)^{-1}0$.

Proof. Let $z' \in (A+B)^{-1}0$. Denote

$$\begin{aligned}
a_n &= D(z', x_n) + 2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z' \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}), \\
b_n &= \left(1 - \tau\mu \frac{\lambda_{n-1}}{\lambda_n} - \tau\mu \frac{\lambda_n}{\lambda_{n+1}}\right) D(x_{n+1}, x_n)
\end{aligned}$$

Inequality from Lemma 4 becomes

$$a_{n+1} \leq a_n - b_n.$$

As $\lim_{n \rightarrow \infty} \lambda_n > 0$ exists, we have

$$1 - \tau\mu \frac{\lambda_{n-1}}{\lambda_n} - \tau\mu \frac{\lambda_n}{\lambda_{n+1}} \rightarrow 1 - 2\tau\mu \in (0, 1) \text{ when } n \rightarrow \infty.$$

Let's show, that $a_n \geq 0$ for all big enough $n \in N$. We have

$$\begin{aligned}
a_n &= D(z', x_n) + 2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z' \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) \geq \\
&\geq \frac{1}{\mu} \|x_n - z'\|^2 - 2\lambda_{n-1} \|Bx_{n-1} - Bx_n\|_* \|x_n - z'\| + \tau \frac{\lambda_{n-1}}{\lambda_n} \|x_{n-1} - x_n\|^2 \geq \\
&\geq \frac{1}{\mu} \|x_n - z'\|^2 - 2\tau \frac{\lambda_{n-1}}{\lambda_n} \|x_n - x_{n-1}\| \|x_n - z'\| + \tau \frac{\lambda_{n-1}}{\lambda_n} \|x_{n-1} - x_n\|^2 \geq \left(\frac{1}{\mu} - \tau \frac{\lambda_{n-1}}{\lambda_n}\right) \|x_n - z'\|^2.
\end{aligned}$$

We can find such $n_0 \in N$ that

$$\frac{1}{\mu} - \tau \frac{\lambda_{n-1}}{\lambda_n} > 0 \text{ for all } n \geq n_0,$$

so $a_n \geq 0$ for all $n \geq n_0$.

Now we can conclude that the next limit exists

$$\lim_{n \rightarrow \infty} \left(D(z', x_n) + 2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z' \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) \right),$$

and

$$\sum_{n=1}^{\infty} \left(1 - \tau\mu \frac{\lambda_{n-1}}{\lambda_n} - \tau\mu \frac{\lambda_n}{\lambda_{n+1}} \right) D(x_{n+1}, x_n) < +\infty.$$

So the generated sequence (x_n) is bounded. Also we have

$$\lim_{n \rightarrow \infty} \phi(x_{n+1}, x_n) = \lim_{n \rightarrow \infty} \|x_{n+1} - x_n\| = 0.$$

From the fact

$$\lim_{n \rightarrow \infty} \left(2\lambda_{n-1} \langle Bx_{n-1} - Bx_n, x_n - z' \rangle + \tau\mu \frac{\lambda_{n-1}}{\lambda_n} D(x_n, x_{n-1}) \right) = 0,$$

we get convergence of sequences $(\phi(z', x_n))$ for all $z' \in (A+B)^{-1}0$.

Let us show that all weak partial limits of (x_n) sequence belong to set $(A+B)^{-1}0$. Consider a subsequence (x_{n_k}) that weakly converges to some point $z \in E$. Let's show that

$$z \in (A+B)^{-1}0.$$

If we take any point $(y, v) \in \Gamma(A+B)$ then $v - By \in Ay$. We have

$$Jx_{n_k} - \lambda_{n_k} Bx_{n_k} - \lambda_{n_k-1} (Bx_{n_k} - Bx_{n_k-1}) - Jx_{n_k+1} \in \lambda_{n_k} Ax_{n_k+1}.$$

Using monotonicity of A operator, we get

$$\langle \lambda_{n_k} v + Jx_{n_k+1} + \lambda_{n_k} (Bx_{n_k} - By) + \lambda_{n_k-1} (Bx_{n_k} - Bx_{n_k-1}) - Jx_{n_k}, y - x_{n_k+1} \rangle \geq 0.$$

And then, using monotonicity of B operator, we can obtain the following estimation

$$\begin{aligned} & \langle v, y - x_{n_k} \rangle + \langle v, x_{n_k} - x_{n_k+1} \rangle = \\ & = \langle v, y - x_{n_k+1} \rangle \geq \left\langle By + \frac{1}{\lambda_{n_k}} (Jx_{n_k} - Jx_{n_k+1} - \lambda_{n_k} Bx_{n_k} - \lambda_{n_k-1} (Bx_{n_k} - Bx_{n_k-1})), y - x_{n_k+1} \right\rangle = \\ & = \frac{1}{\lambda_{n_k}} \langle Jx_{n_k} - Jx_{n_k+1}, y - x_{n_k+1} \rangle - \frac{\lambda_{n_k-1}}{\lambda_{n_k}} \langle Bx_{n_k} - Bx_{n_k-1}, y - x_{n_k+1} \rangle + \\ & \quad + \langle By - Bx_{n_k+1}, y - x_{n_k+1} \rangle + \langle Bx_{n_k+1} - Bx_{n_k}, y - x_{n_k+1} \rangle \geq \\ & \geq \frac{1}{\lambda_{n_k}} \langle Jx_{n_k} - Jx_{n_k+1}, y - x_{n_k+1} \rangle - \frac{\lambda_{n_k-1}}{\lambda_{n_k}} \langle Bx_{n_k} - Bx_{n_k-1}, y - x_{n_k+1} \rangle + \langle Bx_{n_k+1} - Bx_{n_k}, y - x_{n_k+1} \rangle. \end{aligned}$$

From

$$\lim_{n \rightarrow \infty} \|x_n - x_{n-1}\| = 0,$$

and the Lipschitz continuity of B we have

$$\lim_{n \rightarrow \infty} \|Bx_n - Bx_{n-1}\|_* = 0.$$

Due to the uniform continuity on bounded sets of the J [30], we obtain

$$\lim_{n \rightarrow \infty} \|Jx_n - Jx_{n+1}\|_* = 0.$$

Thus

$$\langle v, y - z \rangle = \lim_{k \rightarrow \infty} \langle v, y - x_{n_k} \rangle \geq 0.$$

The maximal monotonicity of operator $A+B$ and arbitrariness of $(y, v) \in \Gamma(A+B)$ imply $z \in (A+B)^{-1}0$ (Lemma 1).

We need to prove that the sequence (x_n) weakly converges to point z . Let's argue by contradiction. Let there exist a subsequence (x_{m_k}) that weakly converges to z' , $z \neq z'$. Obviously that $z' \in (A+B)^{-1}0$. We have the equality

$$2\langle Jx_n, z - z' \rangle = D(z', x_n) - D(z, x_n) + \|z\|^2 - \|z'\|^2.$$

So the next limit exists

$$\lim_{n \rightarrow \infty} \langle Jx_n, z - z' \rangle.$$

Due to the sequential weak continuity of the normalized duality mapping J , we obtain

$$\langle Jz, z - z' \rangle = \lim_{k \rightarrow \infty} \langle Jx_{n_k}, z - z' \rangle = \lim_{k \rightarrow \infty} \langle Jx_{m_k}, z - z' \rangle = \langle Jz', z - z' \rangle,$$

so

$$\langle Jz - Jz', z - z' \rangle = 0.$$

As a result we get the contradiction $z = z'$.

5. Algorithm variants

For completeness, let us formulate a modification of the proposed algorithm with a fixed step size parameter $\lambda > 0$.

Algorithm 2.

Select some points $x_0 \in E$, $x_1 \in E$, and $\lambda \in \left(0, 0.5 \frac{1}{\mu} L^{-1}\right)$. Set $n \leftarrow 1$.

1. Compute

$$x_{n+1} = J_\lambda^A \circ J^{-1} (Jx_n - 2\lambda Bx_n + \lambda Bx_{n-1}).$$

2. If $x_{n-1} = x_n = x_{n+1}$, then $x_n \in (A + B)^{-1} 0$. Else set $n \leftarrow n + 1$ and return to 1.

Consider the problem of finding the zero of a nonlinear monotone Lipschitz continuous operator $B: E \rightarrow E^*$:

$$\text{find } x \in E: Bx = 0.$$

For such case Algorithm 1 becomes

Algorithm 3.

Choose some $x_0 \in E$, $x_1 \in E$, $\tau \in (0, (2\mu)^{-1})$ and $\lambda_0, \lambda_1 > 0$. Set $n \leftarrow 1$.

1. Compute

$$x_{n+1} = J^{-1} (Jx_n - \lambda_n Bx_n - \lambda_{n-1} (Bx_n - Bx_{n-1})).$$

2. If $x_{n-1} = x_n = x_{n+1}$, then $x_n \in B^{-1} 0$. Else go to 3.

3. Compute

$$\lambda_{n+1} = \begin{cases} \min \left\{ \lambda_n, \tau \frac{\|x_{n+1} - x_n\|}{\|Bx_{n+1} - Bx_n\|_*} \right\}, & \text{if } Bx_{n+1} \neq Bx_n, \\ \lambda_n, & \text{otherwise.} \end{cases}$$

Set $n \leftarrow n + 1$ and return to 1.

Weak convergence of Algorithm 3 follows from Theorem 1.

Theorem 2. Let Banach space E be uniformly smooth and 2-uniformly convex, $B: E \rightarrow E^*$ – monotone and Lipschitz continuous operator, $B^{-1} 0 \neq \emptyset$. If J (normalized duality mapping) is sequentially weakly continuous, then the sequence (x_n) converges weakly to some point $z \in B^{-1} 0$.

Remark 3. In [17], the following algorithm was proposed to find the zero of a reverse strongly monotone operator $B: E \rightarrow E^*$:

$$x_{n+1} = J^{-1} (Jx_n - \lambda_n Bx_n), \quad x_1 \in E,$$

where $\lambda_n \in [a, b] \subseteq \left(0, \frac{\alpha}{\mu}\right)$, $\alpha > 0$ – the constant of inverse strong monotonicity of operator B . This algorithm generally does not converge for Lipschitz continuous monotone operators, but it converges weakly ergodic.

Using Theorem 2, let us consider the problem of minimizing a convex continuously Frechet differentiable functional

$$f(x) \rightarrow \min_{x \in E}. \quad (12)$$

We can formulate variant of Algorithm 3 for (12), which is a smooth convex minimization problem.

Algorithm 4.

Choose some $x_0 \in E$, $x_1 \in E$, $\tau \in \left(0, \frac{1}{2\mu}\right)$ and $\lambda_0, \lambda_1 > 0$. Set $n \leftarrow 1$.

1. Compute

$$x_{n+1} = J^{-1} \left(Jx_n - \lambda_n \nabla f(x_n) - \lambda_{n-1} (\nabla f(x_n) - \nabla f(x_{n-1})) \right).$$

2. If $x_{n-1} = x_n = x_{n+1}$, then $x_n \in \arg \min f$. Else return to 3.

3. Compute

$$\lambda_{n+1} = \begin{cases} \min \left\{ \lambda_n, \tau \frac{\|x_{n+1} - x_n\|}{\|\nabla f(x_{n+1}) - \nabla f(x_n)\|_*} \right\}, & \text{if } \nabla f(x_{n+1}) \neq \nabla f(x_n), \\ \lambda_n, & \text{otherwise.} \end{cases}$$

Set $n \leftarrow n + 1$ and return to 1.

For this variant of the algorithm, from Theorem 2 we get

Theorem 3. Let Banach space E be uniformly smooth and 2-uniformly convex, $f : E \rightarrow R$ – convex continuously Frechet differentiable functional with Lipschitz continuous derivative, and $\arg \min f$ is non-empty. Assume that J is sequentially weakly continuous. Then sequence (x_n) generated by Algorithm 4 converges weakly to a point $z \in \arg \min f$.

6. Application to variational inequalities

Consider real Hilbert space H . Let C is a non-empty, convex and closed subset of space H , $B : H \rightarrow H$ is a monotone and Lipschitz continuous operator. Consider the next variational inequality problem:

$$\text{find } x \in C \text{ such that } (Bx, y - x) \geq 0 \quad \forall y \in C, \quad (13)$$

Let $VI(B, C)$ be a solution set of problem (13). Variational inequality (13) is equivalent to the operator inclusion [1]

$$\text{find } x \in H \text{ such that } 0 \in (A + B)x,$$

where $A = N_C$ is a normal cone for convex and closed set C , i.e.

$$N_C x = \begin{cases} \{w \in H : (w, y - x) \leq 0 \quad \forall y \in C\}, & x \in C, \\ \emptyset, & \text{otherwise.} \end{cases}$$

It is known that

$$J_\lambda^A = (I + \lambda A)^{-1} = (I + \lambda N_C)^{-1} = P_C,$$

where P_C is projection operator onto the set C [1].

For variational inequality problem (13), adaptive operator extrapolation method (Algorithm 1) takes the following form:

Algorithm 5.

Choose some $x_1 \in H$, $x_2 \in H$, $\tau \in (0, 0.5)$ and $\mu_1, \mu_2 > 0$. Set $n \leftarrow 2$.

1. Compute

$$x_{n+1} = P_C(x_n - \mu_n Bx_n - \mu_{n-1}(Bx_n - Bx_{n-1})).$$

2. If $x_{n-1} = x_n = x_{n+1}$, then STOP and $x_n \in VI(B, C)$. Else go to 3.

3. Compute

$$\mu_{n+1} = \begin{cases} \min \left\{ \mu_n, \tau \frac{\|x_{n+1} - x_n\|}{\|Bx_{n+1} - Bx_n\|} \right\}, & \text{if } Bx_{n+1} \neq Bx_n, \\ \mu_n, & \text{otherwise.} \end{cases}$$

Set $n \leftarrow n + 1$ and return to 1.

For variational inequality case, from Theorem 1 we get

Theorem 4. Let H be a real Hilbert space, C is a non-empty, convex and closed subset of H , $B: H \rightarrow H$ is a monotone Lipschitz continuous operator, $VI(B, C) \neq \emptyset$. Then the sequence (x_n) generated by Algorithm 5 converges weakly to a point $z \in VI(B, C)$.

7. Numerical experiments

The numerical experiments are performed in Python 3.8.5 with NumPy 1.19 on a 64-bit PC with an Intel Core i7-1065G7 1.3 – 3.9GHz and 16GB RAM. As the test example we consider a toy variational inequality with pseudo-monotone operator.

Example. Let

$$C = \{x \in [-5, 5]^3 : x_1 + x_2 + x_3 = 0\} \subseteq R^3,$$

and operator $B: R^3 \rightarrow R^3$ be defined as

$$Bx = \left(e^{-\|x\|^2} + 0.2 \right) \begin{pmatrix} 2 & 0 & -2 \\ 0 & 3 & 0 \\ -2 & 0 & 4 \end{pmatrix} x$$

The variational inequality problem of finding $x \in C : (Bx, y - x) \geq 0 \ \forall y \in C$ has unique solution $x^* = (0, 0, 0)$. Also, Lipschitz constant for the operator B is known – $L = 10.136$. We compare adaptive and non-adaptive (marked as “lip” on figures) variants of “Extrapolation from the Past” algorithm [9] and Algorithm 5. For stopping criteria and error estimation we use Euclidian distance to known solution $D_n = \|x_n - x^*\|$. For this example, we use $x_0 = (-4, 3, 5)$ as starting point, $\lambda = 0.9(\sqrt{2} - 1)/L$ for non-adaptive version of “Extrapolation from the Past” and $\lambda = 0.9/2L$ for non-adaptive version of Algorithm 5. Also, we use $\tau = 0.3$ and $\tau = 0.45$ (nearly maximal feasible values) for adaptive versions of “Extrapolation from the Past” and Algorithm 5 accordingly.

Time measurements below are got by averaging 100 runs for each algorithm.

Table 1

Time to reach desired error rate D_n , seconds

error\algorithm	Extra Past (lip)	Extra Past	Alg. 5 (lip)	Alg. 5
$\varepsilon = 1 \cdot 10^{-10}$	0.0308	0.0174	0.0181	0.0087
$\varepsilon = 1 \cdot 10^{-13}$	0.0419	0.0251	0.0245	0.0129
$\varepsilon = 1 \cdot 10^{-16}$	0.0555	0.0352	0.0331	0.0182

As we see from Table 1, for this problem Algorithm 5 outperforms other variants.

On the figure below we can see convergence behavior.

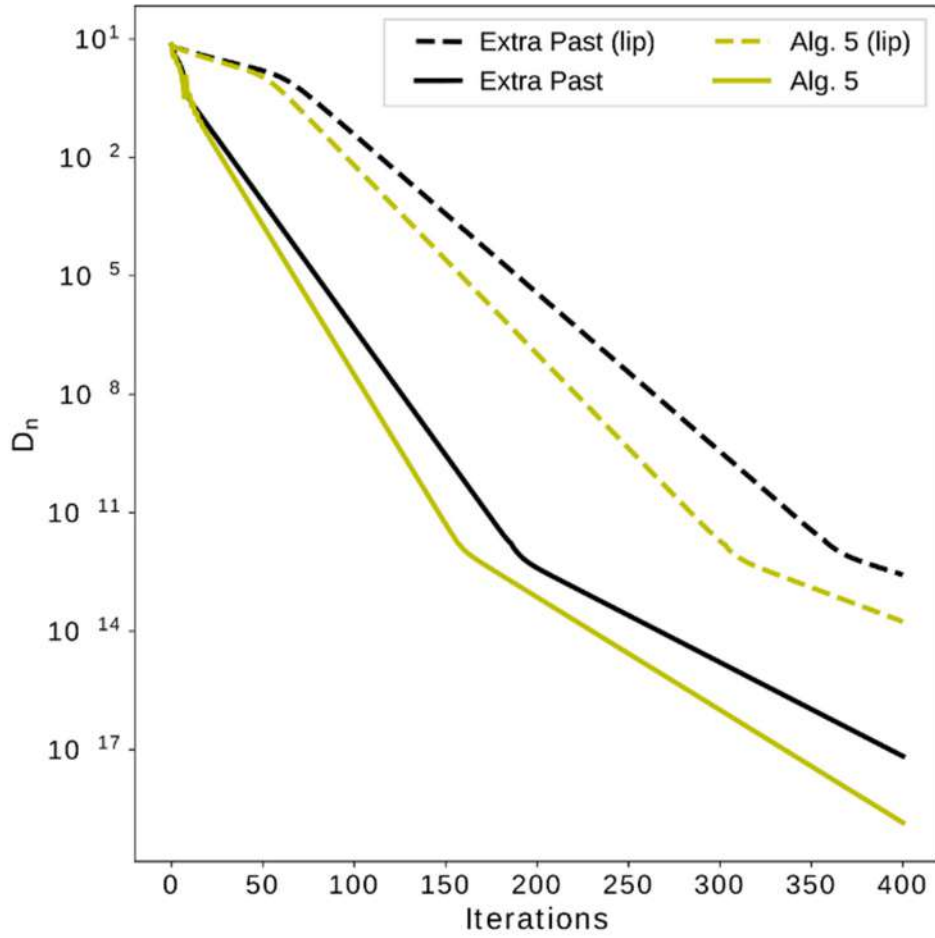


Figure 1: Convergence in terms of iterations

As it can be seen, both adaptive algorithms behave very closely. But it should be noted that Algorithm 5 has only one projection on each iteration instead of two for “Extrapolation from the Past” algorithm [9]

$$\begin{cases} y_n = P_C(x_n - \mu_n B y_{n-1}), \\ x_{n+1} = P_C(x_n - \mu_n B y_n). \end{cases}$$

8. Conclusions

In this paper new iterative splitting algorithm for solving an operator inclusion with the sum of a maximal monotone operator and a monotone Lipschitz continuous operator in a real Banach space is investigated.

Algorithm 1 is an improvement of the well-known “forward-reflected-backward algorithm” of Malitsky–Tam [14] with adaptive step size, where the step update rule does not require a priori knowledge of the Lipschitz constant of operator B [16].

The algorithm advantage is a single computation of the resolvent of the maximal monotone operator A and the value of the monotone Lipschitz continuous operator B at the iteration step.

Method weak convergence theorem is proved for operator inclusions in Banach space with 2-uniform convexity and uniform smoothness [27]. Theoretical applications to operator equations, convex minimization problems, and variational inequalities are presented.

An interesting challenge for the future is the development of a strongly convergent modification for Algorithm 1.

In connection with the study we will point out two topical issues. First, all results are obtained for the class 2-uniformly convex and uniformly smooth real Banach spaces [27], which does not contain important for applications spaces L_p and W_p^m ($2 < p < +\infty$). It is highly desirable to get rid of this limitation. Second, fast and robust algorithms for computing the resolvent for a wide range of maximal monotone operators are needed to effectively apply algorithms for nonlinear problems in Banach spaces.

An interesting question is the study of the behavior of Algorithm 1 in the situation $A = I$. Namely, the question of asymptotic behavior of $\|Bx_n\|_*$. Note that an estimate $\|Bx_n\|_* = O\left(\frac{1}{\sqrt{n}}\right)$ is theoretically possible.

Acknowledgements

This work was supported by the MES of Ukraine (project "Computational algorithms and optimization for artificial intelligence, medicine and defense", 0122U002026).

References

- [1] H. H. Bauschke, P. L. Combettes, *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*. Berlin; Heidelberg; New York, Springer, 2017.
- [2] Y. Alber, I. Ryazantseva, *Nonlinear Ill Posed Problems of Monotone Type*, Springer, Dordrecht, 2006.
- [3] H. Raguet, J. Fadili, G. Peyre, A generalized forward-backward splitting, *SIAM J. Imaging Sci.* 6 (2013) 1199–1226
- [4] D. R. Luke, C. Charitha, R. Shefi, Y. Malitsky, Efficient, Quantitative Numerical Methods for Statistical Image Deconvolution and Denoising. In: Salditt T., Egner A., Luke D.R. (eds.) *Nanoscale Photonic Imaging. Topics in Applied Physics*, vol 134. Springer, Cham, 2020, pp. 313–338.
- [5] R. Tibshirani, Regression shrinkage and selection via the lasso, *J. Royal Statist. Soc.* 58 (1996) 267–288.
- [6] S. I. Lyashko, D. A. Klyushin, D. A. Nomirowsky, V. V. Semenov, Identification of age-structured contamination sources in ground water, in: R. Boucekkine, N. Hritonenko, Y. Yatsenko, Y. (Eds.), *Optimal Control of Age-Structured Populations in Economy, Demography, and the Environment*, Routledge, London–New York, 2013, pp. 277–292.
- [7] V. V. Semenov, A Strongly Convergent Splitting Method for Systems of Operator Inclusions with Monotone Operators, *Journal of Automation and Information Sciences* 46(5) (2014) 45–56. <https://doi.org/10.1615/JAutomatInfScien.v46.i5.40>.
- [8] D. A. Nomirowskii, B. V. Rublyov, V. V. Semenov, Convergence of Two-Stage Method with Bregman Divergence for Solving Variational Inequalities, *Cybernetics and Systems Analysis* 55 (2019) 359–368. doi:10.1007/s10559-019-00142-7.
- [9] L. Chabak, V. Semenov, Y. Vedel, A New Non-Euclidean Proximal Method for Equilibrium Problems, In: O. Chertov et al. (Eds.), *Recent Developments in Data Science and Intelligent Analysis of Information*, volume 836 of *Advances in Intelligent Systems and Computing*, Springer, Cham, 2019, pp. 50–58. doi:10.1007/978-3-319-97885-7_6.
- [10] V. V. Semenov, Modified Extragradient Method with Bregman Divergence for Variational Inequalities, *Journal of Automation and Information Sciences* 50(8) (2018) 26–37. doi:10.1615/JAutomatInfScien.v50.i8.30.
- [11] P. L. Lions, B. Mercier, Splitting algorithms for the sum of two nonlinear operators, *SIAM J. Numer. Anal.* 16 (1979) 964–979.
- [12] G. B. Passty, Ergodic Convergence to a Zero of the Sum of Monotone Operators in Hilbert Spaces, *Journal of Mathematical Analysis and Applications* 72 (1979) 383–390.
- [13] P. Tseng, A modified forward-backward splitting method for maximal monotone mappings, *SIAM Journal on Control and Optimization* 38 (2000) 431–446.

- [14] Y. Malitsky, M. K. Tam, A Forward-Backward Splitting Method for Monotone Inclusions Without Cocoercivity, *SIAM Journal on Optimization* 30 (2020) 1451–1472. doi:10.1137/18M1207260.
- [15] E. R. Csetnek, Y. Malitsky, M. K. Tam, Shadow Douglas-Rachford Splitting for Monotone Inclusions, *Appl Math Optim.* 80 (2019) 665–678. doi:10.1007/s00245-019-09597-8.
- [16] S. Denisov, V. Semenov, Convergence of Adaptive Forward-Reflected-Backward Algorithm for Solving Variational Inequalities, in: *CEUR Workshop Proceedings*, vol. 3106, 2021, pp. 116–127. URL: http://ceur-ws.org/Vol-3106/Paper_11.pdf.
- [17] H. Iiduka, W. Takahashi, Weak convergence of a projection algorithm for variational inequalities in a Banach space, *Journal of Mathematical Analysis and Applications* 339(1) (2008) 668–679. <https://doi.org/10.1016/j.jmaa.2007.07.019>.
- [18] Y. Shehu, Convergence Results of Forward-Backward Algorithms for Sum of Monotone Operators in Banach Spaces, *Results Math.* 74 (2019). doi:10.1007/s00025-019-1061-4.
- [19] Y. Shehu, Single projection algorithm for variational inequalities in Banach spaces with application to contact problem, *Acta Math. Sci.* 40 (2020) 1045–1063. doi:10.1007/s10473-020-0412-2.
- [20] J. Yang, P. Cholamjiak, P. Sunthayuth, Modified Tseng's splitting algorithms for the sum of two monotone operators in Banach spaces, *AIMS Mathematics* 6(5) (2021) 4873–4900. doi:10.3934/math.2021286.
- [21] P. Cholamjiak, Y. Shehu, Inertial forward-backward splitting method in Banach spaces with application to compressed sensing, *Appl. Math.* 64 (2019) 409–435.
- [22] H. A. Abass, C. Izuchukwu, O. T. Mewomo, Q. L. Dong, Strong convergence of an inertial forward-backward splitting method for accretive operators in real Banach space, *Fixed Point Theory* 21 (2020) 397–412.
- [23] S. Takahashi, W. Takahashi, Split common null point problem and shrinking projection method for generalized resolvents in two Banach spaces, *J. Nonlinear Convex Anal.* 17 (2016) 2171–2182.
- [24] K. Aoyama, F. Kohsaka, Strongly relatively nonexpansive sequences generated by firmly nonexpansive-like mappings, *Fixed Point Theory Appl.* 95 (2014). doi:10.1186/1687-1812-2014-95.
- [25] T. Bonesky, K. S. Kazimierski, P. Maass, F. Schopfer, T. Schuster, Minimization of Tikhonov Functionals in Banach Spaces, *Abstr. Appl. Anal.* (2008) Article ID 192679. <https://doi.org/10.1155/2008/192679>.
- [26] J. Diestel, *Geometry of Banach Spaces*, Springer-Verlag, Berlin–Heidelberg, 1975.
- [27] B. Beauzamy, *Introduction to Banach Spaces and Their Geometry*, North-Holland, Amsterdam, 1985.
- [28] I. Cioranescu, *Geometry of Banach Spaces, Duality Mappings and Nonlinear Problems*, Kluwer Academic, Dordrecht, 1990.
- [29] Z. B. Xu, G. F. Roach, Characteristic inequalities of uniformly convex and uniformly smooth Banach spaces, *Journal of Mathematical Analysis and Applications* 157(1) (1991) 189–210.
- [30] M. M. Vainberg, *Variational Method and Method of Monotone Operators in the Theory of Nonlinear Equations*, Wiley, New York, 1974.
- [31] V. Barbu, *Nonlinear Semigroups and Differential Equations in Banach Spaces*, Springer, Netherlands, 1976.
- [32] Y. I. Alber, Metric and generalized projection operators in Banach spaces: properties and applications. in: *Theory and Applications of Nonlinear Operators of Accretive and Monotone Type*, vol. 178, Dekker, New York, 1996, pp. 15–50.
- [33] R. P. Agarwal, D. O'Regan, D. R. Sahu, *Fixed Point Theory for Lipschitzian-type Mappings with Applications*, Springer, New York, 2009.

Entropy Modeling of Optimal Intelligence Development in Regards with the Air Transport Operation

Andriy V. Goncharenko ^{1,2}

¹ Xi'an Jiaotong University, No.28, Xianning West Road, Xi'an Shaanxi, 710049, P. R. China

² National Aviation University, 1, Liubomyra Huzara Avenue, Kyiv, 03058, Ukraine

Abstract

The paper is devoted to the entropy computer modeling of optimal intelligence development in regards with the air transport operation. The subjective analysis theory of the active systems is used as a framework for theoretical elaborations. The contemplations are based upon the subjective entropy paradigm. Several solutions are obtained for a few special cases considered. Conditional optimization of the subjective individuals' preferences functions entropy in conjunction with the proposed hybrid combined relative pseudo-entropy function happened to be helpful in determining the relative certainty/uncertainty degree concerning prevailing/dominating subjective preferences functions. Illustrative examples simulations are performed. Necessary diagrams are plotted.

Keywords

Entropy, preferences, operation, air transport, optimization, intelligence, management, simulation, objective functional.

1. Introduction

Rational methods of the air transportation management systems functioning in the operation are important. Aircraft maintenance and repair [1, 2] procedures are aimed at the aeronautical engineering reliability support, as well as at the decrease of the risks and negative consequences of the aircraft systems, systems' components, and systems' elements failures [3, 4].

The models of intelligence used at the air transportation management systems functioning simulation should take into account the elements of the aviation transportation systems activities usefulness [5, 6].

The uncertainty degree could be evaluated in the framework of the entropy maximum theory terms as in [7 – 9]. The trend of the entropy research in modern science is very popular altogether [10]. Assessing economical parameters and terms [11], it is logically to combine the issues of the publications of [7 – 11] into the subjective analysis theory [12] that has been successfully developed by Professor Kasianov V. A. at the National Aviation University, Kyiv, Ukraine for about the last three decades.

In fact, the intelligence learning aspects might be and should be implemented to the modeling of the different nature system and processes, as for example, for the applicable features of the [13 – 17] publications.

Therefore, the goal of the presented herewith study is to demonstrate the advantages of the entropy paradigm [7 – 9, 12, 17 – 20] applicably to the generalized value of the air transportation management systems functioning learning potential.



2. General approach

Intelligent air transportation management systems' functioning requires that the corresponding information flow be processed quite effectively and adequately. Participants of an economic process should have a possibility, taking into account available resources needed for running their own business, to choose that or another attainable (achievable, reachable) alternative in the problem-resource situation having been formed.

At this, proceeding from some theoretical speculations and [12, 17 – 20], the subject (the active element of the intelligent air transportation management system) distributes the preferences functions in accordance with the postulated optimality [12].

The problem is formulated as to discover the magnitude (possibly some relative value) and direction of the intelligent conflict situation certainty or uncertainty.

For that purpose, it is proposed to apply the entropy by Shannon, which has been developed for the probabilities. It is transformed for the intelligent conflict preferences similar to the references of [12, 17 – 20]:

$$H_{\pi} = - \sum_{i=1}^N \pi_i \ln \pi_i, \quad (1)$$

where i – the subscript that refers to the corresponding conflicting alternative that can be attained; N – the total number of the taken into consideration alternatives deemed to be conflicting; π_i – intelligent preferences functions that are to be found.

For the intelligent conflict management preferences π_i of conflicting alternatives it is going be used canonical expressions of [12, 17 – 20]:

$$\pi_i(t) = \frac{\exp[-\beta_i F_i]}{\sum_{j=1}^N \exp[-\beta_j F_j]}, \quad (2)$$

where β_i and β_j – the structure parameters, their corresponding values relate to the problem setting's objective functional of [12, 17 – 20], F_i and F_j – corresponding intelligent conflict management functions expressing the effectiveness of the i -th and j -th conflicting alternatives.

Modifications on the entropy of the view of (1), required for the stated problem solution, are going to be presented and described below herein.

3. Main Content

The traditional view entropy of (1) has some improperness to a certain degree. Let us say in a two conflicting alternatives situation the distribution of the intelligent conflict preferences are as follows:

$$\pi_1 = 0.351 \text{ and } \pi_2 = 1 - \pi_1 = 0.649. \quad (3)$$

This means that entropy will be the same for the opposite situation situation:

$$\pi_1 = 0.649 \text{ and } \pi_2 = 1 - \pi_1 = 0.351. \quad (4)$$

Thus, the entropy of the view (1) shows no difference with respect to the directions of the certainty or uncertainty of the intelligent conflict management alternating preferences. And that circumstance, as in the case of (3) and (4) with the mirror reflection preferences distribution change, pertains to any distribution of conflicting preferences. Hence, it is impossible to realize which uncertainty or certainty is a “good” one and which is a “bad”. This attitude can be as “right” versus “wrong”.

3.1. General provisions

The relative function [17] as hybrid pseudo-entropy fits the problem solution requirements:

$$\bar{H}_{\max - \frac{\Delta\pi}{|\Delta\pi|}} = \frac{H_{\max} - H_{\pi}}{H_{\max}} \frac{\Delta\pi}{|\Delta\pi|} = \frac{H_{\max} + \sum_{i=1}^N \pi(\sigma_i) \ln \pi(\sigma_i)}{H_{\max}} \frac{\left[\sum_{j=1}^M \pi(\sigma_j^+) - \sum_{k=1}^L \pi(\sigma_k^-) \right]}{\left| \sum_{j=1}^M \pi(\sigma_j^+) - \sum_{k=1}^L \pi(\sigma_k^-) \right|}, \quad (5)$$

where $H_{\max}$ – the maximally possible value of the entropy, in problems formulated in references of [12, 17 – 20] it is

$$H_{\max} = \ln N, \quad (6)$$

$\Delta\pi$ – the factor of the intelligent conflict management alternatives preferences functions domination, proposed in [17]:

$$\Delta\pi = \sum_{j=1}^M \pi(\sigma_j^+) - \sum_{k=1}^L \pi(\sigma_k^-), \quad (7)$$

where σ_j^+ – alternatives that are considered to be positive and σ_k^- – alternatives with the negative content; M – the quantity of the positive alternatives; L – the quantity of the negative subset of the alternatives, [17]:

$$M + L = N. \quad (8)$$

3.2. Specific cases models construction

Let us consider a basic two-alternative situation:

$$v_T(m) = \sqrt[4]{\frac{4}{3} \frac{bm^2 g^2}{C_{x_0} \rho^2 S^2}}, \quad v_L(m) = \sqrt[4]{4 \frac{bm^2 g^2}{C_{x_0} \rho^2 S^2}}, \quad (9)$$

where $v_T(m)$ – the speed of an aircraft horizontal flight that is optimal for the maximal duration, it is obtained as a function of the changeable mass of the aircraft m ; b – the aircraft special aerodynamic coefficient; g – the acceleration of the force of gravity; C_{x_0} – one more aerodynamic coefficient of the aircraft drag when the force of the aircraft wing lift equals “zero”; ρ – the air density at the flight conditions; S – the area that characterizes the aircraft aerodynamics properties; $v_L(m)$ – one more optimal speed, this time for the aircraft flying in a horizontal path and intended for the maximal distance, it is also expressed as a function in the terms of the aircraft mass m when it changes.

The aircraft optimal speeds of (9) are obtained as extremum solutions delivering maximum values to the objective functionals of the aircraft horizontal flights; and these functions are considered in the presented study as the intelligent cyber conflict management effectiveness functions of the aircraft horizontal flight effectiveness.

Another two-alternative situation model is with taking into account the intelligent conflict management effectiveness functions in the view of the solution of the ordinary differential equation systems of the first order:

$$\left. \begin{aligned} \frac{dy_0}{dt} &= \left(1 - \frac{y_0}{Y_0} \right) (k_{00} y_0 - k_{10} y_0 y_1) \\ \frac{dy_1}{dt} &= \left(1 - \frac{y_1}{Y_1} \right) (-k_{11} y_1 + k_{01} y_0 y_1) \end{aligned} \right\}, \quad (10)$$

where y_0 – the first of the two alternative intelligent conflict management effectiveness functions; t – time; Y_0 – marginal value for the first of the two alternative intelligent conflict management effectiveness functions y_0 ; k_{00} – coefficient of the first function value supposed exponential growth; k_{10} – coefficient of the impact of the second of the two alternative intelligent cyber conflict management effectiveness functions upon the first one, which, by assumption, decreases the rate of

the first conflicting function growth; y_1 – the second of the two alternative intelligent conflict management effectiveness functions; Y_1 – marginal value for the second of the two alternative intelligent conflict management effectiveness functions y_1 ; k_{11} – coefficient of the second function value supposed exponential decrease; k_{01} – coefficient of the impact of the first of the two alternative intelligent conflict management effectiveness functions upon the second one, which, by assumption, increases the rate of the second conflicting function growth.

The three-alternative case is like the previous but extended:

$$\left. \begin{aligned} \frac{dy_0}{dt} &= -k_{00}y_0 \left(1 - \frac{y_0}{Y_0}\right) + k_{10}y_1 + k_{20}y_2 \\ \frac{dy_1}{dt} &= k_{01}y_0 + k_{11}y_1 \left(1 - \frac{y_1}{Y_1}\right) - k_{01}k_{21}y_0y_2 \\ \frac{dy_2}{dt} &= k_{02}k_{12}y_0y_1 - k_{22}y_2 \left(1 - \frac{y_2}{Y_2}\right) \end{aligned} \right\}, \quad (11)$$

where designations and interpretations of functions, coefficients, and values are analogous to the previous case (10), but it was customized and extended to the intelligent cyber conflict management situation with some three alternatives.

One more special case to be studied is when there are generalized parameters of the intelligent system learning.

The objective functional is

$$\Phi_\pi = -\sum_{i=1}^2 \text{Pr}_i \ln(\text{Pr}_i) + \beta[\text{Pr}_1 R + \text{Pr}_2 R_0] + \gamma \left(\sum_{i=1}^2 \text{Pr}_i - 1 \right), \quad (12)$$

where Pr_1 and Pr_2 are the generalized perception functions of the two alternatives for resources; R and R_0 are the input learning variable and threshold value resources correspondingly; γ is the coefficient for the subjective preferences normalized assessment; it is likewise β , β_i , and β_j ; they are the corresponding weight coefficients or they might be the structure parameters that are internal, also, these coefficients can be considered as the uncertainty Lagrange multipliers [12]. For the presented study these parameters are interpreted as the internal intelligent control parameters which have some properties of the intelligent object “attitude” to the alternatives [12].

The generalized perceptions functions of Pr_1 and Pr_2 are analogous to the preferences functions of (1) – (4).

Extremizing the objective functional (12) under conditions of

$$\frac{\partial \Phi_\pi}{\partial \pi_i} = 0, \quad (13)$$

one can get the expressions similar to (2):

$$\text{Pr}_1 = \frac{\exp[\beta R]}{\exp[\beta R] + \exp[\beta R_0]}, \quad \text{Pr}_2 = \frac{\exp[\beta R_0]}{\exp[\beta R] + \exp[\beta R_0]}. \quad (14)$$

The corresponding potential for the intelligence growth could be expressed as

$$V(R) = V_0 \text{Pr}_2(R), \quad (15)$$

where V_0 is the amount of the intelligence potential available for the intelligence growth.

The intelligence growth output could be represented with one more generalized function:

$$\text{Output}(R) = V(R)R. \quad (16)$$

3.3. Solutions to the specific cases

In first situation (9) the solution is shown in the Figure 1.

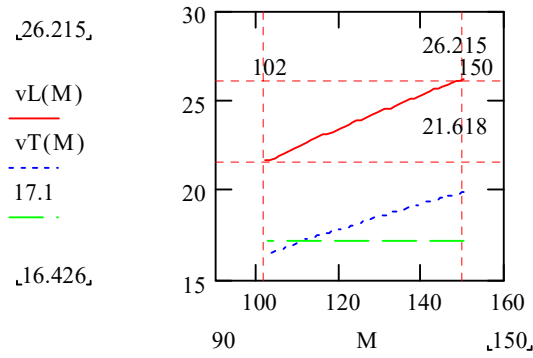


Figure 1: Intelligent conflict of optimal aircraft speeds

The designations in the Figure 1 are as follows: $vL(M)$ is for $v_L(m)$, obtained by the second equation of (9); $vT(M)$ is for $v_T(m)$, obtained by the first equation of (9) respectively.

The results of the system (10) solution are presented in the Figures 2 and 3.

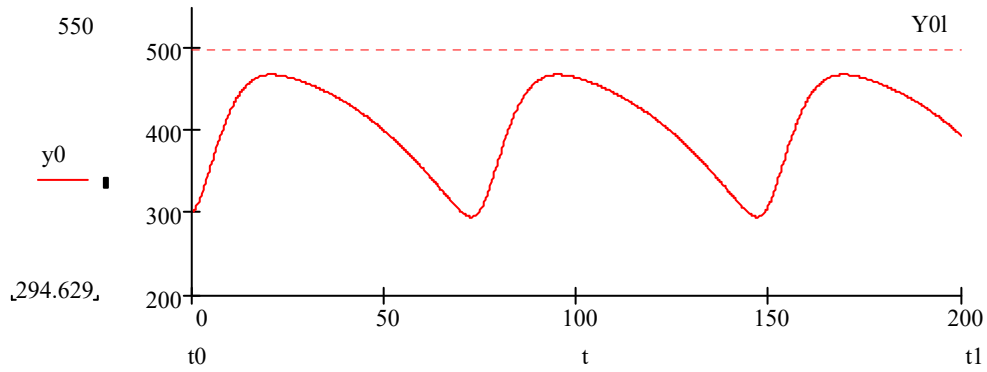


Figure 2: Intelligent cyber conflict system internal self-management for the self-growing effectiveness function

In the Figures 2 and 3 it is designated: y_0 is for y_0 , $Y01$ is for Y_0 ; and y_1 is for y_1 , $Y11$ is for Y_1 correspondingly.

The third case of solution, with the system of equations (11), is demonstrated with the diagrams plotted in the Figures 4 – 6.

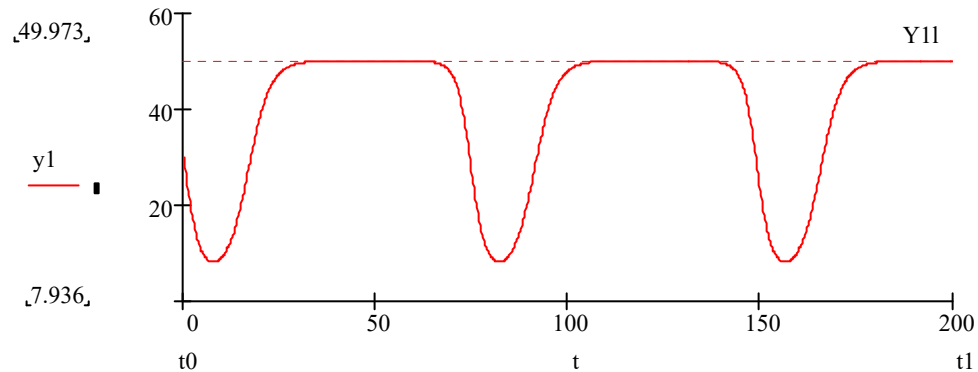


Figure 3: Intelligent conflict system internal self-management for the self-decreasing effectiveness function

The designations in the Figures 4 – 6 are analogous to those for the Figures 2 and 3, and simply extended to the considered three-alternative intelligent conflict situation.

The solutions expressed with the formulae of (14) – (16) to the special case of (12) under conditions of (13) has already been described above.

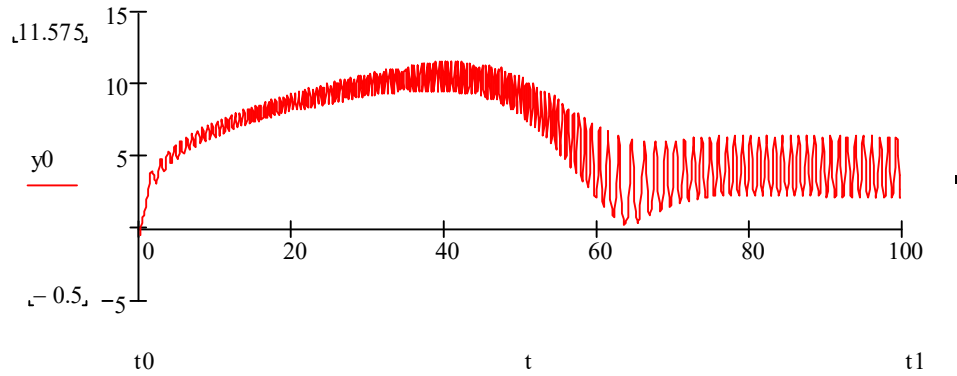


Figure 4: Intelligent conflict system internal self-management for the self-decreasing effectiveness function in case of the three-alternative situation

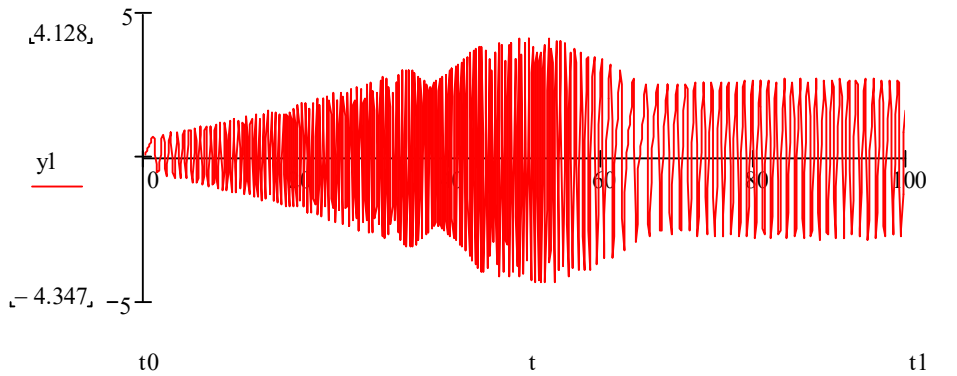


Figure 5: Intelligent conflict system internal self-management for the self-growing effectiveness function in case of the three-alternative situation

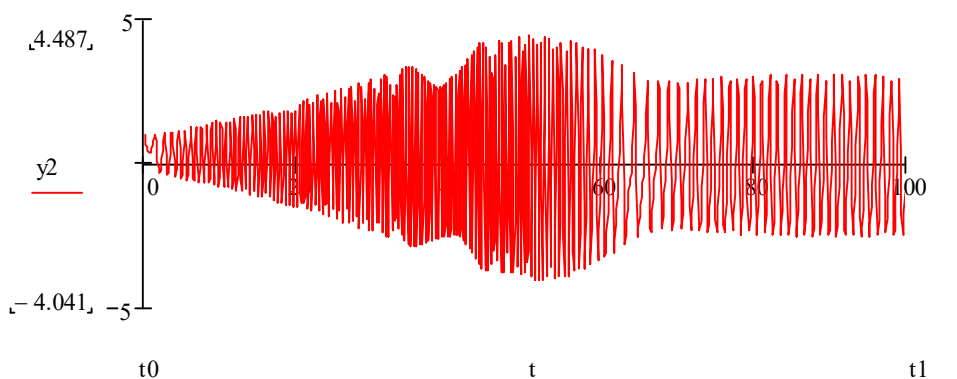


Figure 6: Intelligent conflict system internal self-management for the other self-decreasing effectiveness function in case of the three-alternative situation

3.4. Computer simulation

When the system of equation (11) is adapted to the composition that combines the complete self-management in the intelligent cyber conflict, it gives the results of computer simulation shown in the Figures 7 – 9.

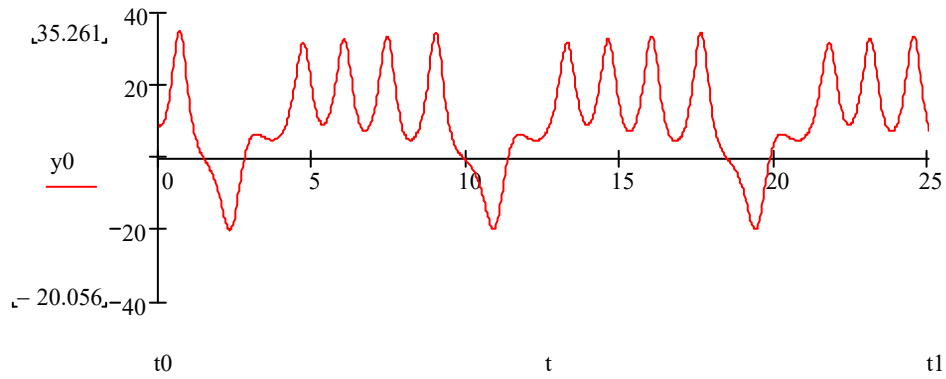


Figure 7: Intelligent conflict system complete internal self-management for the self-decreasing effectiveness function in case of the three-alternative situation

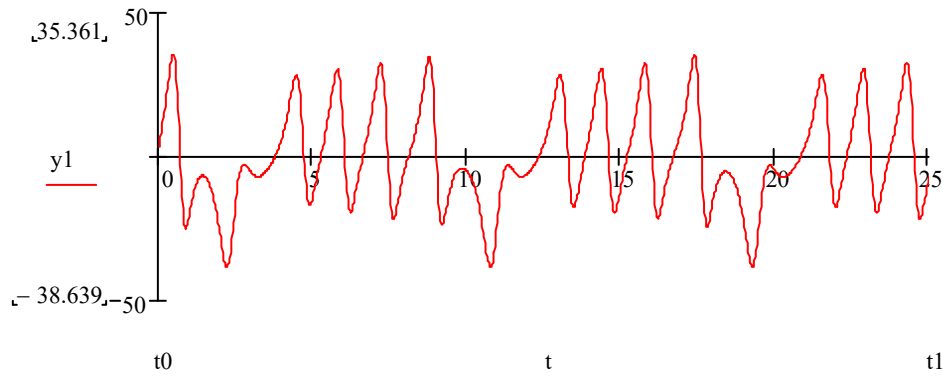


Figure 8: Intelligent conflict system complete internal self-management for the self-growing effectiveness function in case of the three-alternative situation

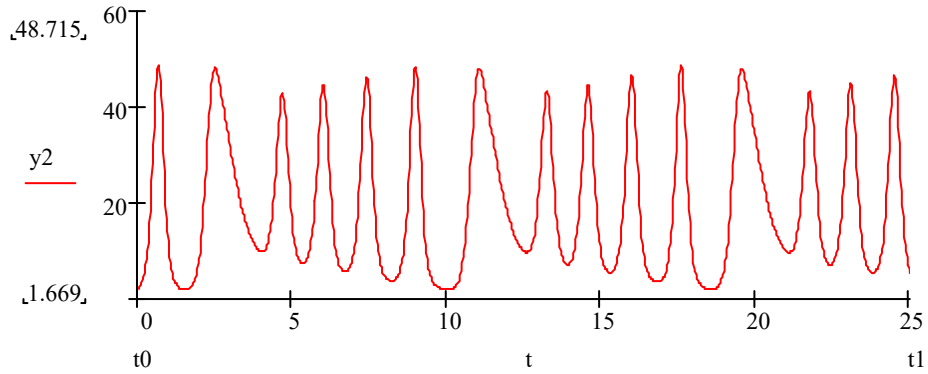


Figure 9: Intelligent conflict system complete internal self-management for the other self-decreasing effectiveness function in case of the three-alternative situation

The intelligent cyber conflict management preferences computed by (2) are illustrated in the Figure 10.

In the Figure 10: p_0 stands for π_0 , p_1 stands for π_1 , and p_2 stands for π_2 , correspondingly.

The traditional entropy of the intelligent conflict management preferences calculated by (1) is plotted in the Figure 11.

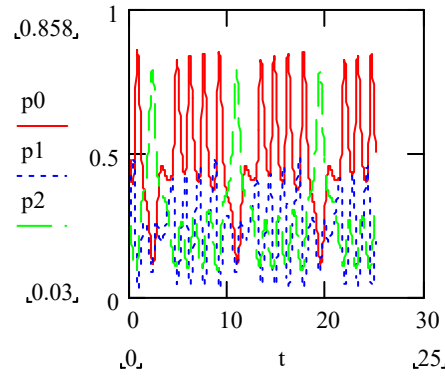


Figure 10: Intelligent conflict system complete internal self-management effectiveness functions preferences in case of the three-alternative situation

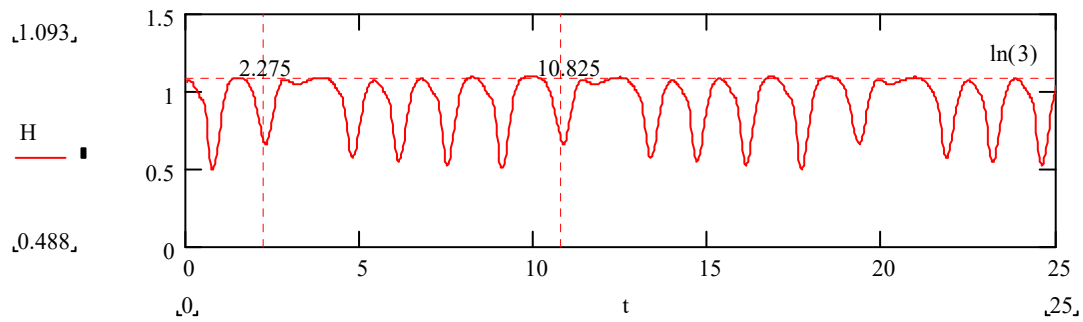


Figure 11: Entropy of intelligent conflict system complete internal self-management effectiveness functions preferences in case of the three-alternative situation

The pointed above hybrid pseudo-entropy relative function (5) of [17], when π_2 is a “positive” preference and π_0 and π_1 are not, that is they are considered as “negative” preferences, is shown in the Figure 12.

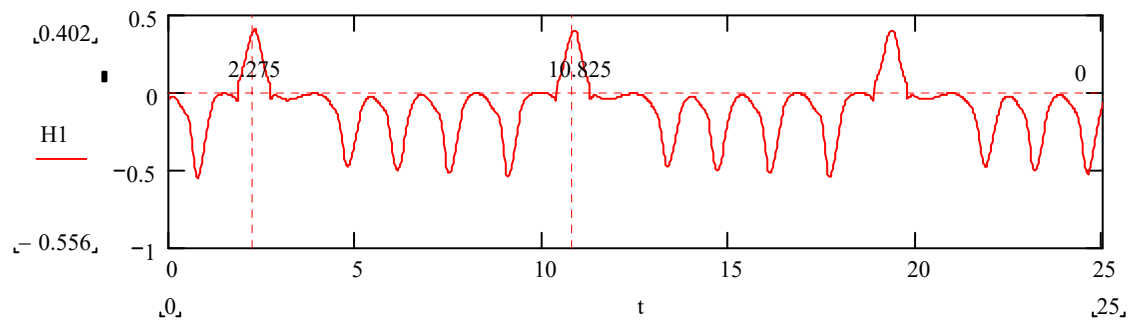


Figure 12: Hybrid pseudo-entropy relative function of intelligent conflict system complete internal self-management effectiveness functions preferences in case of the three-alternative situation

The results of computer simulation in case of (12) – (16) are represented in the Figures 13 – 15.

The generalized intelligence perception function curve, plotted in the Figure 13, is for the calculations by the second equation of (14) with the value of the threshold generalized resource of $R_0 = 100$.

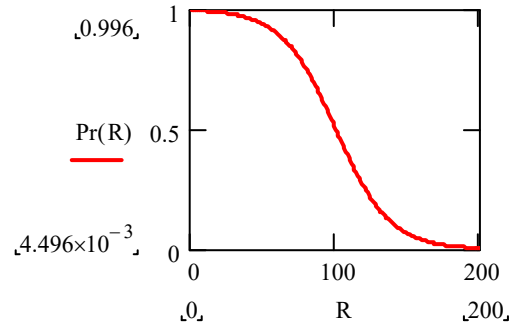


Figure 13: Generalized intelligence perception function

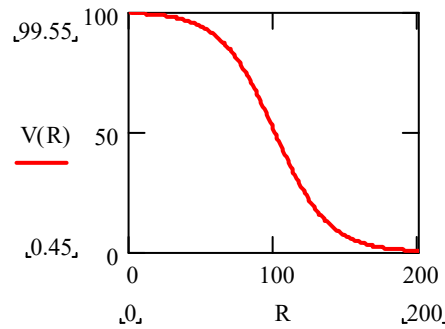


Figure 14: Generalized intelligence potential for the intelligence growth function

The generalized intelligence potential function (see the Figure 14) is used for the intelligence growth function shown in the Figure 15.

It is calculated with the amount of the intelligence potential available for the intelligence growth $V_0 = 100$.

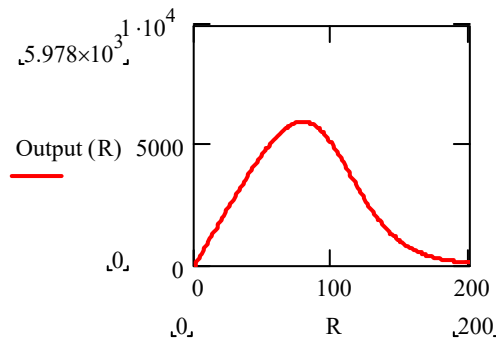


Figure 15: Generalized intelligence growth output function

4. Discussion

As can be seen from the results of the computer simulation (see the Figures 1 – 12) of the intelligent conflict management models involving entropy tools of (1) – (8) in application to (9) – (11), the hybrid pseudo-entropy relative function (5) of [17], of the intelligent conflict management effectiveness functions preferences has advantages over the traditional entropy (1).

The optimal generalized resource value for the intelligence growth output (see the Figures 13 – 15) is also obtainable with the help of the entropy paradigm (12) – (16).

The optimal value, which can be seen in the Figure 15, is lower than the threshold value accepted in the calculation simulations.

4.1. Comparison analysis to the known results

In the case of conflicting alternative speeds of the aircraft horizontal flight (see the curves calculated by the equations of (9) and plotted in the Figure 9), the intelligent cybernetic conflict management function with respect to the available and achievable conflicting alternatives (options) preferences functions entropy in the traditional view of (1) will not show the relative certainty or uncertainty degree for the alternatives, and its direction either. The proposed hybrid pseudo-entropy relative function (5) will show those required qualities and quantities.

Such effects are noticeable when comparing the entropies plotted in the Figures 11 and 12.

The entropy illustrated in the Figure 11 does not represent when, how much, and to which alternative or group of alternatives, that is to “good” or “bad”, “correct” or “wrong”, the intelligent conflict management has its inclination.

Whereas, the proposed combined hybrid pseudo-entropy relative function (5), that takes into account the relative value of the intelligent conflict management situation uncertainty, together with the composition with the intelligent conflict management effectiveness preferences functions index of domination, shows that there are periods of time when the certainty of the considered intelligent cyber conflict management has some definitely “positive” values. These values of time: $t \approx 2.275$ and $t \approx 10.825$ are represented in the Figure 12. Also, there is such effect at the time about $t \approx 19.325$.

The mentioned above points in time portrait practically not that much referred back to the traditional measure of uncertainty pictured in the Figure 11.

In the special case described with the equations of (12) – (16), illustrated in the Figures 13 – 15, the optimal input value of the generalized resources R delivers maximum value to the generalized intelligence growth output.

4.2. Evidently promising investigations

The calculation experimentations with the procedures described with the mathematical expressions of (1) – (11) and their adaptations have been conducted with several abstracted supposed data values and some initial conditions voluntary selected to a certain degree. Therefore, there is no need in their indication herewith.

Though the main ideas and provisions of the presented study are quite comprehensively performed in the models, there is a potential for the further research in the areas of the coefficient estimations. However, some other representations are also possible, as well as some other models of the similar interpretation could be elaborated.

The input value in the special case described with the equations of (12) – (16) is the generalized resources and it is necessary to investigate more complex models of the resource-output relations.

5. Conclusion

The entropy theory of conflicts in general sense can be successfully implemented to the solutions of the intelligent conflict management in the terms of the conflicts genesis, development, transitions, and exodus. The intelligent conflict being a primary reason for the intelligent cyber system motion can be properly managed with the relative function based upon the hybrid pseudo-entropy.

The optimal generalized input intelligent resources value found with the use of the resources generalized perceptions functions entropy conditional extremization ensures the maximum value to the generalized intelligence growth output. Further endeavors in the intelligent conflict management have prospects in the used values estimations and methodological modifications.

6. References

- [1] M. J. Kroes, W. A. Watkins, F. Delp, R. Sterkenburg, Aircraft Maintenance and Repair, 7th. ed., McGraw-Hill, Education, New York, NY, 2013.

- [2] T. W. Wild, M. J. Kroes, Aircraft Powerplants, 8th. ed., McGraw-Hill, Education, New York, NY, 2014.
- [3] B. S. Dhillon, Maintainability, Maintenance, and Reliability for Engineers, Taylor & Francis Group, New York, NY, 2006.
- [4] D. J. Smith, Reliability, Maintainability and Risk. Practical Methods for Engineers, Elsevier, London, 2005.
- [5] R. D. Luce, D. H. Krantz, Conditional expected utility, *Econometrica* 39 (1971) 253–271.
- [6] R. D. Luce, Individual Choice Behavior: A theoretical analysis, Dover Publications, Mineola, NY, 2014.
- [7] E. T. Jaynes, Information theory and statistical mechanics, *Physical review* 106 (4) (1957) 620–630. doi:10.1103/PhysRev.106.620.
- [8] E. T. Jaynes, Information theory and statistical mechanics. II, *Physical review* 108 (2) (1957) 171–190. doi:10.1103/PhysRev.108.171.
- [9] E. T. Jaynes, On the rationale of maximum-entropy methods, *Proceedings of the IEEE* 70 (1982) 939–952. doi:10.1109/PROC.1982.12425
- [10] F. C. Ma, P. H. Lv, M. Ye, Study on global science and social science entropy research trend, in: *Proceedings of the IEEE International Conference on Advanced Computational Intelligence, ICACI*, Nanjing, Jiangsu, China, 2012, pp. 238–242. doi:10.1109/ICACI.2012.6463159
- [11] E. Silberberg, W. Suen, *The Structure of Economics. A Mathematical Analysis*, McGraw-Hill Higher Education, New York, NY, 2001.
- [12] V. Kasianov, Subjective Entropy of Preferences. Subjective Analysis, Institute of Aviation Scientific Publications, Warsaw, Poland, 2013.
- [13] O. Solomentsev, M. Zaliskyi, T. Herasymenko, O. Kozhokhina, Yu. Petrova, Data processing in case of radio equipment reliability parameters monitoring, in: *Proceedings of the International Conference on Advances in Wireless and Optical Communications, RTUWO*, Riga, Latvia, 2018, pp. 219–222. doi:10.1109/RTUWO.2018.8587882.
- [14] V. Marchuk, M. Kindrachuk, Ya. Krysak, O. Tisov, O. Dukhota, Y. Gradiskiy, The mathematical model of motion trajectory of wear particle between textured surfaces, *Tribology in Industry (Tribol. Ind.)* 43 (2) (2021) 241–246. doi: 10.24874/ti.1001.11.20.03.
- [15] D. Shevchuk, O. Yakushenko, L. Pomytkina, D. Medynskyi, Y. Shevchenko, Neural network model for predicting the performance of a transport task, *Lecture Notes in Civil Engineering*. 130 LNCE (2021) 271–278. doi:10.1007/978-981-33-6208-6_27.
- [16] S. Subbotin, The neuro-fuzzy network synthesis and simplification on precedents in problems of diagnosis and pattern recognition, *Opt. Mem. Neural Networks* 22 (2013) 97–103. <https://doi.org/10.3103/S1060992X13020082>.
- [17] A. V. Goncharenko, Multi-optional hybridization for UAV maintenance purposes, in: *Proceedings of the IEEE International Conference on Actual Problems of UAV Developments, APUAVD*, Kyiv, Ukraine, 2019, pp. 48–51. doi:10.1109/APUAVD47061.2019.8943902.
- [18] A. V. Goncharenko, Active systems communicational control assessment in multi-alternative navigational situations, in: *Proceedings of the IEEE International Conference on Methods and Systems of Navigation and Motion Control, MSNMC*, Kyiv, Ukraine, 2018, pp. 254–257. doi:10.1109/MSNMC.2018.8576285.
- [19] A. Goncharenko, A multi-optional hybrid functions entropy as a tool for transportation means repair optimal periodicity determination, *Aviation* 22 (2) (2018) 60–66. doi: 10.3846/aviation.2018.5930.
- [20] A. V. Goncharenko, Airworthiness support measures analogy to the prospective roundabouts alternatives: theoretical aspects, *Journal of Advanced Transportation* Article ID 9370597 2018 (2018) 1–7. doi: 10.1155/2018/9370597.

Integration of Decision-Making Stochastic Models of Air Navigation System Operators in Emergency Situations

Tetiana Shmelova¹, Kseniia Lohachova² and Maxim Yatsko³

¹National Aviation University University, Liubomyra Huzara ave. 1, Kyiv, 03058, Ukraine

²National Aviation University University, Liubomyra Huzara ave. 1, Kyiv, 03058, Ukraine

³National Aviation University University, Liubomyra Huzara ave. 1, Kyiv, 03058, Ukraine

Abstract

The authors present a new objective-subjective approach to conflict management to ensure proper collaboration of different aviation personnel using decision-making (DM) methods in Certainty, Risk, and Uncertainty. For improvement of outcomes of the collective solutions the collaborative decision-making (CDM) models are used, when is difficult to make the final right decision based on many factors influencing DM of different aviation specialists (pilot, air traffic controller, flight dispatcher) during the development of an emergency in flight. The Integration of Uncertainty Models to the Collective Model of CDM has been presented. An example is presented - the building of individual models and CDM models for the pilot-in-command, flight dispatcher, and air traffic control officer in bad weather conditions (lightning strikes and thunderstorm activity) during the approach of aircraft. The optimal landing aerodrome was chosen for landing in bad weather conditions.

Keywords

Integration of Stochastic and Non-stochastic Models, Collaborative Decision Making, Decision Making under Uncertainty, Emergency, objective-subjective approach, individual and collaborative models, pilot, flight dispatch, air traffic controller, objective and subjective factors

1. Introduction

In the process of aviation industry development, safe air transportation is the core objective of the functioning of the aviation system. But despite the rapid and constant growth from the point of developed technologies applied in the operational processes of air traffic control, flight planning, modern airplanes manufacturing, and improvement of hours flown by pilots and flight performance, the number of aviation accidents does not decrease. Detailed aviation accident statistics are presented in the Annual safety reviews of the European Aviation Safety Agency (EASA) and the guide to European statistics "Statistics Explained" [1; 2]. According to the published data from the EASA, the number of aviation accidents during the last years has grown [2]. There were 789 fatalities in aviation accidents involving European Union (EU) - registered aircraft over the period 2016-2020 [1]. Most of the air accidents dealt with general aviation, near 80% (as known, general aviation consists of all civil aviation aircraft other than commercial air transport) [1]. As per EASA annual safety review, the number of non-fatal accidents in 2019 was higher than the average 10-year period before [2]. The number of occurrences depended on the size and load of the aircraft involved in the accident, and the complexity of emergency situations. Most of the occurrences in 2019 were related to difficult meteorological conditions [2]. Besides, the possible sources for such statistics may have next reasons [3; 4]:

1. Insufficient pre-flight preparation
2. Inadequate training of aviation personnel
3. Gaps in developed procedures and manuals, situational unawareness

International Workshop on Computer Modeling and Intelligent Systems (CMIS), May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: shmelova@ukr.net (T. Shmelova); kseniialogachova@gmail.com (K. Lohachova); maxim_yatsko@i.ua (M. Yatsko)
ORCID: 0000-0002-9737-6906 (T. Shmelova); 0000-0002-3802-5540 (K. Lohachova); 0000-0003-0375-7968 (M. Yatsko)



© 2020 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

4. Availability of single guidance for multiple users for non-conflict decision-making (DM), especially in emergency situations
5. The influence of the socio-technical environment and non-professional factors (psychophysiological, individual psychological, socio-psychological) on the person
6. Psycho-emotional states of operators.

Most of these reasons belong to human factor. Human error is considered the first cause of accidents [2; 3; 4]. Indeed, as many safety scholars affirm, human error is only an epiphenomenon, that is a secondary phenomenon that occurs alongside or in parallel to a primary event. The real cause of accidents which induces an error are human performance and limitations, poor teamwork, organizational pressure on the professionals in the team to obtain unreasonable performance, faulty human-machine interaction, and conflict interaction of professionals in important DM too [3; 4]. It's clear, human factor should be included for progressive and sustainable development and safety enhancement of the complex aviation system. Safety is the functioning of certain operational processes in the aviation systems with the objective of controlling the safety risks of the outcomes of hazards during operation. International Civil Aviation Organization (ICAO) developing proactive approaches for the safety provision based on risk and human performance assessment [5]. The one from the modern concepts is the information and intelligent support of collective solutions of different specialists in an emergency [5]. There are a lot of professionals involved in the provision of safety during flight planning and operations process. They are flight dispatchers, flight crews, air traffic controllers, maintenance staff, ground handling personnel, etc. Each of them plays a major role at different stages, as the safe flight starts not only from the aircraft departure. They strictly follow the manuals and legal documents approved in the field of their professional activity [6; 7; 8; 9].

The main actions of aviation specialists in accordance with the stage aircraft operation are presented in the Table 1.

Table 1

The main actions of aviation specialists according to the stage of aircraft operation

Aviation specialist	Stage of flight	Operations
Flight dispatch (FD)	Planning	Responsible for flight planning and organization, for choice of the optimal flight route, alternate aerodromes, and proper fuel amount calculation for definite flights, regulated by international and national documents, orders, and instructions for the normal and abnormal operational environment of aircraft [6]
Pilot-in-command (PIC)	Flight	Is holding the full right of DM before departure and taking all responsibility in flight, following existing aircraft flight and operations manuals, quick reference handbooks (QRH) in case of emergency and abnormal conditions according to the type of aircraft [7]
Air traffic control officer (ATCO)	Safety of flights of aircraft in the control sector of airspace	Ensures required aircraft separation minima established in each sector of airspace and provides flight crews with assistance in emergencies, according to the instructions defined and approved within particular air traffic control sector, by national laws, letters of agreement between neighboring countries, and handbooks in case of aircraft emergencies [8; 9]

Sometimes specialists from other fields, such as emergency and ground services, and medical staff may be involved in joint DM (Figure 1). For example, the digital health and telemedicine in emergencies are applied, allowing to invite qualified medical personnel for a consultation in case of incapacitation of one of the pilots or one of the passengers [10; 11].

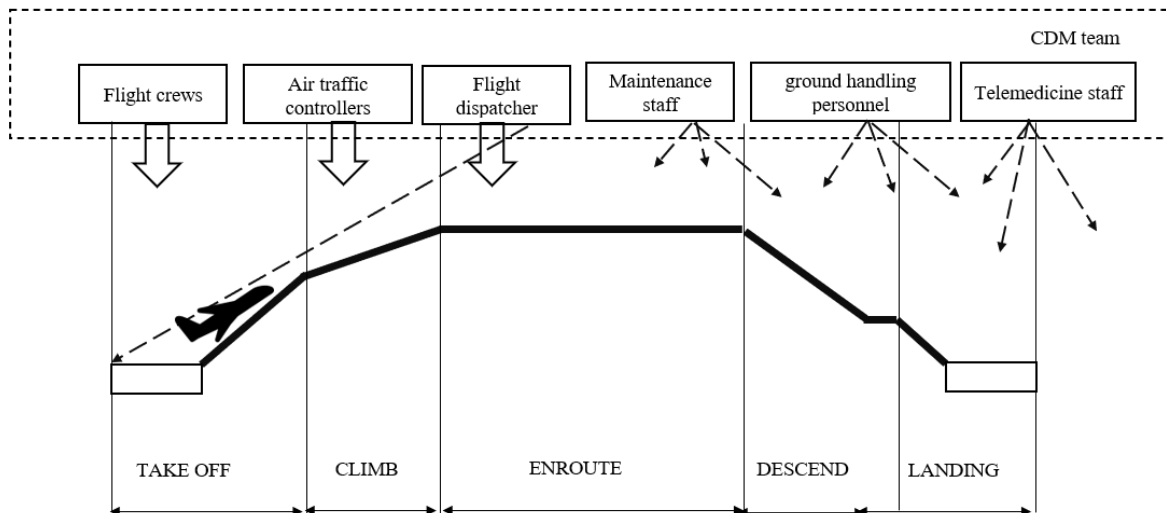


Figure 1: The team of human-operators involved in Collaborative Decision-Making during the flight

Mostly, the complexity, content, particularity of documents, regulations of the professional activity of each aviation personnel are different, which does not allow to develop the general algorithm of actions for all aviation staff in a specific situation, especially for difficult in-flight conditions, where the uncertainty, lack of information and time for DM take place. That's why arises the conflict between actions and decisions of involved staff who making the decisions at the same time for one situation. Therefore, to ensure proper collaboration and conflict resolution between different aviation personnel for improvement of outcomes of group decisions, it is necessary to implement new approaches for conflict management and integrate DM stochastic and non-stochastic models for Air Navigation System (ANS) operators, especially in an emergency [12; 13].

In the concept of the Global Air Navigation Plan developed by ICAO [14], the proper collaboration is possible through the provision of air traffic management system (ATM) members with an environment that ensures enough storage of significant information and its proper usage among ATM system. In this environment, all members closely interact with each other in making common decisions and achieve so-called Collaborative Decision-Making (CDM). This concept foresees the improvement of the whole performance of the ATM system in general taking into consideration the individual performance of ATM system members. This approach allows choosing the direction of action with respect to selected objectives, based on decisions made by each participant, and information exchange influenced those decisions, applying main DM principles [15]. That is why ICAO promotes the global implementation of performance management principles, gradually transferring the existing ATM systems to performance-based global ANS. The performance-based approach (PBA) is a mean of establishing the performance management process and is based on three principles of obtaining effectiveness of solutions [16]:

- A strong and competent focus on desired results of DM
- Conscious and rightly of DM
- Reliance on real facts and data for rational DM.

Hence, according to the ICAO requirements, the meeting of the day-to-day performance in operations may be achieved through the mechanism of Flight and Flow Information for a Collaborative Environment (FF-ICE) Concept [17]. Concept FF-ICE defines air navigation information requirements for flight planning, flow management, air traffic management, and trajectory management, and to be a basis of the performance-based ANS [17].

In the monograph, "Intelligent Automated System for Supporting the Collaborative Decision Making by operators of the Air Navigation System during Flight Emergencies", was researched the PIC and ATCO CDM during flight emergencies using deterministic models in order to achieve the maximum synchronization of operators' technological procedures. Collaborative Decision Models are created for only two operators such as PIC and ATCO [18]. The ability to reach a general consent on the desired outcome to be achieved by all ANS operators (FDs, flight crews, ATCOs, maintenance staff,

ground handling personnel) in terms of performance results is the basic condition for the successful application of the approach for conflict management, especially in emergency situations [13]. The application of classical DM criteria under uncertainty makes it possible to take into account the factors influencing the choice and find collective solutions for several participants in the process [19].

The purposes of the work are:

1. Integration of DM stochastic and non-stochastic models of ANS operators in emergency
2. To build the individual DM models for the PIC, FD, and ATCO in an emergency and to determine a collective solution for all operators-participants of the process.
3. To consider an example of building CDM models for the PIC, FD, and ATCO in conditions of uncertainty when it is required to choose an optimal alternate aerodrome in bad weather conditions before the approach.

2. The Integration of Uncertainty Models to Collective Model of Collaborative Decision-Making

For the effectiveness of DM of human-operators (H-O)'s of ANS in emergency situations the integration of Stochastic, and Non-stochastic models has been proposed [4; 12; 13]. There are the next main steps of the Method of the Integration of Stochastic models (DM under Risk and Uncertainty) and Non-stochastic models (DM in Certainty):

1. Analysis of the development of the emergency situation (ES) and synthesis of DM models in accordance with conditions, factors, potential participants, and stages of the development of the event (Figure 2).
2. Building deterministic DM models using Network Planning methods and determining output results (Figure 2):
 - critical time T_{cr} ; T_{mid} ; T_{min} ; T_{max}
 - critical paths of performance of actions of ANS operators in ES
 - ambiguous situations S_1 , S_2 , and synchronization of H-Os actions using stochastic models
 - alternate situations S and optimal CDM using stochastic models.

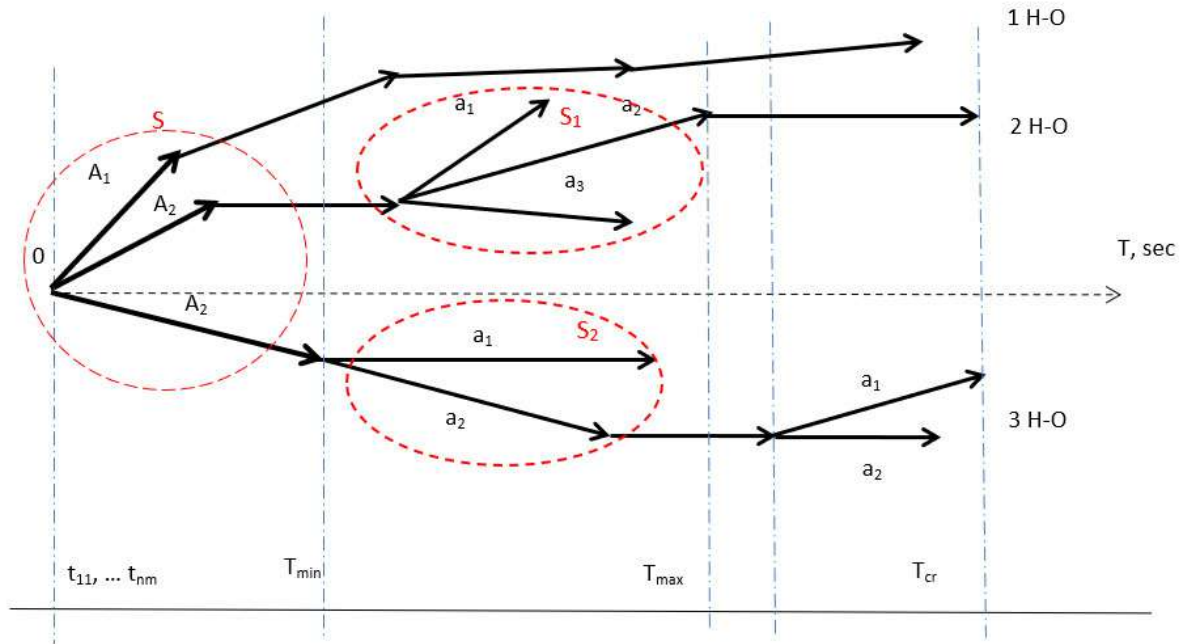


Figure 2: The deterministic DM models in ES for 3 human-operators

3. Analysis of procedures of the main technology, conditions of development of ESs, information about problem and conflict situations, and using the effective integration of DM models:

- If there is statistical information about the conditions of the development of the ES (the presence of probabilities of the development of the ES), then DM models under Risk conditions are applied and a decision tree is built (Stochastic Uncertainty DM models)
- In the case of absence of statistical information and probabilities, there are the factors which influence on the development of the situation, then DM models under Uncertainty are applied and a DM matrix is built (Non-Stochastic Uncertainty DM models)

4. Structural analysis of the ES and building of the Stochastic Uncertainty DM model (decision tree) according to conditions the development situation. Definition of the number of DM stages and time of stage t_i ; relevant output data: A_i ; p_i ; u_i ; β_i (alternatives at each stage A_i ; probabilities p_i and outcomes u_i for each alternative A_i ; added risks β_k influence on DM according to stages of the development situation, increasing threats due DM of H-Os. (Figure 3). Optimal solutions according to conditions of development of ESs. Risk R defined as:

$$R = A_{opt} = \min\{A_i\} = \min\left\{t_i\left(\sum_{i=1}^n p_i u_i - \beta_k\right)\right\} \quad (1)$$

where ...

- A_i - alternative decisions, $A = \{A_1, A_2, \dots, A_i, \dots, A_m\}$;
- p_j - probability of development situation, $p = \{p_1, p_2, \dots, p_i, \dots, p_m\}$;
- u_{ij} - expected outcomes of alternative actions, $U = \{u_1, u_2, \dots, u_i, \dots, u_m\}$;
- β_i - added risk (increasing threats), depend from stages of the development situation; $B = \{\beta_1, \beta_2, \dots, \beta_i, \dots, \beta_m\}$;
- t_i - time of stage of the development situation, $T = \{t_1, t_2, \dots, t, \dots, t_m\}$.

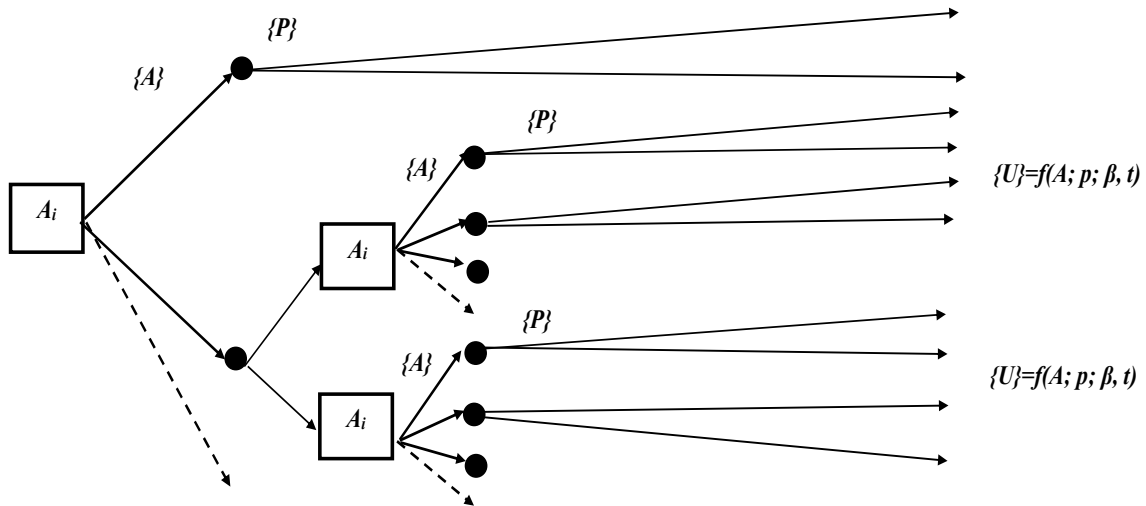


Figure 3: The Stochastic Uncertainty DM models in ES with outcomes $\{U\}=f(A; p; \beta, t)$

4. Factor analysis of the ES and building non-stochastic DM model (Table 2) in accordance with factors that influence DM as DM matrix:

- Alternative actions $A = \{A_1, A_2, \dots, A_i, \dots, A_m\}$
- Factors influence on DM according to situation $\Lambda = \{\lambda_1, \lambda_2, \dots, \lambda_j, \dots, \lambda_n\}, j = \overline{1, n}$
- Outcomes u_{ij} depending on actions / alternatives A_i , influence of variable (u_v) or static (u_s) factors - λ_j

For variable factors (u_v) is required to consider the development of the situation - time t_i (from Non-Stochastic DM model) and probability p_j (from Stochastic Uncertainty DM model). For static factors (u_s) is required to consider the normative documents [6 - 9], procedures connected with development of situation, type of aircraft, runway technical characteristics, characteristics of operators [20]. DM matrix $[U(A; \Lambda)]$ with H-O alternative actions $\{A\}$, factors $\{\Lambda\}$ and outcomes $\{U\}$ include (Table 2):

$$\{U(A; \Lambda)\} = \{U_s; U_v\} = \Lambda(t, p, N),$$

where

U_s – set of outcomes, influenced by the static factors;
 U_v – set of outcomes, influenced by the variable factors;
 t – time of development of ES;
 p – probabilities of development ES;
 N – normative documents according to development of ES.

Table 2
Matrix of DM in uncertainty, with variable or static factors

Alternative actions in ES	Variable factors influence DM in critical situation	Static factors influence DM in critical situation
A_1	$U_{v11}, \dots, U_{v1n}$	$U_{s1}, \dots, U_{sm}$
A_2	$U_{v2}, \dots, U_{v2n}$	$U_{s2}, \dots, U_{sm}$
A_3	$U_{v3}, \dots, U_{v3n}$	$U_{s3}, \dots, U_{sm}$

5. The optimal solutions are found using classical DM criteria under uncertainty (criterion Wald (maximin), criterion Laplace, criterion Hurwitz, criterion Savage) [19]. The choice of criterion depends on the type of aircraft, type of flight, the conditions for the development of the situation, the nature of the emergency. The short characteristic of criterion and flights:

- Criterion of Wald (maximin) - if this flight is performed for the first time:

$$A^* = \max_{A_i} \left\{ \min_{B_j} u_{ij}(A_i, B_j) \right\} \quad (2)$$

where

A_i – alternative solution from set $\{A\}$;
 B_j – factor from set of factors $\{A\}$.

- Criterion of Laplace - if this flight is regular:

$$A^* = \max_{A_i} \left\{ \frac{1}{m} \sum_{j=1}^n u_{ij}(A_i, B_j) \right\} \quad (3)$$

- Criterion of Hurwitz – different approach using optimism-pessimism coefficient α .
For charter flight:

$$A^* = \max_{A_i} \left\{ \alpha \max_{B_j} u_{ij}(A_i, B_j) + (1 - \alpha) \min_{B_j} u_{ij}(A_i, B_j) \right\} \quad (4)$$

where

α - optimism-pessimism coefficient, $0 \leq \alpha \leq 1$, 0 – extreme of pessimism and 1 - extreme of optimism.

- Criterion of Savage – recalculation result after flight:

$$A^* = \min_{B_j} \max_{A_i} r_{ij}(A_i, B_j), \quad (5)$$

where

r_{ij} – loss matrix for recalculations after DM:

$$r_{ij}(A_i, B_j) = \Delta = \max_{B_k} u_{ij}(A_i, B_k) - u_{ij}(A_i, B_j) \quad (6)$$

6. Integration of the deterministic, stochastic uncertainty (DM in Risk) and non-stochastic uncertainty (DM in uncertainty) models and determining optimal solution.

The flight pattern of an aircraft, the occurrence of an emergency (for example, when approaching the destination airfield, weather conditions have changed, bad weather conditions (BWC) is shown in Figure 4 with the following output data:

- Flight – performed for the first time
- Aircraft – Boeing -737
- Aerodromes – take-off (A_1), landing (A_2) and alternate aerodromes (A_2 ; A_3 ; A_4)
- Emergency – BWC (lightning strike and thunderstorm activity)
- The place of the event in-flight – emergency before approach

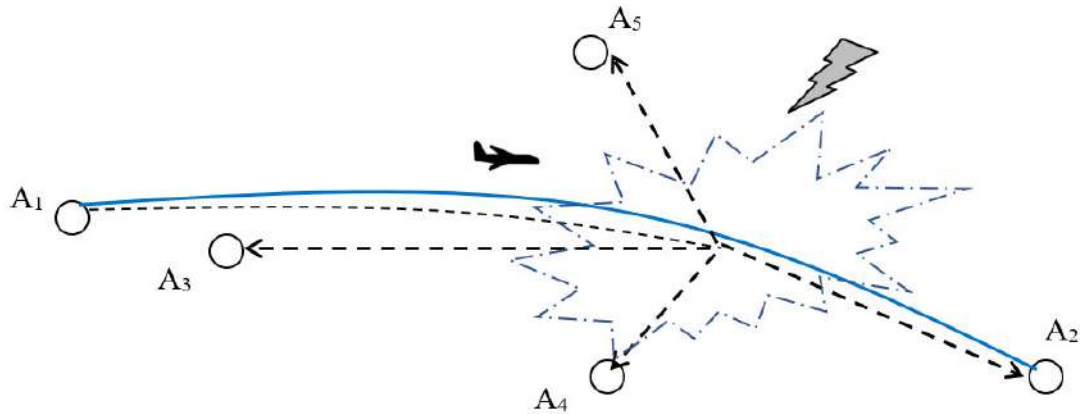


Figure 4: The scheme of aircraft flight route with possible alternative aerodromes

The decision for the selection of the alternate aerodrome in an emergency was implemented with the following output data (Table 3):

- The decision-making matrix
- Alternative actions - $\{A\} = \{A_1, A_2, \dots, A_i, \dots, A_m\}$ of DM H-O in ES and aerodromes of take-off (A_1), landing (A_2) and alternate aerodromes ($A_3; A_4$)
- States of nature or factors $\{A\} = \{\lambda_1, \lambda_2, \dots, \lambda_j, \dots, \lambda_n\}$ – variable or conditionally static factors $\{A\}$ and outcomes $U_s(A)$ influence on DM in ES. For example, DM in BWC (lightning strike and thunderstorm activity); “variable factor” - are meteorological conditions on aerodromes; “static factor” - available approach systems and aerodrome, crew, aircraft capabilities
- Outcomes of DM matrix $\{U\} = u_{11}, u_{12}, \dots, u_{ij}, \dots, u_{nm}$.
- Conditions of DM under uncertainty if this flight is performed for the first time - Criterion Wald (maximin)

Table 3

Matrix of DM in uncertainty for solution optimal aerodrome in BWC

Alternative aerodromes	Meteorological situation	Distance	Fuel reserve	Aerodrome capability	Crew capability	Aircraft capabilities
	λ_1	λ_2	λ_3	λ_4	λ_5	λ_6
Take-off aerodrome A_1	u_{11}	u_{12}	u_{13}	u_{14}	u_{15}	u_{16}
Landing aerodrome A_2	u_{21}	u_{22}	u_{23}	u_{24}	u_{25}	u_{26}
Alternate aerodrome A_3	u_{31}	u_{32}	u_{33}	u_{34}	u_{35}	u_{36}
Alternate aerodrome A_4	u_{41}	u_{42}	u_{43}	u_{44}	u_{45}	u_{46}
Alternate aerodrome A_5	u_{51}	u_{52}	u_{53}	u_{54}	u_{55}	u_{56}

The expected outcomes of the decision matrix are formed based on the influence of factors on DM, and the requirements of regulatory documents [6 – 9; 14 - 17], according to the Aeronautical Information Publication (AIP) too. It is proposed to evaluate the outcomes of the decision matrix using the Expert Judgment Method (EJM) [4; 19]. Each participant in the event - an aviation specialist fills in the decision matrix based on personal opinion. All matrices have the same factors that influence DM.

2.1. The Algorithm of the Collaborative Decision-Making in conflict / emergency situation

The Algorithm of the CDM in conflict/emergency was obtained using the methods of DM under uncertainty for effective collective solutions of different aviation specialists in an emergency:

1. Calculation of route direction
2. Building of DM matrix with:
 - Alternative solutions $\{A\}$
 - Factors, influencing on DM $\{A\}$
 - Outcomes of choosing of alternative solutions caused by factors influencing on DM $\{U\}$
3. Alternative solutions $\{A\}$ - the list of alternate aerodromes (AA):

$$\{A\} = \{A_{Dest} \cup A_{Dep} \cup \{AA\}\} = \{A_1, A_2, \dots, A_i, \dots, A_n\},$$

where

alternate aerodrome - an aerodrome of departure ($A_{Dep} - A_1$) and it's characteristics;

alternate aerodrome - an aerodrome of destination ($A_{Dest} - A_2$) and it's characteristics;

other alternate aerodromes and it's characteristics according to the calculated route - $A_3; A_4; A_5, \dots$

4. Factors $\{A\}$ influencing on DM for each operator (O_1 - PIC, O_2 - ATCO, O_3 - FD, and O_i - other aviation specialists). These factors may be original or identical and objective (Table 3). For example, next factors:

$$\{A\} = \lambda_1, \lambda_2, \dots, \lambda_j, \dots, \lambda_m,$$

where

λ_1 – fuel reserve on board;

λ_2 – remoteness of alternate aerodrome;

λ_3 – runway technical characteristics;

λ_4 – meteorological conditions on alternate aerodromes;

λ_5 – the approach lighting system;

λ_6 – available approach system;

λ_7 – available navigation aids;

λ_8 – aircraft performance characteristics;

λ_9 – the presence of radiocommunication;

λ_{10} – air traffic intensity, etc.

5. Outcomes $\{U\}$ - formation of possible consequences $\{U\}$ influencing on the selection of AA in the case of emergency:

$$\{U\} = U_{11}, U_{12}, \dots, U_{ij}, \dots, U_{nm},$$

where

$\{U\}$ - set of outcomes of DM matrix $U_{ij} (i=1, \dots, m; j=1, \dots, n)$.

The possible consequences U_{ij} are defined by means of using the EJM according to data from the regulatory documentation and opinions of O_i operators (PIC, FD, and ATCO and other aviation specialists) [4; 6 – 9; 14 - 17].

5. Formation of the matrixes of solutions for each operator. Formation of the matrix 1 of solutions for the first operator (O_1 - PIC) (Table 4).

Table 4

The DM matrix of DM in Uncertainty for operator O_1

The matrix 1		Factors influencing on DM for operator O_1 - PIC					
	$\{A\}$	λ_1	λ_2	...	λ_j	...	λ_n
Alternative actions in critical situation	A_1	u_{11}	u_{12}	...	u_{1j}	...	u_{1n}
	A_2	u_{21}	u_{22}	...	u_{2j}	...	u_{2n}
	...	...	...	...	...	...	...
	A_i	u_{i1}	u_{i2}	...	u_{ij}	...	u_{in}
	...	...	...	...	...	...	...
	A_m	u_{m1}	u_{m2}	...	u_{mj}	...	u_{mn}

Analogically, DM matrix for the second operator (O_2 - ATCO); the third operator (O_3 - FD) and other operators who are involved in this situation (Figure 1).

6. To choose the methods of DM under uncertainty (equations (2) – (6)) considering the conditions of DM under uncertainty (type of flight, emergency, conditions of development of situation).

7. Finding optimal solutions for each operator using the criteria of DM under uncertainty: Wald, Laplace, Savage, Hurwitz (equations (2) – (6)):

- $A_1^* = A_j(O_1)$ - solutions of PIC $A(O_1)$,
- $A_2^* = A_j(O_2)$ - solutions ATCO,
- $A_3^* = A_j(O_3)$ - solutions FD.

Analogically, DM for other aviation specialists.

8. Formation of the collective matrix of solutions (Table 5), where:

- $\{A\}$ – alternate aerodromes;
- $\{\lambda\}$ – optimal opinions of all operators (O_1 - PIC, O_2 - ATCO, O_3 - FD, and O_i - other aviation specialists) from individual matrixes.
- $\{u\}$ – outcomes - optimal decisions of operators in accordance with the selected criterion / flight conditions from individual matrixes $A_j(O_1)$; $A_j(O_2)$; $A_j(O_3)$

Table 5

The DM matrix of DM in Uncertainty for operators

The collective matrix	Results of optimal solutions by all operators						
	$\{A\}$	$A_j(O_1)$	$A_j(O_2)$	$A_j(O_3)$	$A_j(O_j)$	...	$A_n(O_n)$
Alternate aerodromes	A_1	u_{11}^*	u_{12}^*	u_{13}^*		...	u_{1n}^*
	A_2	u_{21}^*	u_{22}^*	u_{23}^*		...	u_{2n}^*
	...	...	...	...	...	...	...
	A_i	u_{i1}^*	u_{i2}^*	u_{i3}^*	u_{ij}^*	...	u_{in}^*
	...	...	...	...	...	...	...
	A_m	u_{m1}^*	u_{m2}^*	u_{m3}^*		...	u_{mn}^*

9. Finding of optimal solutions for all operators using the criteria of DM under uncertainty: Wald, Laplace, Savage, Hurwitz. For example, for criteria of Wald (2):

$$A^* = \max_{A_i} \left\{ \min_{B_j} u_{ij} \left(\min A_i, A_j(O_j) \right) \right\}$$

where

A_i – alternative solution from set $\{A\}$;

$A_j(O_j)$ – factors $\{A\}$ – opinions of operators from individual matrixes.

The factors in CDM matrix are objective. For each case, depending on the conditions of the situation and priorities of DM, a specific criterion has chosen.

2.1.1. The illustrative example of the Collaborative Decision-Making in emergency situation

Lightning strikes may affect airline operations because of costly delays and serious disruptions to flight schedules [20; 21]. The Aviation Safety Network safety database gives the following lightning strike accident statistics [22]: 29 occurrences with the contributory cause of the lightning strike, including 14 losses of control. A lot of jet airplane lightning strikes occur while in clouds, during the climb and descent phases of flight than in any other flight phase. The reason is that lightning activity is more prevalent between 5,000 to 15,000 feet altitude. The probability of a lightning strike decreases significantly above 20,000 feet. That's why airplanes flying short routes in areas with a high incidence of lightning activity are likely to be struck more often than long-haul airplanes operating in more benign lightning environments [22; 23; 24; 25].

A lightning strike may result in [25]:

- communication failure
- electrical problems
- emergency descends
- pilot incapacitation.

Expect events:

- pilot may be blinded by a lightning strike
- navigation system problems.

The situation considered in this work: bad weather conditions - lightning strike and thunderstorm activity closely to the destination aerodrome.

Initial data:

1. Airplane: Boeing 737
2. Route (Figure 5): Kyiv (Boryspil) (A_1) - Odessa (A_2).
3. Actual and forecasted weather conditions at the time of arrival to Odessa - thunderstorm activity in the vicinity of aerodrome, heavy shower rain, hail.
4. Alternate aerodromes [19]:
 - Chisinau (A_4);
 - Kharkiv (A_5);
 - Alternate aerodrome, which can be chosen along the route: Dnipro (A_3).
5. Factors influencing on DM for each operator:
 - $\{F\}$ - factors considered by operator O_1 (PIC);
 - $\{L\}$ - factors considered by operator O_2 (ATCO);
 - $\{\Lambda\}$ - factors considered by operator O_3 (FD).

Considering the interaction of the PIC, FD, and ATCO, it is necessary to understand their roles and responsibilities. The FD is responsible for the planning of the flight, the PIC is directly responsible for the safe execution of this flight, ATCO provides ATC service, information, and assistance to the crews. Therefore, when making a decision, they analyze a general group of factors, cause the main purpose is the completion of the task (flight from point A to point B), including risk analysis, but from different points of view. And the final decision in flight makes the PIC.

For rational CDM, each operator involved in DM has analyzed and considered the current situation. There are 3 operators in the CDM process: PIC (O_1), ATCO (O_2), FD (O_3) (Figure 1, Table 1). Each operator composed a matrix of decisions, where alternative solutions are alternate aerodromes for the route Kyiv - Odesa, and each operator considers the same factors in the current situation, but with different priorities.

The common factors for each operator (f_j, l_j, λ_j) taken into consideration when making a decision when approaching the destination aerodrome:

- f_1, l_1, λ_1 – fuel reserve on board of the aircraft (always controlled);
- f_2, l_2, λ_2 – meteorological situation (of departure, destination, alternate aerodromes, en-route);
- f_3, l_3, λ_3 – aircraft capabilities (available equipment on board, MEL peculiarities, existing operating limitations);
- f_4, l_4, λ_4 – aerodrome capabilities (available approach systems, technical characteristics of the runways and taxiways, lightning system, service hours restrictions, aerodrome category, firefighting and search and rescue category, emergency service);
- f_5, l_5, λ_5 – crew capability (crew operating minima, crew duty time);
- f_6, l_6, λ_6 – location of obstacles in approach, missed approach and departure sectors;
- f_7, l_7, λ_7 – air situation (intensity of air traffic control sector, radio frequency overload);
- f_8, l_8, λ_8 – commercial point (airport charges, distance from destination aerodrome, passenger and cargo service facilities, the presence of contracts with handlers, the presence of customs, border and migration control service, etc.).

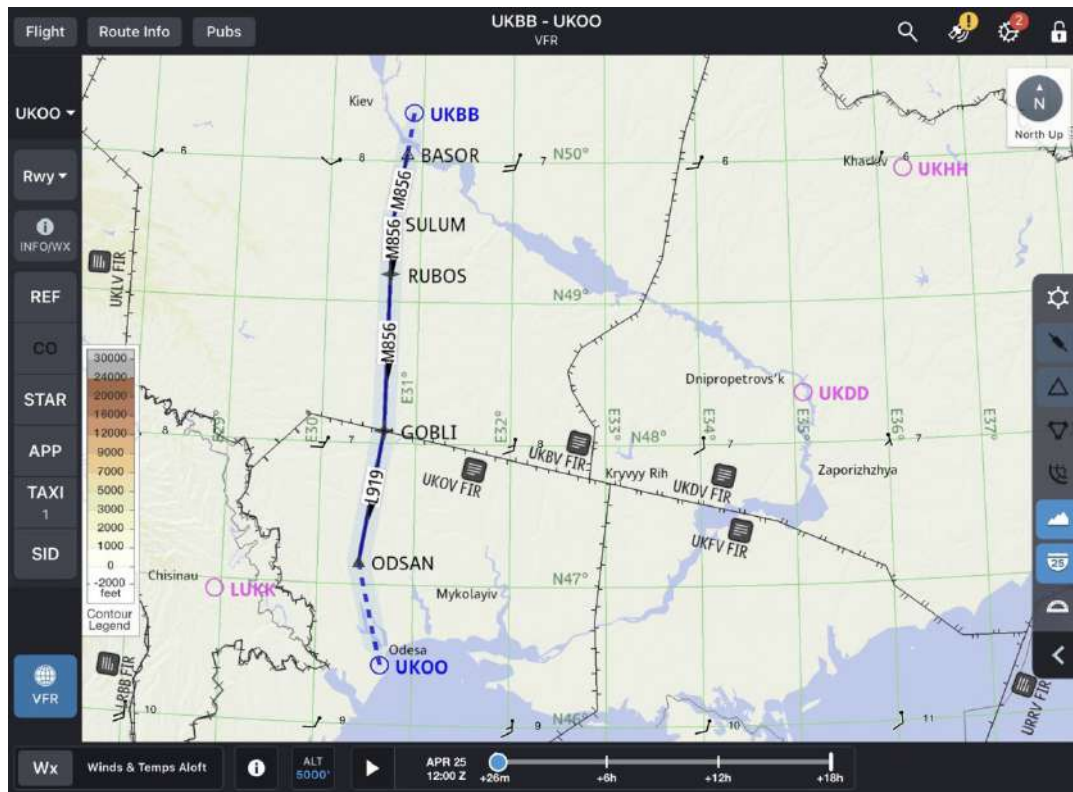


Figure 5: The route of flight Kyiv (Boryspil) - Odessa (UKBB - Kyiv (Boryspil); UKOO – Odesa; LUKK – Chisinau; UKHH - Kharkiv (Osnova); UKDD- Dnipro)

These factors are objective. Composed the operators' DM matrixes when approaching to the destination aerodrome in BWC (Tables 6-9). Expected outcomes considered by the PIC (operator O_1) represented in Table 6. The optimal solutions were determined according to the criteria Wald, Laplace, Horvitz, Savage (2) - (5). The solution according to the criterion of Savage was obtained from the loss matrix (6).

Table 6
The DM matrix of DM in Uncertainty for operator O_1 (PIC)

The matrix 1		Factors $\{F\}$, influence DM for operator O_1 - PIC								Solutions			
Alternative decisions $\{A\}$		f_1	f_2	f_3	f_4	f_5	f_6	f_7	f_8	W	L	$H, \alpha=0,5$	S
Departure Destination	Kyiv (A_1)	4	10	10	10	8	10	9	9	4	8,8	7	6
	Odessa (A_2)	10	1	6	7	1	10	2	7	1	5,5	5,5	9
Alternate aerodromes	Dnipro (A_3)	5	9	6	7	7	9	6	7	5	7,0	7	4
	Chisinau (A_4)	6	7	8	9	8	9	7	9	6	7,9	7,5	3
	Kharkiv (A_5)	4	10	7	7	6	5	6	7	4	6,5	7	6

The optimal landing aerodrome during the approach on the route "Kyiv - Odessa" in accordance with the PIC's decision is as follows:

- by Wald criterion - Chisinau (A_4)
- by Laplace criterion - Kyiv (A_1)
- by the Hurwitz criterion - Chisinau (A_4)

- according to the Savage criterion - Chisinau (A_4)
- Expected outcomes considered by the ATCO (operator O_2) represented in Table 7.

Table 7
The DM matrix in Uncertainty for operator O_2 (ATCO)

The matrix 2		Factors $\{L\}$, influencing on decision-making of ATCO (O_2)								Solutions			
Set of alternative decisions $\{A\}$		l_1	l_2	l_3	l_4	l_5	l_6	l_7	l_8	W	L	H, $\alpha=0,5$	S
Departure	<i>Kyiv (A_1)</i>	3	10	10	10	9	10	5	7	3	8,0	6,5	7
Destination	<i>Odessa (A_2)</i>	10	1	10	10	1	10	8	10	1	7,5	5,5	9
	<i>Dnipro (A_3)</i>	4	8	10	10	6	8	6	5	4	7,1	7	6
Alternate aerodromes	<i>Chisinau (A_4)</i>	8	5	10	10	7	7	7	9	5	7,9	7,5	5
	<i>Kharkiv (A_5)</i>	8	4	10	10	7	7	4	8	4	7,3	7	6

The optimal landing aerodrome during the approach on the route "Kyiv - Odessa" in accordance with the ATCO's decision as follows:

- by Wald criterion - Chisinau (A_4)
- by Laplace criterion - Kyiv (A_1)
- by the Hurwitz criterion - Chisinau (A_4)
- according to the Savage criterion - Chisinau (A_4)

Evaluation of optimal alternate aerodrome for landing in difficult meteorological conditions by FD at the stage of flight planning. Matrix of possible outcomes of DM by FD during choosing of alternate aerodromes at the stage of flight planning represented in Table 8.

Table 8
The DM matrix for operator O_3 (FD)

The matrix 3		Factors $\{A\}$, influencing on decision-making of FD (O_3)								Solutions			
Alternative decisions $\{A\}$		λ_1	λ_2	λ_3	λ_4	λ_5	λ_6	λ_7	λ_8	W	L	H, $\alpha=0,5$	S
Departure	<i>Kyiv (A_1)</i>	3	10	10	10	7	10	7	5	3	7,8	6,5	7
Destination	<i>Odessa (A_2)</i>	10	1	7	7	1	10	3	10	1	6,1	5,5	9
	<i>Dnipro (A_3)</i>	3	8	6	5	5	8	5	8	3	6,0	5,5	5
Alternate aerodromes	<i>Chisinau (A_4)</i>	8	9	8	7	5	5	6	9	5	7,1	7	4
	<i>Kharkiv (A_5)</i>	6	4	9	8	6	8	4	10	4	6,9	7	6

The optimal landing aerodrome during the approach on the route "Kyiv - Odessa" in accordance with the FD's decision is as follows:

- by Wald criterion - Chisinau (A_4)
- by Laplace criterion - Kyiv (A_1)
- by the Hurwitz criterion - Chisinau (A_4) and Kharkiv (A_5)

- according to the Savage criterion - *Dnipro* (A_3)

To determine the consistency of operators, collective matrices were constructed, in which the factors in the decision matrixes for the operators (PIC (O_1), ATCO (O_2), FD (O_3)) are identical, the solutions of the operators and taken from matrices, presented in Tables 6; 7 and 8. In the CDM matrixes used subjective factors – opinions of operators. The optimal CDM if this flight is performed for the first time is presented in the Table 9 (Wald criterion). In this case, the optimal landing aerodrome, determined by objective (fuel reserve on board of the aircraft; meteorological situation; crew, aircraft and aerodrome capabilities; location of obstacles in approach; air situation and commercial point) and subjective factors (PIC, FD, and ATCO) is alternate aerodrome Chisinau (A_4).

Table 9

The CDM matrix for all operators (criteria Wald)

	PIC	ATCO	FD	CDM
Alternate aerodromes	O_1	O_2	O_3	Criterion Wald
<i>Kyiv</i> (A_1)	4	1	5	1
<i>Odessa</i> (A_2)	1	1	1	1
<i>Dnipro</i> (A_3)	5	4	5	4
<i>Chisinau</i> (A_4)	6	5	5	5
<i>Kharkiv</i> (A_5)	4	4	4	4

The optimal CDM if this flight is regular (Criterion of Laplace) presented in the Table 10 – is *Kyiv* (A_1).

Table 10

The CDM matrix for all operators (criteria Laplace)

	PIC	ATCO	FD	CDM
Alternate aerodromes	O_1	O_2	O_3	Laplace
<i>Kyiv</i> (A_1)	8,8	7,8	8,0	8,2
<i>Odessa</i> (A_2)	5,5	6,1	6,1	5,9
<i>Dnipro</i> (A_3)	7,0	6,0	6,0	6,3
<i>Chisinau</i> (A_4)	7,9	7,1	6,8	7,3
<i>Kharkiv</i> (A_5)	6,5	6,9	6,9	6,8

The optimal CDM in different approaches using optimism-pessimism coefficient $\alpha=0,5$ (criterion of Hurwitz) presented in Table 11 – is *Dnipro* (A_3).

Table 11

The CDM matrix for all operators (criteria Hurwitz)

	PIC	ATCO	FD	CDM
Alternate aerodromes	O_1	O_2	O_3	Hurwitz
<i>Kyiv</i> (A_1)	5,5	5,5	5,5	5,5
<i>Odessa</i> (A_2)	7	7	5,5	6,3
<i>Dnipro</i> (A_3)	7,5	7,5	7	7,3
<i>Chisinau</i> (A_4)	7	7	7	7,0
<i>Kharkiv</i> (A_5)	5,5	5,5	5,5	5,5

Collective decisions were analyzed with varying degrees of optimism-pessimism coefficient $\alpha=0,5$. The consistency of decisions increases with an increase in the coefficient of optimism, with a decrease in the coefficient in the direction of pessimism, the mismatch increases.

The consistency of decisions determined using the Savage criterion (the recalculation after a flight), were determined for cost initial matrix (Table 12):

$$A^* = \min_{B_j} \max_{A_i} r_{ij}(A_i, B_j) = \min_{B_j} (3; 6; 1; 0; 3) = 0 = A_4,$$

where

$$r_{ij}(A_i, B_j) = \Delta_{A_i} = u_{ij}(A_i, B_j) - \min_{B_k} u_{ij}(A_i, B_k) = (3; 6; 1; 0; 3)$$

Table 12

The CDM matrix for all operators (criteria Savage) – recalculation)

	PIC	ATCO	FD	CDM
Alternate aerodromes	O_1	O_2	O_3	Savage
<i>Kyiv (A₁)</i>	6	7	7	3
<i>Odessa (A₂)</i>	9	9	9	6
<i>Dnipro (A₃)</i>	4	6	5	1
<i>Chisinau (A₄)</i>	3	5	4	0
<i>Kharkiv (A₅)</i>	6	6	6	3

The loss matrix presents on the Table 13. The loss matrix shows risks if operators do not choose the optimal collective solution. The maximum risks are selected, which then are minimized.

Table 13

The loss CDM matrix for all operators (criteria Savage))

	PIC / loss	ATCO/ loss	FD / loss	Max loss
Alternate aerodromes	O_1	O_2	O_3	Savage
<i>Kyiv (A₁)</i>	3	2	3	3
<i>Odessa (A₂)</i>	6	4	5	6
<i>Dnipro (A₃)</i>	1	1	1	1
<i>Chisinau (A₄)</i>	0	0	0	0
<i>Kharkiv (A₅)</i>	3	1	2	3

The optimal landing aerodrome, determined by objective and subjective factors is alternate aerodrome Chisinau (A_4) as in Wald criterion. The calculations showed a balance between safety and cost of the flight using criteria Wald (maximum safety) and criteria Savage (minimal cost and loss).

3. Acknowledgements

Thanks to participants of the scientific research - professional aviation specialists such as pilots and dispatches from airline.

To demonstrate the new method of Integration of Decision-Making Stochastic Models and CDM used the individual science works of aviation students in education (courses “Informatic of DM in ANS” in National Aviation University, Kyiv). Students of different qualifications (PIC, FD, and ATCO, operators of Unmanned Air Vehicles, engineers, and technical personnel) study theoretical courses and then use their knowledge for scientific research. In the chapter "Collaborative Decision Making in Emergencies by the Integration of Deterministic, Stochastic, and Non-Stochastic Models" [13] was presented example of the emergency situation "Lightning strike" from a master's diploma of Zhanna Maksymchuk's “Dynamic models of air traffic controller decision-making in difficult meteorological conditions "Lightning strike" [25]. The diploma considers an example of choosing the optimal landing aerodrome for the ATCO. The authors proposed to find joint solutions for the PIC, FD, and ATCO in an emergency situation.

4. Conclusion

Collaborative Decision-Making is a process of presenting individual and collaborative information and decisions made by various interacting participants, such as PIC, FD, and ATCO in professional solutions, providing synchronization of decisions taken by participants, the exchange of information between them, the effective balancing between safety and cost in collective solutions. It is important to ensure the possibility of making a joint, integrated solution by aviation specialists at an acceptable level of efficiency with minimal risk and maximum safety [5]. This is achieved by the completeness and accuracy of the available information, and by the well-coordinated interaction between specialists, their clear and correct understanding of job duties, and their roles in the process of completing a common task. A Collaborative Decision-Making process in uncertain situations should be provided by using DM stochastic models (stochastic uncertainty and non-stochastic uncertainty models) to improve and simplify the deterministic model.

The objective-subjective approach is effective for determining the optimal solution in important and difficult situations and in conflict interactions between operators. After analysis of the situation firstly the performance of the synthesis (aggregation) of individual models (with objective factors), the next step is to determine collective solutions (with subjective factors). The example of choosing an aerodrome in case of an emergency (for example, in difficult meteorological conditions) using the methods of DM under uncertainty was presented. The calculations showed a balance between safety and cost using criteria Wald (maximum safety) and criteria Savage (minimal cost and loss).

The direction of further research is to develop the DM models for all CDM participants within the Airport CDM (A-CDM) concept may aggregate the solutions of partners (airport operators, manned aircraft operators, unmanned aircraft operators, ground handling agents, and air traffic services) in joint work within non-segregated airspace [26]. For more accurate and timely information provision all partners at the airport are required to create a single operational picture of air traffic [15; 17]. The research was carried out during the simulator training of air traffic controllers [27], it is planned to build the same models for other aviation specialists, operators of manned and unmanned aircraft, flights in integrated airspace [28]. Specialists from other fields, for example, emergency and ground services, and medical staff may be involved in joint DM. Priorities and effectiveness of connecting specialists in the group for CDM will be in the next research. For effective CDM is proposed to include the participant (Artificial Intelligence) in the group of participants in a joint solution and organize the research of a hybrid optimal solution.

5. References

- [1] Air safety statistics in European Union (EU), 2021. URL: https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Air_safety_statistics_in_the_EU
- [2] EASA Annual Safety review 2020. European Union Aviation Safety Agency, 2020. URL: <https://www.easa.europa.eu/downloads/117065/en>
- [3] A. Chialastri, AF 447 as a Paradigmatic Accident: The Role of Automation on a Modern Airplane, IGI Global, USA, Pennsylvania, 2021. doi:10.4018/978-1-5225-7709-6.ch006
- [4] T. Shmelova, Yu. Sikirda, N. Rizun, Abdel-Badeeh M. Salem, Yu. N. Kovalyov (Eds.), Socio-Technical Decision Support in Air Navigation Systems: Emerging Research and Opportunities, IGI Global, USA, Pennsylvania, 2017. doi: 10.4018/978-1-5225-3108-1
- [5] Safety Management Manual, 3rd ed. Doc. ICAO 9859-AN 474. ICAO, Canada, Montreal, 2013.
- [6] Flight Operations Officer/Flight Dispatcher Training Manual, Doc ICAO 7192. ICAO, Canada, Montreal, 1998.
- [7] Procedure of making decision of departure and arrival of aircraft of Ukraine civil aviation according instrument flight rules: State aviation authority of Ukraine Order No 295, 2005.
- [8] Air Traffic Management. Doc. 4444-RAC/501. Fifteenth ed. ICAO, Canada, Montreal, 2007
- [9] Guidelines for Controller Training in the Handling of Unusual/Emergency Situations (ASSIST). – EUROCONTROL, Brussels, Belgium, 2003.

- [10] Abdel-Badeeh Salem (Ed.), *Innovative Smart Healthcare and Bio-Medical Systems: AI, Intelligent Computing and Connected Technologies*. 1st ed. CRC Press, 1st edition: December 28, 2020, ISBN 9780367490614
- [11] Emergency telemedicine and mission safety support, 2022. URL: <https://www.aircareinternational.com/inflight-emergency-telemedicine>
- [12] T. Shmelova, Yu. Sikirda, C. Scarponi, A. Chialastri, Deterministic and Stochastic Models of Decision Making in Air Navigation Socio-Technical System, in *Proceedings of the 14th International Conference on ICT in Education, Research and Industrial Applications. Integration, Harmonization and Knowledge Transfer*, Kyiv, Ukraine 2018. URL: http://ceur-ws.org/Vol-2104/paper_221.pdf
- [13] T. Shmelova, Collaborative Decision Making in Emergencies by the Integration of Deterministic, Stochastic, and Non-Stochastic Model, in: Jon W. Beard (Ed.), *Information Technology Applications for Crisis Response and Management*, IGI Global, USA, Pennsylvania. 2021, doi: 10.4018/978-1-7998-7210-8.ch010
- [14] Global Air Navigation Plan 2016-2030. Doc. ICAO 9750. Fifth ed., ICAO, Canada, Montreal, 2016.
- [15] Manual on Collaborative Decision-Making (CDM). Doc. 9971. Second ed. ICAO, Canada, Montreal, 2014.
- [16] Manual on Global Performance of the Air Navigation System (PBA). Doc. 9883. First ed. ICAO, Canada, Montreal, 2009.
- [17] Manual on Flight and Flow Information for a Collaborative Environment (FF-ICE). Doc. 9965. First ed., ICAO Canada, Montreal: ICAO. 2012.
- [18] Yu. Sikirda, M. Kasatkin, D. Tkachenko, Intelligent Automated System for Supporting the Collaborative Decision Making by Operators of the Air Navigation System during Flight Emergencies, in: T. Shmelova, Yu. Sikirda, A. Sterenharz (Eds.), *Handbook of Artificial Intelligence Applications in the Aviation and Aerospace Industries*, IGI Global, USA, Pennsylvania. 2021, doi: 10.4018/978-1-7998-1415-3.ch003
- [19] Hamdy A. Taha, *Operations Research: An Introduction*, Global Edition, 10th Ed., University of Arkansas, 2019, ISBN-13: 9781292165554
- [20] Rules of information provision of flight safety management system of aircraft of civil aviation of Ukraine. Ministry of transport and communications Order No 295. 2009.
- [21] Operation of Aircraft. ICAO Annex 6. ICAO, Canada, Montreal, 2018.
- [22] The Aviation Safety Network, an exclusive service of Flight Safety Foundation, 2022. URL: <https://aviation-safety.net/>
- [23] G. Sweers, B. Birch, J. Gokcen, Lightning Strikes: Protection, Inspection, and Repair, Boeing AERO Magazine. QTR. 2022. URL: https://www.boeing.com/commercial/aeromagazine/articles/2012_q4/4/
- [24] SkyBrary, 2022. URL: <https://www.skybrary.aero/index.php/Lightning>
- [25] Zh. Maksymchuk, Dynamic models of air traffic controller decision-making in difficult meteorological conditions "Lightning strike", Master's thesis, National Aviation University. Kyiv, Ukraine, 2020
- [26] Airport collaborative decision-making, 2022. URL: <https://www.eurocontrol.int/concept/airport-collaborative-decision-making>
- [27] T. Shmelova, Yu. Sikirda, T. Jafarzade, Artificial Neural Network for Pre-Simulation Training of Air Traffic Controller, in: Mehdi Khosrow-Pour (Ed.), *Research Anthology on Artificial Neural Network Applications (Volumes 3)*, IGI Global, USA, Pennsylvania, 2022. doi: 10.4018/978-1-6684-2408-7.ch065
- [28] T. Shmelova, V. Lazorenko, O. Burlaka, Unmanned Aerial Vehicles for Smart Cities: Estimations of Urban Locality for Optimization Flight in: José António Tenedório (Ed.), *Methods and Applications of Geospatial Technology in Sustainable Urbanism (Chapter 15)*, IGI Global, USA, Pennsylvania, 2021. doi: 10.4018/978-1-7998-2249-3.ch015

Stability Verification of Self-Organized Wireless Networks with Block Encryption

Volodymyr Sokolov^a, Pavlo Skladannyi^a, and Hennadii Hulak^a

^a *Borys Grinchenko Kyiv University, 18/2 Bulvarno-Kudriavska str., Kyiv, 04053, Ukraine*

Abstract

Sensor networks allow the transmission of data for medical and research purposes. They are essential for building IoT wireless networks. The widespread use of such networks requires simplification of their construction and administration. The easiest way to develop and scale such a network is to use self-organizing networks. Since IoT networks operate in the available ISM frequency range, it constantly requires monitoring the load on the air. The paper shows the process of selecting free wireless channels using spectrum analysis. After the channel is selected, the throughput of the self-organizing system (botnets) is analyzed. The architecture and analysis of research results for botnets based on available hardware are shown. In addition, the issue of secure data transmission using fast and simple XTEA block encryption is considered.

Keywords

Spectrum analysis, bandwidth analysis, self-organizing network, wireless, botnets, IoT, block encryption, XTEA.

1. Introduction

The development of network technologies is increasing, affecting the number of digital devices connected to the network, including IoT. With this rapid development, there is a need to develop effective data exchange protocols and devices that will provide this exchange because the standard protocols used in traditional networks can not fully meet the needs of a new type of network, which became known as sensory. These networks should self-organize, as setting up a network manually on each new device is quite problematic. In addition, they should automatically repair broken connections and find new routes when the network topology changes dynamically, as most of the devices that the IoT concept combines are pretty mobile, so it is not possible to secure them in the network [1]. In addition to ensuring the smooth operation of the network, you need to provide cryptographic protection of data transmission [2]. Such encryption requires additional energy costs, which is critical for sensor networks and IoT in general. Thus, this paper considers the properties of such networks and ways to optimize them to increase the level of fault tolerance, data exchange rate, their protection, and performance control [3].

Sect. 2 analyzes the available research and highlights aspects that need further analysis. Spectrum and bandwidth analysis are provided in Sect. 3 and 4. Sect. 5 and 6 are devoted to the construction of the self-organizing network and their experimental study. Also, Sect. 7 uses block cipher for the self-organizing network. The final section outlines the directions for future research in this area.

2. Related Works

An ad hoc wireless peer-to-peer network is a decentralized wireless network [4, 5]. The network can be called “special” for a particular case (according to its name) because it does not rely on existing infrastructure, such as routers in wired networks or access points in the routed network [6].

CMIS-2022 : Fifth International Workshop on Computer Modeling and Intelligent Systems, May 12, 2022, Zaporizhzhia, Ukraine
EMAIL: v.sokolov@kubg.edu.ua (V. Sokolov); p.skladannyi@kubg.edu.ua (P. Skladannyi); h.hulak@kubg.edu.ua (H. Hulak)
ORCID: 0000-0002-9349-7946 (V. Sokolov); 0000-0002-7775-6039 (P. Skladannyi); 0000-0001-9131-9233 (H. Hulak)



© 2022 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

Instead, each node participates in routing by sending data to other nodes, thus determining which nodes the data will be transmitted dynamically based on the network connection and the routing algorithm used in a particular case.

The decentralized nature of wireless networks makes them suitable for various applications where central nodes cannot be laid and improve networks' scalability compared to wireless managed networks [7]. Minimal configuration and rapid deployment make peer-to-peer networks suitable for emergencies such as natural disasters or military conflicts. The presence of dynamic and adaptive routing protocols allows peer-to-peer networks to be quickly configured and modified [8].

Concerning the security of data transmitted by sensor networks, it should be noted that most ad hoc networks do not exercise any control over access to the network. These networks are vulnerable to consumer attacks, where a malicious node enters packets into the network. Source [9] prevent such attacks. It is necessary to use authentication mechanisms that ensure that only authorized nodes can enter traffic into the network [10]. However, it should be understood that even with authentication, these networks are vulnerable to de-authorization packets or retard attacks, resulting in the intermediate node rejecting the package or delaying it rather than immediately sending it to the next node.

Thus, based on the dependence of the principles of the network on its purpose, it is necessary to identify the requirements for the designed network before starting work on it to avoid situations that would require further changes to the overall architecture of the program.

3. Spectrum Analysis

Spectrum analyzers do not increase the availability of information in the network but contribute to its improvement by detecting weaknesses and filling the spectrum with other networks (collision prevention/interference). Unlike standard tools available in network cards, the spectrum analyzer collects information about noise levels, for example, from magnetrons (in microwave ovens), can detect stationary interference and the operation of other equipment (Bluetooth, ZigBee, children's toys, video transceivers, medical sensors, etc.) [11]. Also, the spectrum analyzer "sees" not only the service packets, which assess the signal level in the network cards but also all other packets. Therefore, the topic of the paper is essential and highly relevant. Previous versions of industrial spectrum analyzers TI CC25xx [12] and Cypress CYRF693x were used by the authors as sensors in the study of antenna technology [13, 14] and the construction of sensor networks [15].

Fig. 1a and 1b show the results of reception at the location of the spectrum analyzer in the near zone of the transmitter for data transmission over the Bluetooth wireless channel, and in Fig. 1c and 1d—Wi-Fi band 2.4 GHz. From the figures, it is easy to see that the solid assembly showed much better results since it is better assembled using the screen and external antenna [8].

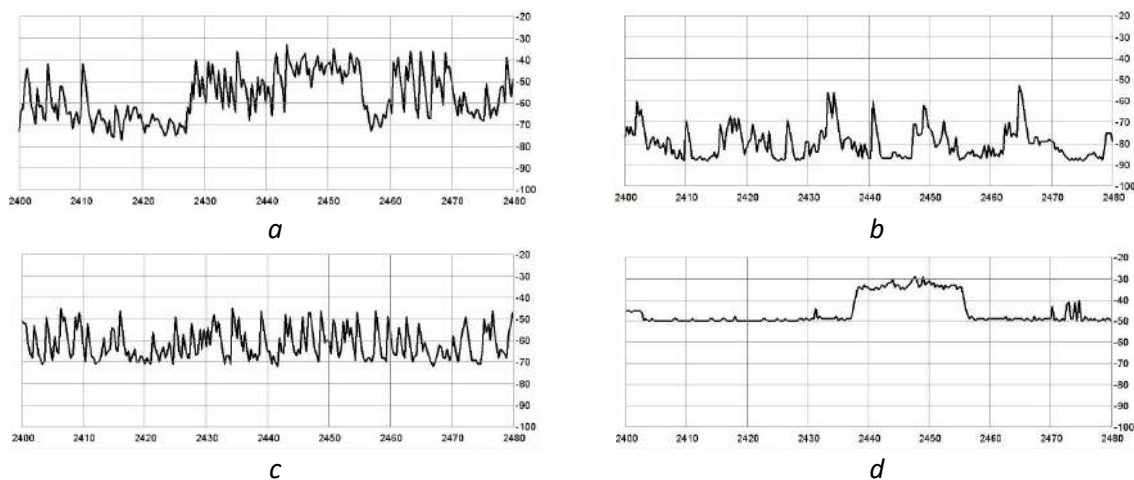


Figure 1: Examples of spectra for different devices: modular for Bluetooth (a); single-board for Bluetooth (b); modular for Wi-Fi (c); single-board for Wi-Fi (d)

After assembly, testing, and debugging, the devices are ready for technical and scientific tasks. The general view of all three devices is shown in Fig. 2.



Figure 2: General view of different devices

The paper presents the results of the design and manufacture of spectrum analyzers on finished components (transceiver chips for the IEEE 802.15.4/ZigBee standard) [16]. Designing, manufacturing printed circuit boards, assembling devices, and programming microcontrollers are described in detail. Testing and improvement of existing devices. When wiring, the board revealed the dependence of the quality of the device on the quality of its assembly, the presence of an electromagnetic screen, and the type of antenna [8].

4. Bandwidth Analysis

For the purity of the experiment, both machines are configured identically with the same operating system and the same settings. Therefore, everything written as Raspberry Pi (RPi)—applies to RPi Zero and RPi 3 (Fig. 3). The operating system (OS) used an OS without a graphical interface—Raspbian OS Lite to have a minimum load on the system. The peculiarity of this OS is the absence of “extra” programs that slow down the computer. In this version of the OS, there is a package for working with GPIO. RPi act as a wireless Wi-Fi access point, configured by example from [9]. The access points have a *vsftpd* FTP server, which is easy to set up, fast, and secure. As a result, the average of all tests will be taken into account. Every five seconds, the readings are taken: the load and temperature of the server processor, the current and voltage consumed by the server, and the download speed of the file on the FTP client.

The second stage was the automation of the process of writing history files. This will minimize human error during the recording of results and thus increase the accuracy of measurements and recording speed. A digital INA219 sensor was used to measure voltage and current. In addition to the built-in processor temperature sensor, an external DS18B20 sensor is additionally used [17]. Both sensors are connected to the GPIO interface, which allows you to receive information in text format, which is used to automate measurements.

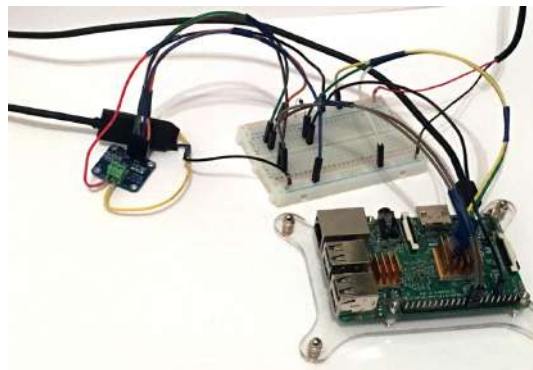


Figure 3: Test bench for Raspberry Pi 3

The *sensors.py* script is running on the Server-PC, which receives data from the DS18B20 and INA219 sensors. The *ina219* library is used to work with INA219. When the DS18B20 temperature sensor is connected, a file is automatically created in which the temperature is displayed in text format. Therefore, to get the ambient temperature in *sensors.py*, it is enough to open and expand this file using the built-in OS library in Python. The general scheme of operation of the *sensors.py* script is shown in Fig. 4.

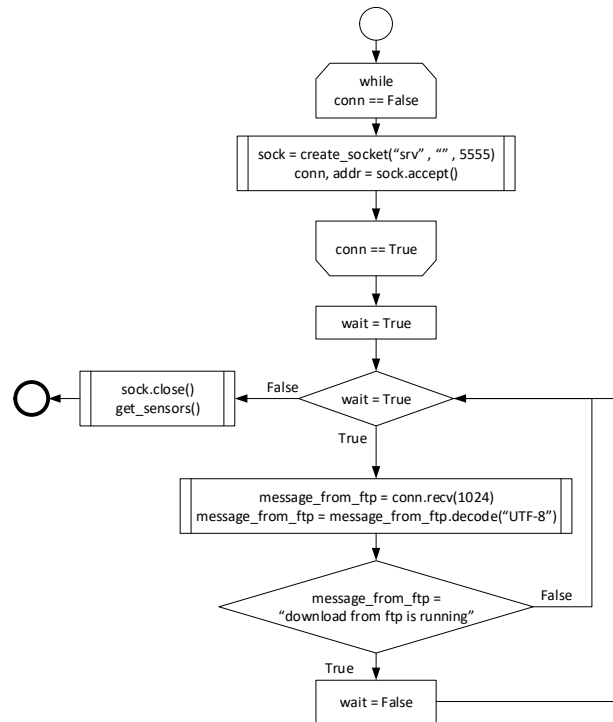


Figure 4: Sensor read algorithm

The processor on RPi 3 was less loaded than on RPi Zero, so that you can run more resources on RPi 3 than on RPi Zero. Interestingly, at a longer distance from the client from the access point, the CPU 3 CPU load was lower. RPi Zero at a longer length periodically decreased sharply but tended to the maximum value (Fig. 5).

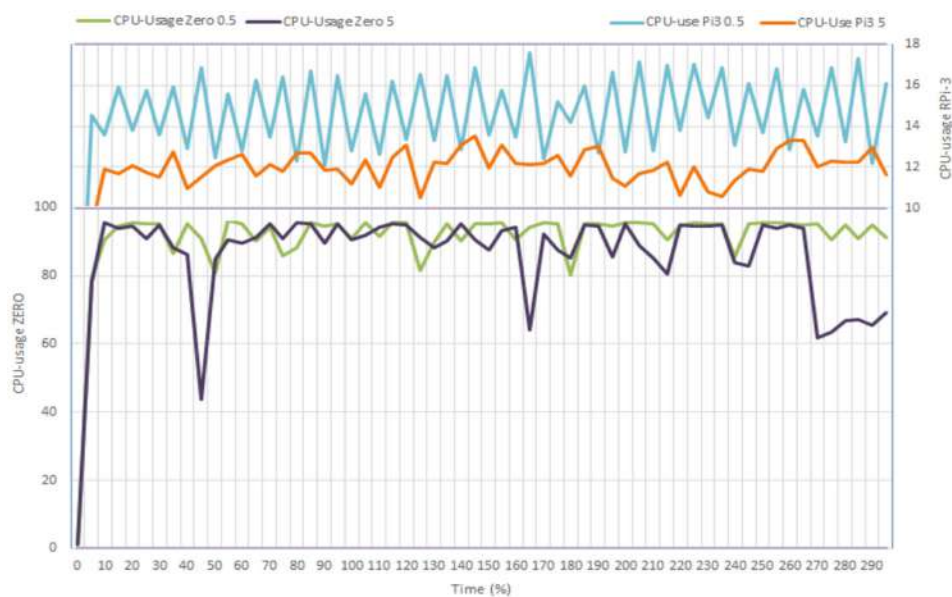


Figure 5: Comparison of RPi Zero and RPi 3 CPU usage

The injection rate decreased most markedly with a sharper decrease in current consumption. Current consumption increased sharply when necessary to reduce the CPU load to reduce the CPU temperature (Fig. 6).



Figure 6: Comparison of CPU usage, power consumption, and file download speed of RPi 3

Using Pearson's correlation, you can calculate the relationship between the factors:

$$r = \frac{n\sum xy - (\sum x)(\sum y)}{\sqrt{(n\sum x^2 - (\sum x)^2)(n\sum y^2 - (\sum y)^2)}} \quad (1)$$

where x is the value of one factor; y is the value of another factor to which the ratio refers.

Table 1 shows the results of calculating the Pearson correlation coefficient for different distances by expression (1). The correlation coefficient varies from +1 to -1. In the presence of a positive correlation, an increase in one indicator contributes to the rise. With a negative correlation, an increase in one indicator entails a decrease in another. The larger the modulus of the correlation coefficient, the more noticeable the change in one hand affects the difference in the other. At a factor of 0, the relationship between them is absent [18].

Table 1

The results of the calculation of the Pearson correlation coefficient for RPi Zero / RPi 3

	Speed	CPU load	CPU temp.	Temp.	Voltage	Amperage
Speed	1.00/1.00					
CPU load	0.55/0.38	1.00/1.00				
CPU temp	0.30/0.24	0.26/0.46	1.00/1.00			
Temp.	0.11/0.02	-0.02/0.05	0.87/0.19	1.00/1.00		
Voltage	-0.16/-0.04	0.06/-0.17	-0.18/-0.08	-0.18/-0.04	1.00/1.00	
Amperage	0.05/0.32	0.01/0.43	0.22/0.17	0.21/-0.07	-0.12/-0.10	1.00/1.00

This technical complex in RPi and the DS18B20 temperature sensor and the INA219 current, voltage, and power sensor can be integrated into such systems as monitoring systems in data centers, intelligent home systems, etc. This technical complex is already used for temperature monitoring servers in the data center. For visualization, monitoring, and analysis of the received data, the Grafana platform is shown in Fig. 7.



Figure 7: Grafana dashboard with temperature and voltage indicators

5. Self-Organizing Network

To test the capabilities of the ad hoc sensor network based on the IEEE 802.11 standard, we will try to build our network, which will be stolen from the modules on the ESP8266 microcontroller with a self-developed firmware [19].

Since the most widely used such networks are used in IoT solutions, we can conclude that the main function of modules in the network is to obtain a small amount of data from devices to which they are connected via serial interface, transferring this data over Wi-Fi between modules and return the result to the parent device via the serial interface. Accordingly, the system “module device” must both receive and process the received commands. Therefore, NodeMCU and ESP-12E were used to simulate these conditions for convenience.

As a result, we have approximately such a test bench necessary to begin developing and implementing the first tests, shown in Fig. 8.



Figure 8: Test benches with NodeMCU 1.0 (a) and ESP-01 (b) modules

The only disadvantage of this solution is the voltage given by this battery: up to 9 V. The ESP8266 can operate at voltages from 3 to 3.6 V. However, the problem can be solved with a suitable voltage regulator. In this case, the stabilizer LD1117 was used. You can also use similar stabilizers, such as AMS1117 or LM1117. As a result, the test device shown in the figure was assembled, which is entirely autonomous, so when implemented on an entire board, it can be used in conditions with difficult access, thus using the device to obtain data from instruments located at remote points, while without having to go to them.

The next step was to develop firmware for the assembled system. Already at the stage of developing the program architecture, some difficulties appeared. Since full-fledged ad hoc networks

typically run on multifunctional equipment designed directly to build networks (Wi-Fi access points and routers), which has significant computing capabilities and the ability to multithreaded data processing, the problem of how to implement the same functions on a microcontroller. Significantly inferior to its technical characteristics. Several open-source solutions have been considered, including the aforementioned ESP8266WiFiMesh library. As a result, it was found that all such solutions are implemented on the alternate change of modes of operation of the module, as shown in Fig. 9.

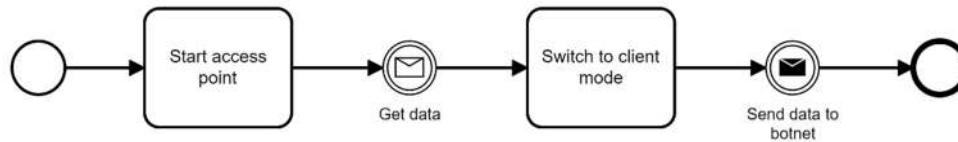


Figure 9: Mesh network operation algorithm

The exchange algorithm was implemented using the ESP8266Mesh library. Therefore the mode of operation of the module is controlled by the algorithm shown above, but with a small change, which was described above. At each stage of the cycle, the firmware listens for incoming connections. If one of the modules is connected to it processes the incoming message and sends back a response. The next step is to check the messages in the queue and, if available, switch to the client mode, send the first message from the line to related modules. Then the procedure is repeated.

Because the size of the queue can vary depending on the node load and the speed of sending messages, data is stored in dynamic memory. As a result, there were some difficulties in writing this part of the program. Due to a small error in the work with the pointers, the program suffered a fatal error, and the microcontroller rebooted in an emergency. The problem was related to a pointer that connected messages in a queue, causing the program to reference an already occupied area of memory. After correcting the error, the first working prototype of the program was obtained, thanks to which it was possible to implement data transfer between two test NodeMCU 1.0 modules [20].

6. Experimental Assessment of Self-Organizing Network

The experimental setup consists of a working network model, which includes three nodes (two NodeMCU boards from the test bench and a test device with the ESP-01 module).

The first test was the direct transfer of data between nodes and the node's received messages. For this purpose, two test messages were sent from node Node14754480, which were addressed to node Node52126. In the test message, a command was sent to turn on the LEDs on the Node52126 module. The module received the message in two ways: directly from the sender's module and through the "intermediary" Node1592748, which also received a message from Node14754480 and forwarded it. The second message was rejected as a duplicate. The system worked correctly. The following are the messages to the console from modules Node14754480 and Node52126. The message, in this case, was received directly from the sending module and processed. The red LED was lit. The message has reached the recipient. No further mailing is required. The system has switched to listening mode. The test was passed successfully.

The message transmission time was approximately measured. Depending on the conditions, the data transmission in this network can take from one to five seconds, depending on the message path and the response rate of the modules.

A test device based on an Arduino Nano board was also used for testing, to which the module must send a message for processing, after which the LEDs should light upon this device. This function works in the same way as when processing commands with LEDs on test modules with NodeMCU. The Arduino also generates a message for other modules every five seconds and transmits it to the ESP-01 module for further transmission. During testing, the problem with the incorrect reading of messages by the module was revealed. The reason for this was the special symbols provided by the transfer fee. After eliminating this problem, messages began to be sent correctly, and the board could communicate with other devices.

7. Implementation of XTEA Encryption Protocol

Based on the hardware and software features of Pololu Wixel (TI CC25xx microcontroller) devices [21, 22], the following problems were identified during the implementation of this project:

- Compilation of types.
- Buffer overflow.
- Work with pointers.
- Variables in XDATA memory.
- Looping with minimal delay.
- Compiler errors.

Casting types. Because the C programming language used in SDCC is hard-typed [23], there is a problem with typesetting. For example, the compiler forbids equating uint8 and char, although char also requires one byte. It is also impossible to equate specific arrays. Hence the need to compile data types between variables and intentionally specify a summary when passing a function parameter. Compilation of types, in most cases, does not give the desired result because the programming language does not involve the conversion of certain data types.

The user enters data into the client program; the program processes and sends this data to the sender via USB (Fig. 10). If necessary, the program sends the keys to the sender and receiver before sending the processed data. The sender encrypts the received data via XTEA, adds its unique MAC address, and sends it all over the air [24–26]. The receiver receives the sent data from the air, decrypts them, and transmits them to the client program via USB. The sender and receiver can transfer open data via USB.

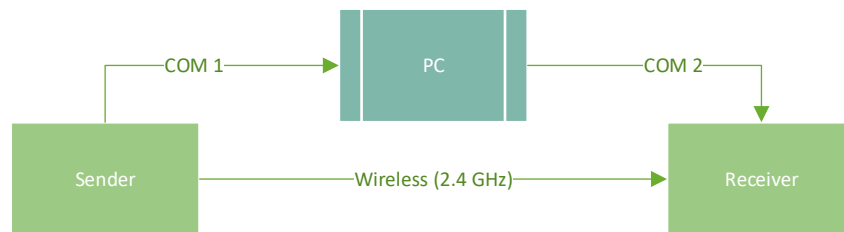


Figure 10: Schematic diagram of the project data

For the device to understand that a key will be sent now, which will need to be used instead of the standard one, we send it a group of eight identical bytes containing the ENQUIRY control symbol (number 05) from the ASCII table, which is intended for use in teletype communication and this situation can be used for its purposes (Fig. 11).

When dividing data on the computer side, the following are used: the number of iterations, the remainder of the last iteration, and the number of elements that are not enough to complete the last iteration.

First, we check whether the multiple lengths of the entered data are eight. From here, we determine how many iterations we need to perform. Then we get how many elements remain on the last (possibly incomplete) iteration. Then we calculate the number of missing elements and increase the array by this number. When you resize a .NET array, it automatically initiates all undefined elements with zeros.

You can receive data on a computer without distribution because its memory is much larger than the devices' memory [27]. In general, the integration was successful with some side issues that arose during the project implementation on the platform used, which were identified, researched, and resolved. The devices have proven to be very promising and durable, so that this project will be continued with new goals and objectives [28].

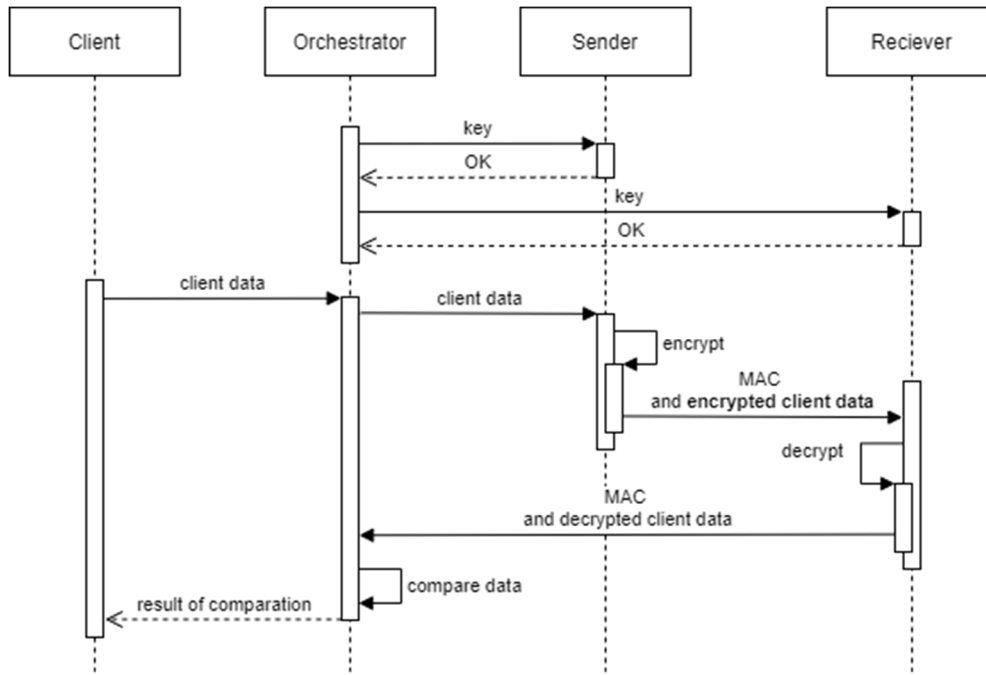


Figure 11: Key and data exchange scheme

8. Conclusions and Further Research

Based on the results, we can conclude that all the studied factors, namely the server CPU load, CPU temperature, current consumption, affect the download speed of files. When the processor is running at peak temperatures, there is a process of protection against overheating, called clock throttling. At this time, the processor's clock speed is reduced, and its performance and efficiency are reduced. This reduces the file download speed. According to the obtained data, it is seen that in RPi Zero W, the process of clock throttling occurred with a higher decrease in clock frequency and processor performance than in RPi 3.

For large amounts of wireless data transmission, RPi 3 is more suitable than RPi Zero W because the third version is more productive to support more simultaneous tasks. However, RPi Zero is more suitable for simple tasks because this version requires less power. For IoT devices, autonomy and power consumption is a significant indicator.

Based on the developed experimental layout and experimental results, we developed and created our ad hoc sensor network implementation based on the ESP8266 microcontroller. For this purpose, third-party open solutions were used in combination with our developments. Several tests were conducted, which allowed us to understand its weaknesses better, pay attention to them in the future, and think about ways to correct them. The obtained results can be effectively used for further work in this direction, which will improve the solution and further apply them in real platforms, built on both small controllers and full-fledged network devices, which is now quite relevant. Therefore, in the future, after some refinements, this system can be used in current conditions to form an intelligent home system to obtain information from specific sensors.

In this paper, there were problems in the implementation of encryption algorithms on low-power processors. Due to lack of device resources, the block XTEA encryption protocol was chosen, which, although not yet widely recognized, has shown promising results in our specific task. Several problems related to type matching, addressing, memory spaces, buffer overflow, and others were solved during the work. Resolved issues with the coordination of the receiver and transmitter. The main task of the work—building a secure communication channel by encrypting data in the channel—was solved and received firmware and application software for a full test of the devices. Also, this work has great application potential. The introduction of encryption in existing systems will have little effect on the implementation and will not affect the cost estimate of projects but will dramatically increase the security of data transmission in these networks.

Continuation of this work may verify the operability of other protocols on this and similar hardware for systems that can be embedded in data exchange projects in short-range wireless communication channels. Also of particular interest is the determination of the dependence of the information rate of data transmission on the length of the key. In the future, it is planned to check packages and implement an algorithm for retrieving lost packages; translate the client part to Python for cross-platform use; make one universal firmware for all devices (TX and RX); provide for the possibility of using many (more than two) devices on one channel; implement, test and investigate other encryption algorithms.

9. References

- [1] IEEE Standard for Local and Metropolitan Area Networks. Part 15.4: Low-Rate Wireless Personal Area Networks (LR-WPANs). doi:10.1109/ieeestd.2011.6012487.
- [2] S. Gnatyuk, Critical aviation information systems cybersecurity, in: Meeting Security Challenges Through Data Analytics and Decision Support, NATO Science for Peace and Security Series, D: Information and Communication Security. IOS Press Ebooks 47(3) (2016) 308–316.
- [3] S. Gnatyuk, et al., New secure block cipher for critical applications: Design, implementation, speed and security analysis, *Advances in Artificial Systems for Medicine and Education III* (2020) 93–104. doi:10.1007/978-3-030-39162-1_9.
- [4] F. Stajano, R. Anderson, The resurrecting duckling: Security issues for ad-hoc wireless networks, *Lecture Notes in Computer Science* (2000) 172–182. doi:10.1007/10720107_24.
- [5] S. Zhu, et al., LHAP: A lightweight hop-by-hop authentication protocol for ad-hoc networks, in: 23rd International Conference on Distributed Computing Systems Workshops, 2003. doi:10.1109/icdcs.2003.1203642.
- [6] M. M. Zanjireh, H. Larijani, A survey on centralised and distributed clustering routing algorithms for WSNs, in: 2015 IEEE 81st Vehicular Technology Conference (VTC Spring) (2015). doi:10.1109/vtc.2015.7145650.
- [7] V. Y. Sokolov, D. M. Kurbanmuradov, Method of counteraction in social engineering on information activity objectives, *Cybersecurity: Education, Science, Technique 1* (2018) 6–16. doi:10.28925/2663-4023.2018.1.616.
- [8] V. Y. Sokolov, Comparison of possible approaches for the development of low-budget spectrum analyzers for sensory networks in the range of 2.4–2.5 GHz, *Cybersecurity: Education, Science, Technique 2* (2018) 31–46. doi:10.28925/2663-4023.2018.2.3146.
- [9] Oestoidea, Access Point on Raspberry Pi 3 with Parameter Display (2017). URL: https://github.com/Oestoidea/Adafruit_Python_SSD1306.
- [10] Python Software Foundation, pi-ina219 1.2.0. Project Description (2018). URL: <https://pypi.org/project/pi-ina219/>.
- [11] I. Bogachuk, V. Sokolov, V. Buriachok, Monitoring subsystem for wireless systems based on miniature spectrum analyzers, in: 2018 International Scientific-Practical Conference Problems of Infocommunications. Science and Technology (2018). doi:10.1109/infocommst.2018.8632151.
- [12] Texas Instruments, CC2510Fx, CC2511Fx Silicon Errata (2015). URL: <http://www.ti.com/lit/er/swrz014d/swrz014d.pdf>.
- [13] V. Astapenya, et al., Last mile technique for wireless delivery system using an accelerating lens, in: 2020 IEEE International Conference on Problems of Infocommunications. Science and Technology (2020). doi:10.1109/picst51311.2020.9467886.
- [14] V. Astapenya, V. Sokolov, D. Ageyev, Experimental Evaluation of an Accelerating Lens on Spatial Field Structure and Frequency Spectrum, in: 2020 IEEE Ukrainian Microwave Week (2020). doi:10.1109/ukrmw49653.2020.9252755.
- [15] V. Sokolov, A. Carlsson, I. Kuzminykh, Scheme for dynamic channel allocation with interference reduction in wireless sensor network, in: 2017 4th International Scientific-Practical Conference Problems of Infocommunications. Science and Technology (2017). doi:10.1109/infocommst.2017.8246463.

- [16] N. Sastry, D. Wagner, Security considerations for IEEE 802.15.4 networks, in: Proceedings of the 2004 ACM Workshop on Wireless Security—WiSe'04 (2004). doi:10.1145/1023646.1023654.
- [17] Les Pounder, DS18B20 Temperature Sensor With Python (Raspberry Pi) (2017). URL: <https://bigl.es/ds18b20-temperature-sensor-with-python-raspberry-pi/>.
- [18] V. Sokolov, B. Vovkotrub, Y. Zotkin, Comparative bandwidth analysis of lowpower wireless IoT-switches. *Cybersecurity: Education, Science, Technique* 5 (2019) 16–30. doi:10.28925/2663-4023.2019.5.1630.
- [19] Aeron, Devyte, ESP8266 WiFi Mesh (2018). URL: <https://github.com/esp8266/Arduino/tree/master/libraries/ESP8266WiFiMesh>.
- [20] M. Vladymyrenko, V. Sokolov, V. Astapenya, Research of stability in ad hoc self-organized wireless networks, *Cybersecurity: Education, Science, Technique* 3 (2019) 6–26. doi:10.28925/2663-4023.2019.3.626
- [21] Pololu Corporation, Pololu Wixel User's Guide (2015). URL: <https://www.pololu.com/docs/0J46/all>.
- [22] Pololu Corporation, Wixel SDK Documentation (2015). URL: <http://pololu.github.io/wixel-sdk/>.
- [23] S. Dutta, SDCC Compiler User Guide, ver. 3.9.5, rev. 11239 (2019). URL: <http://sdcc.sourceforge.net/doc/sdccman.pdf>.
- [24] Amandeep, Implications of bitsum attack on tiny encryption algorithm and XTEA, *Journal of Computer Science* 10(6) (2014) 1077–1083. doi:10.3844/jcssp.2014.1077.1083.
- [25] Amandeep, G. Geetha, On the complexity of algorithms affecting the security of TEA and XTEA, *Far East Journal of Electronics and Communications* (2016) 169–176. doi:10.17654/ecsv3pi16169.
- [26] O. Arsalan, A. I. Kistijantoro, Modification of key scheduling for security improvement in XTEA, in: 2015 International Conference on Information & Communication Technology and Systems (2015). doi:10.1109/icts.2015.7379904.
- [27] I. Kuzminykh, et al., Investigation of the IoT device lifetime with secure data transmission, *Internet of Things, Smart Spaces, and Next Generation Networks and Systems* (2019) 16–27. doi:10.1007/978-3-030-30859-9_2.
- [28] D. Kurbanmuradov, V. Sokolov, V. Astapenya, Implementation of XTEA encryption protocol based on IEEE 802.15.4 wireless systems, *Cybersecurity: Education, Science, Technique* 2(6) (2019) 32–45. doi:10.28925/2663-4023.2019.6.3245.

Comparison of the Efficiency of Parallel Algorithms KNN and NLM Based on CUDA for Large Image Processing

Lesia Mochurad^a, Roman Bliakhar^a

^a Lviv Polytechnic National University, 12 Bandera street, Lviv, 79013, Ukraine

Abstract

Digital image processing is widely used in various fields of science, such as medicine – X-ray analysis, magnetic resonance imaging, computed tomography, cosmology – collecting information from satellites, their transmission and analysis. Image noise accompanies any of the stages of image processing – from obtaining them to segmentation and object recognition. In order to process large images in real time, the paper proposes to use CUDA technology to parallelize KNN and NLM algorithms. The subject area of the satellite images is selected for noise reduction. A comparative analysis of the effectiveness of the proposed approach. It is investigated how the computation time of successive versions of algorithms changes with increasing the size of the image itself. Based on a series of numerical experiments, it was possible to achieve an acceleration of 40 times using CUDA technology. It is shown that the calculations on the CPU significantly exceed the time spent on execution compared to the GPU. This is due to the fact that in tasks of this type, several threads on the CPU are not able to compete with thousands of threads of the graphics core. We can also observe that as the number of GPU threads increases, the time decreases significantly. This study is especially relevant in current trends in video cards. that in tasks of this type, several threads on the CPU are not able to compete with thousands of threads of the graphics core. We can also observe that as the number of GPU threads increases, the time decreases significantly. This study is especially relevant in current trends in video cards. that in tasks of this type, several threads on the CPU are not able to compete with thousands of threads of the graphics core. We can also observe that as the number of GPU threads increases, the time decreases significantly. This study is especially relevant in current trends in video cards.

Keywords

Computer vision, noise reduction, graphics processor, computational process optimization, acceleration.

1. Introduction

Most images are affected by various types of noise caused by equipment failure, problems with data transmission or compression, and natural factors. Therefore, the first stage of image processing is filtering. The presence of noise in the image can cause inaccuracies and distortions in the segmentation and recognition phase. For example, the system may perceive noise for individual objects, which, in turn, will negatively affect further research. That is why noise removal is an important problem in the field of image processing, which has many applications in various fields, such as medicine (magnetic resonance imaging, computed tomography), industry, military applications, space research, communications (satellite television) [1-3].

Most often, noise reduction is used to improve visual perception, but can also be used for specialized purposes – for example, in medicine to increase the clarity of images on X-rays [4], in pre-processing images for further recognition and more. In addition, noise reduction plays an important role in compressing video and images.

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: lesia.i.mochurad@lpnu.ua (L. Mochurad); bliakharr@gmail.ua (R. Bliakhar)

ORCID: 0000-0002-4957-1512 (L. Mochurad); 0000-0001-8478-5308 (R. Bliakhar)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

Despite the fact that noise reduction is a classic problem and has been studied for a long time, this task is still difficult and open. The main reason is that from a mathematical point of view, noise reduction is an inverse problem, so its solution is not unique. As a result, there are many different methods, the advantages and disadvantages of which will be discussed and analyzed in the next section.

Today there is a relevant area of image processing from satellites [5]. They are used by both cartographers and private companies to analyze land data. Since these images are used in further analysis by artificial intelligence methods, in particular such as computer vision, an important step is part of their pre-processing. This process often involves reducing noise in the image, which could be caused by various interferences, and therefore lose some information. That is why the subject area of satellite images is chosen for this work, because these images are large in size, and therefore require large computing power. Therefore, in this case it will be effective to take advantage of parallel calculations.

Based on the above, we can formulate the purpose of this work, namely, to conduct a comparative analysis of the effectiveness of parallel algorithms KNN and NLM based on CUDA technology to accelerate the noise of large-scale images obtained from satellite.

2. Analysis of literature sources

Significant results in the field of image noise reduction have been achieved in recent decades. This is due to the wide range of applications in everyday life, as we are surrounded by more and more digital devices: smartphones, webcams, surveillance cameras, drones and more. Accordingly, this problem occurs in such areas as: satellite television, computed tomography [4], magnetic resonance imaging, medicine in general, and in the field of research, which is actually the basis of many applications, such as object recognition and technology, such as geographic information systems, astronomy and many others [2].

Thus, image noise has remained a pressing issue for many researchers in recent times. To date, many classical filtration algorithms have been proposed (Median Filtering, Linear Filtering, Anisotropic Filtering) [3, 6, 7]. Median, linear, anisotropic and other classical filtering algorithms are usually able to remove noise, but have disadvantages, because their use loses a lot of information and suffers image geometry, resulting in objects in the image may take a completely different shape.

All linear filtering algorithms are optimal for Gaussian distribution of signals, interference and observed data, and lead to the smoothing of sharp differences in the brightness of the processed images. However, real images, strictly speaking, are usually not subject to this probability distribution.

Some of these problems are solved using non-parametric image processing methods. Studies have shown that good results in image processing were obtained using median filtering [8]. Note that the median filtering is a heuristic method of processing, its algorithm is not a mathematical solution of a strictly formulated problem. In the last two decades, nonlinear algorithms based on ranking statistics for the recovery of images damaged by various types of noise have been actively developed in digital image processing.

One of the algorithms without the above disadvantages is Non-Local Means (NLM) [8]. Because the principle is simple and straightforward, and the information contained in the image is better stored, NLM has become a popular noise reduction algorithm in recent years. Although NLM has achieved good results, there are some areas for improvement [9]. Mainly in the following aspects: first, although good results are achieved, the efficiency of the algorithm is slightly lower than traditional algorithms; secondly, it is difficult to determine the weight of similarity between image blocks, which in turn affects the efficiency of noise reduction [10].

Another effective algorithm, in terms of reducing noise and maintaining the informativeness of the original image, is K-Nearest Neighbors [11]. This algorithm is conceptually similar to the above NLM, but, as indicated in studies, with lower computational complexity [12]. Therefore, in theory, it should work faster than NLM, but the result of noise reduction should be compared in more detail.

So let's highlight the problem statement for our study:

- Develop a parallel algorithm of KNN and NLM methods based on CUDA technology and implement them programmatically.
- Test and compare the noise reduction results for each algorithm, as well as compare both algorithms with each other.
- Measure the execution time of each algorithm with a different configuration of the computer system for further calculation of acceleration and comparative analysis of efficiency.

3. Methods and tools

Mathematically, the problem of image noise reduction can be modeled as follows: $y = x + n$, where y – noisy input image, x – unknown clear image without noise, and n – is an additive of white Gaussian noise with standard deviation σn , which can be estimated in practice by various methods, such as the median of absolute deviations (MAD) [Ошибка! Источник ссылки не найден.3] or the principal component algorithm (PCA) [Ошибка! Источник ссылки не найден.4].

The purpose of image noise reduction is to reduce distortion while minimizing the loss of original features of the input image and increase the signal-to-noise ratio (SNR) [15].

The following main challenges can be identified for this task:

- flat areas of the image should be smooth;
- the edges must be protected from erosion;
- textures must be preserved;
- new distortions and artifacts should not be formed.

The principle of the first methods of noise reduction of images was quite simple: the color of the pixel was replaced by the average color of neighboring pixels. The law of variance in probability theory ensures that if nine pixels are averaged, the standard deviation of the mean noise is divided by three [8]. Thus, if we can find for each pixel nine other pixels in the image with the same color (before distortion due to noise), we can divide the noise into three (or four with 16 similar pixels, etc.) [9].

In fact, the most similar pixels will not always be the closest, on the contrary almost never. For example, you can give images with periodic patterns (when a certain fragment may occur several times). So the solution is to scan a huge part of the image for all the pixels that really resemble the pixel we want to replace (mute). The noise reduction task is then performed by calculating the average color of the most similar pixels found. This results in much greater clarity after filtering and less loss of image detail compared to average local algorithms. This algorithm is called non-local means and its definition can be written as follows:

Let Ω – image area, and p and q are two dots (pixels) within the image, then:

$$NLu(p) = \frac{1}{C(p)} \int_{\Omega} f(d(B(p), B(q))) u(q),$$

where $u(p)$ is the filtered value of the image at the point p ; $u(q)$ is the unfiltered value of the image at the point q ; $C(p)$ – this is the normalizing factor; f is a weighing function in which $d(B(p), B(q))$ is the Euclidean distance between image patches focused on p and q in accordance.

When calculating the Euclidean distance $d(B(p), B(q))$ all pixels in the patch $B(q)$ have the same weight (importance), so the function $f(d(B(p), B(q)))$ can be used to mute all pixels in the patch $B(p)$ and not just for p .

This algorithm can be performed in two implementations: pixelwise and patchwise. In our work we use pixel implementation, because in fact there is not much difference between the two options, so the overall quality in terms of saving details is not improved with the use of patchwise implementation.

Let's introduce the noise reduction of the color image $u = (u_1, u_2, u_3)$ and a specific pixel p as follows: $\hat{u}_i(p) = \frac{1}{C(p)} \sum_{q \in B(p, r)} u_i(q) w(p, q)$, $C(p) = \sum_{q \in B(p, r)} w(p, q)$, where $i = 1, 2, 3$ and $B(p, r)$ denote the environment centered at the point p with dimensions $(2r + 1) \times (2r + 1)$ pixels. The search area is limited by a square circumference of a fixed size, which is due to the limitation of

calculations. This is a window 21×21 for small and medium values σ . The window size increases to 35×35 for large values σ due to the need to detect more similar pixels to reduce additional noise.

Weight $w(p, q)$ will depend on the Euclidean distance $d^2 = d^2(B(p, f), B(q, f))$ for $(2r + 1) \times (2r + 1)$ color patches centered accordingly p and q :

$$d^2(B(p, f), B(q, f)) = \frac{1}{3(2f + 1)^2} \sum_{i=1}^3 \sum_{j \in B(0, f)} (u_i(p + j) - u_i(q + j))^2.$$

That is, each pixel value is restored as the average of the most similar pixels, where this similarity is calculated in the color image. Therefore, for each pixel, each channel value is the result of the average of the same pixel [15].

To calculate the weight $w(p, q)$ we will use the exponential function: $w(p, q) = e^{-\frac{\max(d^2 - 2\sigma^2, 0.0)}{h^2}}$, where σ indicates the standard deviation of the noise and h is the filtering parameter set relative to the value σ [12, 15]. The weight of the function is set to average similar noise patches. That is, patches with square distances less than $2\sigma^2$ are installed to 1, while larger distances decrease faster due to the exponential function.

The weight of the reference pixel p on average, is set as the maximum weight in the vicinity $B(p, r)$. This parameter avoids excessive weighing of the reference point in the middle. Otherwise $w(p, q)$ should be equal to 1 and as a result, more value will be required to reduce noise h . Therefore, using the above procedure, we will be able to restore the noiseless value for each pixel p .

To objectively assess the basic metrics of our noise reduction using the NLM algorithm, consider an additional algorithm – KNN [Ошибка! Источник ссылки не найден.6]. The main principle of KNN is that the category of the data point is determined according to the classification of the nearest K neighbors. The KNN (k-Nearest Neighbors) filter was developed to suppress white noise, which is based on the Gaussian Blur Filter algorithm [17]. Consider how it works.

Let $u(x)$ – the value of the input noisy image at the point x , a $f(y)$ – the result obtained by the KNN filter with parameters h and r at some point y . Ω is the space of a given size around the selected pixel, i.e. it is a block of pixels in size $N \times N$, so that the selected pixel (dot y) is located in the center of this block. Then the filtration formula can be represented as: $f(y) = \frac{1}{C(x)} \sum_{\Omega(x)} u(y) e^{-\left(\frac{(y-x)^2}{r^2}\right)} e^{-\frac{(u(y)-u(x))^2}{h^2}}$, where $C(x)$ – normalizing factor.

Since the purpose of our study is to accelerate the noise reduction of images using the GPU, for such purposes we will use CUDA [Ошибка! Источник ссылки не найден.]. Is a software-hardware architecture of parallel computing, which allows you to significantly increase computing performance through the use of graphics processors from Nvidia. CUDA allows you to reduce the load on the processor, which gives significant benefits over time. Given the power of modern video cards, we can get a significant acceleration in the calculations. And because we will be working with images, this technology should be ideal for this type of task.

Both methods can be easily implemented using CUDA technology. Since this technology uses the SIMD parallel computing model [19], we do not need to parallelize the algorithm itself (NLM, KNN). It is enough to implement a standard version of the algorithm in the form of a kernel function. The kernel function is actually a normal sequential program, without organizing the launch of threads. Instead, the GPU runs many copies of a given function in different threads.

Consider in more detail the mechanism for organizing parallel calculations:

- Host – CPU. Performs a control role – runs tasks on the device, allocates memory on the device, moves memory to/from the device.
- Device – GPU. Performs the role of "subordinate" – does only what the CPU tells him.
- Kernel – a task (NLM, KNN), which runs the host on the device.

Threads are grouped into blocks, in the middle of which they have shared memory and run synchronously. The blocks are combined into grids, the main task of which is to serve as an isolated container to which the kernel function applies (a separate function can be applied to each grid). As part of our problem, the calculation creates three grids for a color image, and one for black and white (because algorithms split images into RGB color vector). The number of threads is equal 2^n , where $1 < n \leq 10$ ($2^{10} = 1024$, the maximum number of threads per block). The input noisy image is

divided by k blocks where $k_{ij} = n_i \times n_j$, $i \in \frac{N}{n}$, $j \in \frac{M}{n}$, where N – width, a M – image height, respectively. Everyone k_{ij} the grid block corresponds k_{ij} an image block, where for each pixel of the image the weight in a separate stream is calculated.

To conduct experiments, it was decided to develop a software product in the Python programming language. Accordingly, the PyCuda library was used to use the CUDA parallel computing architecture. As well as an OpenCV library for image processing. All calculations were performed in Google Colab [20], which allows you to write and execute Python code in a browser with a minimum of settings and free access to powerful computing capabilities. The configuration includes a graphics professor NVIDIA Tesla T4 with 2560 CUDA cores and 16 Gb of video memory and an Intel Xeon CPU with a clock speed of 2.20Ghz with two cores and two threads.

4. Results of numerical experiments

First, we evaluate the quality of both algorithms and evaluate the effect of noise reduction of various images. First we test the work of KNN.



Figure 1: Comparison of 1 image fragment before (left) and after (right) noise reduction using KNN algorithm



Figure 2: Comparison of 5 fragments of the image before (left) and after (right) noise reduction using the KNN algorithm

Now let's perform the same tests for the NLM algorithm.



Figure 3: Comparison of 5 fragments of the image before (left) and after (right) noise reduction using the NLM algorithm



Figure 4: Comparison of 6 fragments of the image after noise reduction using KNN (left) and NLM (right)



Figure 5: Comparison of 8 fragments of the image after noise reduction using KNN (left) and NLM (right)

In addition, studies were performed on images of different dimensions and with different configurations of computing power to measure the runtime of algorithms with CPU/GPU comparison. The calculations were performed for images of different dimensions, where P – the number of pixels in the image $n \times m$ pixels, and for GPUs with different number of threads per block (4, 16, 64, 256 and 1024 threads, respectively).

Table 1

Program execution time in seconds for KNN algorithm with different CPU/GPU startup configuration

P	CPU	GPU 4 tpb	GPU 16 tpb	GPU 64 tpb	GPU 256 tpb	GPU 1024 tpb
175k	0.056724	0.009772	0.005263	0.005370	0.005531	0.004922
700k	0.204899	0.030974	0.014606	0.011442	0.013098	0.011224
3kk	0.799452	0.102840	0.048849	0.038600	0.037616	0.033931
11kk	3.366155	0.405643	0.144978	0.118630	0.124715	0.121243
45kk	12.349476	1.169624	0.490328	0.418505	0.433192	0.396563
81kk	21.875035	2.154644	0.869768	0.721381	0.657146	0.594174

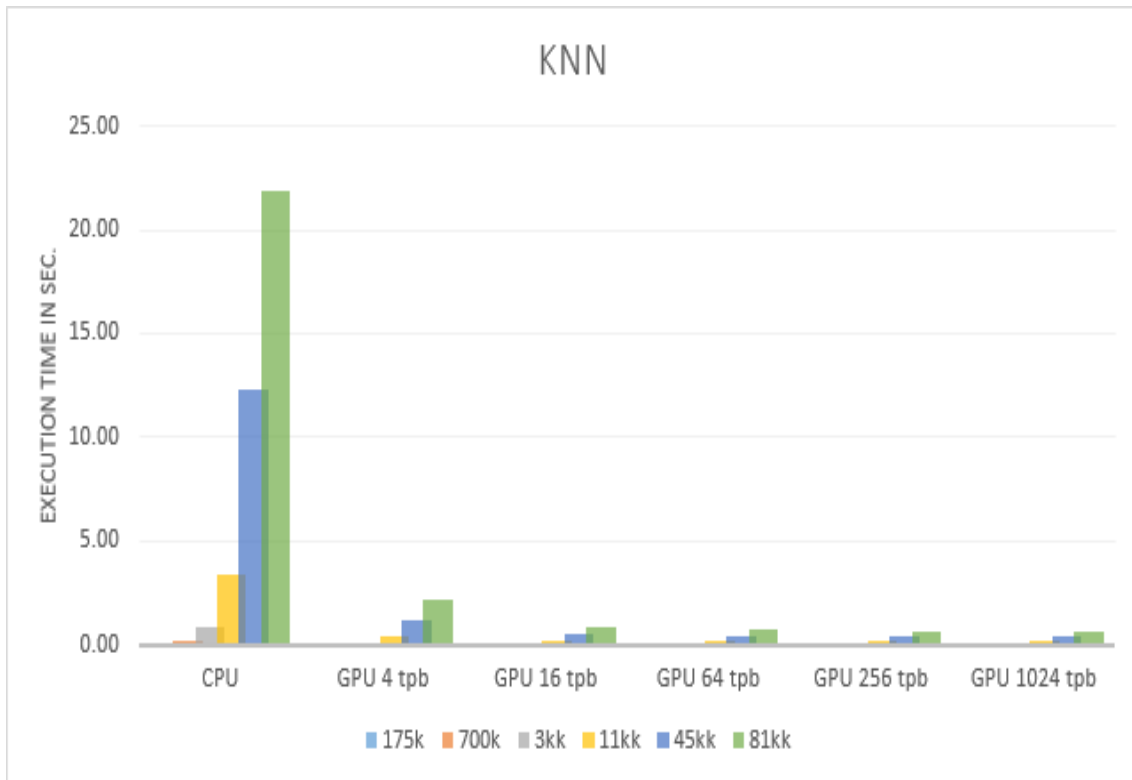


Figure 6: Execution time of parallel KNN algorithm using CPU/GPU in seconds for different image dimensions and with different program startup configuration

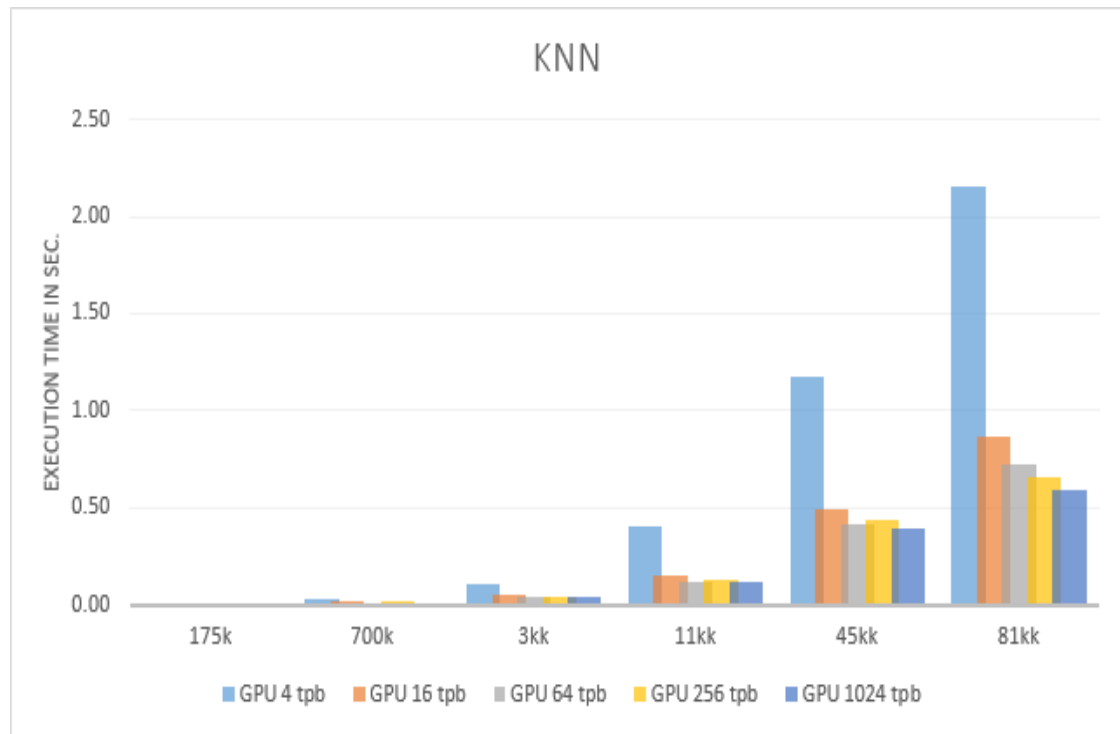


Figure 7: Comparison of runtime results in seconds for KNN algorithm executed using GPU

Table 2

NLM algorithm execution time in seconds with different CPU/GPU startup configuration

P	CPU	GPU 4 tpb	GPU 16 tpb	GPU 64 tpb	GPU 256 tpb	GPU 1024 tpb
175k	0.062403	0.023235	0.009618	0.007488	0.005829	0.004538
700k	0.230475	0.083275	0.028574	0.017382	0.010574	0.006732
3kk	0.853824	0.316450	0.066979	0.047274	0.033365	0.023549
11kk	3.264484	0.800226	0.204644	0.157260	0.120847	0.087866
45kk	13.059060	2.331046	0.811966	0.607054	0.453854	0.339316
81kk	23.748522	3.996524	1.333901	1.009624	0.720390	0.543901

We calculate the acceleration for the measurements of each algorithm. To do this, divide the execution time by CPU by the execution time by GPU.

Table 3

Acceleration factor for the KNN algorithm performed using the GPU

P	GPU 4 tpb	GPU 16 tpb	GPU 64 tpb	GPU 256 tpb	GPU 1024 tpb
175k	5.804566	10.777158	10.562708	10.255388	11.524214
700k	6.615167	14.028813	17.907669	15.643454	18.254952
3kk	7.773751	16.365716	20.711111	21.252815	23.561176
11kk	8.298323	23.218389	28.375263	26.990804	27.763618
45kk	10.558499	25.186156	29.508552	28.508098	31.141290
81kk	10.152507	25.150428	30.323849	33.287953	36.815880

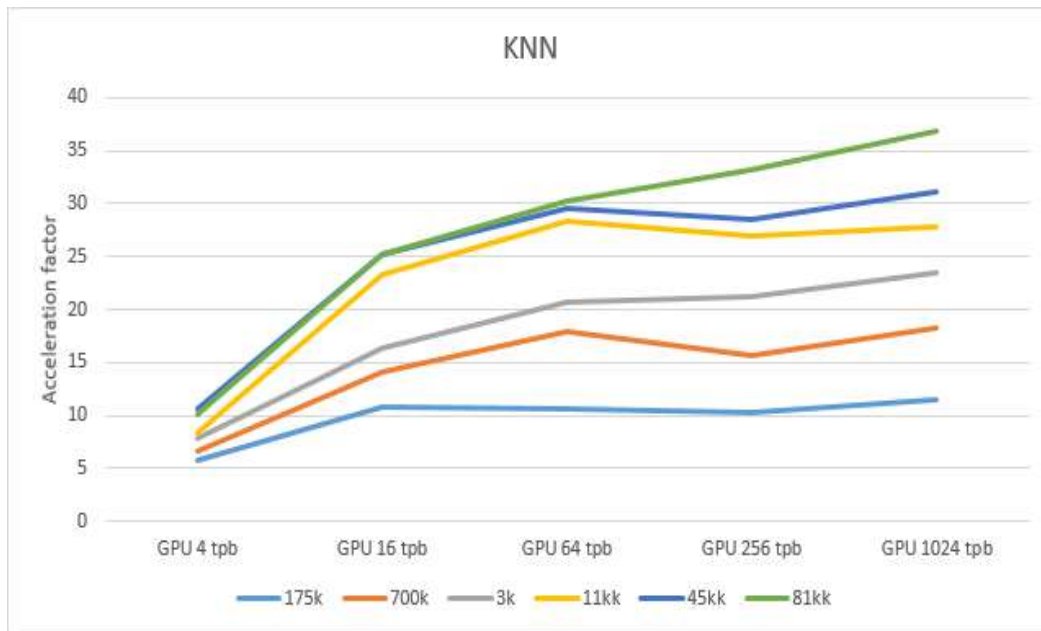


Figure 8: Dependence of the acceleration factor on the GPU on the number of threads per block at different image dimensions for the KNN algorithm

Table 4

Acceleration factor for the NLM algorithm performed using the GPU

P	GPU 4 tpb	GPU 16 tpb	GPU 64 tpb	GPU 256 tpb	GPU 1024 tpb
175k	2.685736	6.488113	8.334059	10.705201	13.750961
700k	2.767646	8.065955	13.259593	21.797397	34.235835
3kk	2.698137	12.747567	18.061340	25.590138	36.257286
11kk	4.079453	15.952002	20.758530	27.013322	37.153126
45kk	5.602232	16.083252	21.512204	28.773715	38.486371
81kk	5.942295	17.803809	23.522147	34.966196	43.663300

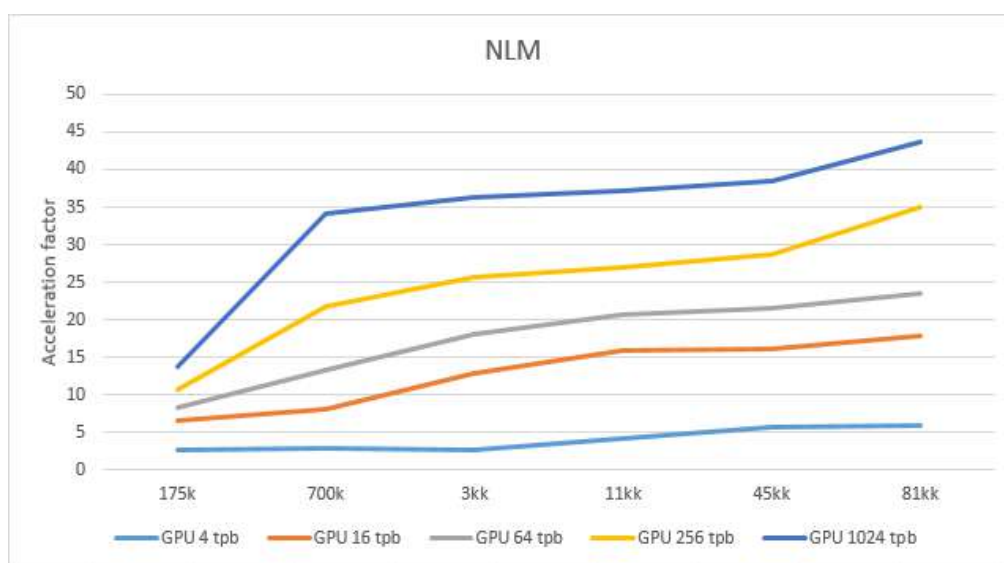


Figure 9: Dependences of the acceleration factor on the GPU on the number of threads per block at different image dimensions for the NLM algorithm

5. Discussion

Now let's analyze our experiments. To begin with, let's summarize the results of the noise reduction using each of the NLM and KNN algorithms.

As we can see from the above Figures 1-3, the algorithms reduce noise well and restore the original image even with loud noise. For comparison, in Figure 2 using the KNN algorithm, the image fragment after clearing loses the clarity of the contours, the image becomes slightly pixelated. At the same time in Figure 3 algorithm NLM visually makes less noise, so the image is still quite blurred, but the boundaries of the contours are not as violated as when noise reduction using the KNN algorithm. However, this is just one example, so don't judge by just one piece.

In Figure 4 shows that the image obtained with NLM, less clear compared to KNN. This is best seen in detail when looking at cars parked in the parking lot in the upper right corner of the image. Similar results can be observed in Figure 5. The image obtained during the KNN algorithm is clearer and richer. An example is the road marking lines at the bottom of the image, which in the first case, retained their geometry, not to mention the image on the right, where the lines are almost invisible, the image obtained by NLM was "blurred".

A similar problem is mentioned in [21], where the authors emphasize that in high-detail images, after the noise reduction with NLM, the geometry of objects deteriorates.

The next step is to analyze our program execution time metrics and GPU acceleration factors. For clarity, tests were performed on different image sizes. Therefore, it is clear from Table 1 and Table 2, that the largest time difference between CPU/GPU we get for scale images. This confirms the correctness of the choice of subject area for our task of accelerating noise reduction.

Considering the visualization in Figure 6 and Figure 7, respectively, we can see that the calculations on the CPU significantly exceed the time spent on execution compared to the GPU. This is due to the need for a large number of parallel calculations for image processing. For tasks of this type, multiple threads on the CPU are not able to compete with thousands of graphics core threads. We can also observe that as the number of GPU threads increases, the time decreases significantly.

In Table 3 and Table 4 there is a tendency to increase the rate of acceleration with increasing number of flows per unit. This is because the more threads, the fewer blocks, and therefore less time spent on memory operations.

Comparing the acceleration rates of both algorithms. We can see that the parallel NLM algorithm loses with a small number of threads, but wins with the largest number of threads per block. We can also see that the KNN algorithm does not receive such a large increase in acceleration as the NLM algorithm when increasing the image size for noise reduction from 45 million pixels to 81. This is well observed in Figure 8 and Figure 9, where the acceleration coefficient graph for the NLM algorithm grows faster than the KNN. Although KNN at 4 threads gets twice as fast as NLM.

If we compare the time spent on both algorithms, we can conclude that there is no clear favorite, although KNN shows better results with a small number of threads per block, due to the peculiarity of the algorithms. This can affect the results if the calculations are performed using older generation video cards, where the maximum possible number of threads per unit is many times less. In the conditions of use of modern video cards NLM, with the maximum number of threads, is ahead of the competitor.

6. Conclusions

This paper compares KNN and NLM algorithms in the context of parallel computations to accelerate the noise reduction of large images. To do this, using the Python programming language, implemented two algorithms with versions for CPUs and GPUs using the CUDA architecture.

In the course of the work, the results of each algorithm were analyzed and it was found that both algorithms cope well with the task of noise reduction and recovery of the original images obtained as satellite images.

Comparative tables are also constructed for each algorithm with a different configuration of the program start, in which you can evaluate the obtained time measurements and calculated acceleration rates. These data were visualized on graphs, which allowed us to objectively compare the obtained

metrics. Regarding the obtained time metrics and acceleration metrics, we can conclude that parallel computing with CUDA on GPUs is many times better than computing on CPU. It is for these algorithms that the GPU and CUDA technology can achieve an acceleration of approximately 40 times. Thanks to these results, we can use algorithms such as KNN and NLM in real-time video processing. It will also save a lot of time on processing images obtained from telescopes and satellites, which are high resolution.

7. References

- [1] L. Guo, H. Zhang, H. Zhai, Ch. Gong, B. Song, Sh. Tong, Zh. Cao, X. Xu, Research on image processing and intelligent recognition of space debris, Proceedings Volume 10988, Automatic Target Recognition XXIX;1098817 (2019). doi:10.1117/12.2525058.
- [2] Unsupervised Denoising for Satellite Imagery using Wavelet Subband CycleGAN [Electronic resource]. – 2020. – Resource access mode:<https://arxiv.org/pdf/2002.09847.pdf>.
- [3] L. Fan, F. Zhang, H. Fan, et al, Brief review of image denoising techniques, Vis. Comput. Ind. Biomed. Art 2, 7 (2019). doi:10.1186/s42492-019-0016-7.
- [4] Sameera V. Mohd Sagheer, Sudhish N. George, A review on medical image denoising algorithms, Biomedical Signal Processing and Control, volume 61 (2020), 102036. doi:10.1016/j.bspc.2020.102036.
- [5] A. Asokan, J. Anitha, M. Ciobanu, A. Gabor, A. Naaji, D. Hemanth, J. Image Processing Techniques for Analysis of Satellite Images for Historical Maps Classification — An Overview. Appl. Sci. (2020), 10, 4207. doi:10.3390/app10124207.
- [6] C. He, K. Guo, H. Chen, An Improved Image Filtering Algorithm for Mixed Noise, Appl. Sci. (2021), 11, 10358. doi: 10.3390 / app112110358.
- [7] S. Kushwaha, R. Singh, Performance Comparison of Different Despeckled Filters for Ultrasound Images, Biomed Pharmacol J., (2017), 10(2). doi:10.13005/bpj/1175.
- [8] G. Shweta, Meenaksh, A Review and Comprehensive Comparison of Image Denoising Techniques, International Conference on Computing for Sustainable Global Development. (2014), P. 975.
- [9] K. Chen, X. Lin, X. Hu, et al, An enhanced adaptive non-local means algorithm for Rician noise reduction in magnetic resonance brain images, BMC Med Imaging 20, 2 (2020). doi:10.1186/s12880-019-0407-4.
- [10] X. Zhang, Two-step non-local means method for image denoising, Multidim Syst Sign Process (2021). doi:10.1007/s11045-021-00802-y.
- [11] Haiyan Wang, Peidi Xu, Jinghua Zhao, Improved KNN Algorithm Based on Preprocessing of Center in Smart Cities, Complexity, vol. 2021, Article ID 5524388, 10 pages, (2021). doi:10.1155/2021/5524388.
- [12] M. Colom, A. Buades, Analysis and Extension of the Ponomarenko et al. Method, Estimating a Noise Curve from a Single Image [Electronic resource], IPOL. (2013) Resource access mode:http://www.ipol.im/pub/art/2013/45/article_lr.pdf.
- [13] Chunjie Wu, Yi Zhao & Zhaojun Wang, The median absolute deviations and their applications to shewhart $\bar{X}$ control charts, Communications in Statistics – Simulation and Computation, 31: 3, 425-442, (2002). doi:10.1081/SAC-120003850.
- [14] Sidharth Mishra, Uttam Sarkar, Subhash Taraphder et al, Principal Component Analysis, International Journal of Livestock Research. (2017). doi:10.5455 / ijl.20170415115235.
- [15] Suarez-Perez Alex, Gabriel Gemma, Rebollo Beatriz et al, Quantification of Signal-to-Noise Ratio in Cerebral Cortex Recordings Using Flexible MEAs With Co-localized Platinum Black, Carbon Nanotubes, and Gold Electrodes. Frontiers in Neuroscience, 12. (2018). doi:10.3389 / fnins.2018.00862.
- [16] Yanhui Guo, Siming Han, Ying Li, Cuifen Zhang, Yu Bai, K-Nearest Neighbor combined with guided filter for hyperspectral image classification, Procesia Computer Science, volume 129. (2018), pp. 159-165. doi:10.1016/j.procs.2018.03.066.
- [17] P. Singhal, A. Verma and A. Garg, "A study in finding the effectiveness of Gaussian blur filter over bilateral filter in natural scenes for graph based image segmentation," 2017 4th International

- Conference on Advanced Computing and Communication Systems (ICACCS). (2017), pp. 1-6. doi:10.1109 / ICACCS.2017.8014612.
- [18] L. Mochurad, G. Shchur, Parallelization of Cryptographic Algorithm Based on Different Parallel Computing Technologies,. Proceedings of the Symposium on Information Technologies & Applied Sciences (IT&AS 2021). Bratislava, Slovak Republic, March 5. (2021), Vol-2824, 20-29 p.
 - [19] L. Mochurad, N. Boyko, Technologies of distributed systems and parallel computation: monograph. Lviv: Publishing House “Bona”, (2020), 261 p.
 - [20] Google Colaboratory [Electronic resource] – Resource access mode:<https://research.google.com/colaboratory/faq.html>.
 - [21] A. Buades, B. Coll, and J. Morel, On image denoising methods, Technical Report 2004-15, CMLA. (2004).

Data Mining Information System for Complex Technical Systems Failure Risk Evaluation

Vladimir Vychuzhanin^a, Nickolay Rudnichenko^a, Tetiana Otradska^a, Igor Petrov^b

^a Odessa National Polytechnic University, Shevchenko Avenue 1, Odessa, 65001, Ukraine

^b National University Odessa Maritime Academy, Didrichson street 8, Odessa, 65029, Ukraine

Abstract

The article presents the results of the development of a data mining information system for complex technical systems failure risk evaluation and modelling, which allows taking into account the probabilities of system elements failure with the analysis of the numerical values of damage to interelement connections. The proposed system implements a risk assessment method, as well as intelligent models for assessing and analyzing the risk of elements failures and components of complex technical systems of various topologies and hierarchies based on simulation modeling, soft computing theory and data mining models. For the holistic storage of the obtained and revealed knowledge from the data in the system, a distributed data warehouse module configuration has been developed, based on the use of trained decision tree models, which makes it possible to automate the processes of generating and analyzing brief metadata on the key parameters of the system elements in dynamics when assessing the risks of their failure. The use of the proposed method and models within the framework of the developed system makes it possible analysis process automating, assessing the performance and reliability of complex technical systems elements and interelement connections in heterogeneous emergency scenarios for various target risk management programs. The use of the developed data mining information system for complex technical systems failure risk evaluation has made it possible to increase the efficiency and effectiveness of managerial decision-making by 10% compared to the existing cognitive-simulation approach without intelligent models.

Keywords

Complex technical system, imitational modeling, data mining system, cognitive and fuzzy models, expert systems, knowledge base.

1. Introduction

The efficiency of functioning of complex technical systems (CTS) largely depends on the reliability of their interconnected, interacting elements, inter-element connections (IEC). Examples of CTS are telecommunications, energy, transport systems, oil and gas production, mining and processing industries. Ensuring reliability is possible by using methods, diagnostic tools and predicting the technical state of systems at the stages of CTS design and operation [1,2]. Assessment of the technical state of elements and IEC is based on the analysis of the CTS failure risk, which combines the probabilities and possible damage from failures [3-6]. The problems arising in assessing the risk of CTS failures are associated [2] with: a variety of equipment with structural and functional features, differing in physical principles and modes of operation, making it difficult to use universal methods, means of diagnosing the technical state of systems; a composition of numerous interconnected, interacting subsystems, blocks and nodes, with non-trivial IEC, between which there is a mutual exchange of energy, matter and information (EMI); incompleteness and vagueness of information, accounting complexity and interaction of interrelated subsystems and their elements.

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: vint532@yandex.ua (V.Vychuzhanin); nickolay.rud@gmail.com (N.Rudnichenko); tv_61@ukr.net (T. Otradska), firmn@list.ru (I. Petrov);

ORCID: 0000-0002-6302-1832 (V.Vychuzhanin); 0000-0002-7343-8076 (N.Rudnichenko); 0000-0002-5808-5647 (T. Otradska), 0000-0002-8740-6198 (I. Petrov);



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

To ensure high indicators of efficiency, reliability, environmental performance, it is required to use methods of diagnostics and forecasting the technical state of CTS [7], which allow timely identification of the presence and location of failed components, determine the degree of operability under changes in operating conditions and predict possible system failure.

Taking into account the described specifics of the CTS, one should take into account the heterogeneity and large amount of diagnostic data obtained for various subsystems and elements [4]. This increases the relevance of automating the search and identifying hidden patterns in the identified failures and violations of the operating modes of the components of technical systems, which becomes possible through the use of modern methods and technologies of data mining built into the process of analyzing and evaluating data.

2. Description of Problem

Analysis of publications showed the predominant use of deterministic, probabilistic, expert, combined methods for assessing the failure risk of various types of CTS [8]. The methods based on probabilistic safety analyzes consider only individual sequences of events during the operation of the equipment. Many methods are based on the assumption that all elements of the system are operating normally. Often, methods for assessing the risk of CTS failures are based on engineering, model, expert and other approaches that have a narrow industry focus, which limits the possibility of their widespread use. However, the significant advantages of the probabilistic method of Bayesian networks and its further development make it promising to use the method for assessing the structural and functional risk of failures of interconnected, interacting CTS components. Bayesian networks combine empirical frequencies of occurrence of various values of variables, subjective estimates of expectations and theoretical ideas about the mathematical probabilities of certain consequences from a priori information [8].

The use of information, analytical and intellectual capabilities in the systems of diagnostics and forecasting with the determination of reliability level allows for high-quality design and maintenance CTS equipment [2,9,10]. The use of a data mining information system (DMIS) of the risk of CTS failures, based on knowledge, fast processing and analysis of large volumes of heterogeneous information provides an adequate approach to design, reliable operation, and high-quality diagnostics of systems [11]. The basis for the intelligent solution of the problems CTS state diagnostics and forecasting are traditional for the class of unstructured and poorly formalized problems [11, 12]: the lack of the possibility of obtaining complete and reliable information for making informed decisions; the presence of uncertainty in the data used, as well as multivariance in the search for a solution. Such problems lead to an increase in the labor intensity and financial costs of diagnostics, forecasting at the stages of CTS operation and design. The solution of problems within the framework of the DMIS is possible by using methods and models that use the results of simulation and fuzzy modeling [2,9,10]. The successful implementation of diagnostics and forecasting CTS risk of failures is achieved by the use of achievements in the ISDF, in particular, the technologies of expert systems (ES) using knowledge bases (KB) [2].

The analysis of publications [27,28] showed also that the development of information technologies allows solving the problem of assessing the failure risk of CTS with a large elements number, with the display of systems in the directed graphs form with structural and functional relationships and interactions of elements [2,13,14]. To study the CTS model, among the many existing methods of reliability modeling, it is recommended to use a relatively developed technology of cognitive simulation (CS) of the CTS failure risk [15,16]. Each element of the CTS corresponds to the vertex of the digraph, to each IEC, a directed arc that coincides with the direction of transmission of the EMI. The vertex of the directed graph and the branch connecting the vertices of the model graph can have certain characteristics expressed in numerical values (branch weights, branch transformation coefficients, and so on).

Within the framework of the cognitive approach, the basis of the model is a cognitive map in the form of a directed graph, which includes the elements of the CTS [17-19]. The advantage of CS over analytical methods for studying the risk assessments of CTS failures is the use of partially reliable and incomplete data about the object of research in modeling.

The developed theoretical base and the availability of a large range of simulation software with AnyLogic, Extend, Arena and other famous systems [20] contributes to the use of CS. Considering the existing uncertainty, incompleteness and vagueness of the information received by the systems for diagnostics and forecasting CTS technical state, it is advisable in the ISDF, in addition to cognitive model (CM), to use fuzzy-probabilistic models (Fig. 1) [20-23], which will allow the joint use of heterogeneous data for obtaining a reliable assessment failure risk.

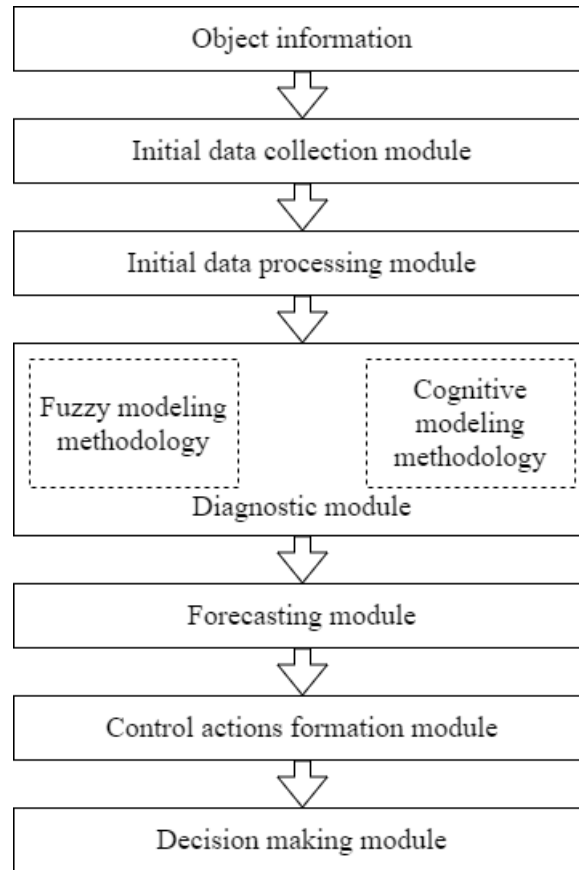


Figure 1: DMIS Structure

The initial data processing module (Fig. 1) in DMIS will allow processing data for further clear cognitive and fuzzy modeling. The diagnostics module will provide the construction of clear cognitive and fuzzy models for diagnostics and predicting in the design and operation CTS failures risk.

There are various approaches to the formation of fuzzy probabilistic models for assessing the risk of failures of functionally interconnected and interacting CTS. To construct such models, it is recommended to use the probabilities of loss of performance as input parameters, as well as damage from the consequences of a risk event in the event of failures of elements and IEC of the CTS [24]. The paper [25] describes the use of fuzzy-probabilistic models for assessing the risk of CTS failures, which allow early identification of a possible transition of a particular process to an emergency mode to prevent and risk unacceptable scales. In [26], a model for identifying pre-emergency situations of a technological process based on a fuzzy logic apparatus with additions by elements is used as a similar model, which allows taking into account the probability of CTS equipment failure. The use of fuzzy logic for the construction of fuzzy-probabilistic models for assessing the risk of failure of interconnected and interacting elements and CTS IEC in various operating conditions allows to reduce the time spent on researching systems [15].

When assessing the risks of failure of CTS elements as a component of the analysis of the overall level of its reliability [29], it is important to take into account that the effectiveness of DMIS largely depends on the results of using the technologies of expert systems using knowledge bases [2]. A typical diagram of an ES with a knowledge base (KB) and a database (DB) is shown in Fig. 2. When constructing a knowledge base, the choice of knowledge representation models is essential [24]. Such

models for the representation of knowledge: production, based on the rules, allowing to represent knowledge in a sentence: if a condition, then an action (the disadvantage of the model is that products contradict each other when a large number of them accumulate); frame (a frame contains a finite number of slots); ontological - conceptual constructions based on the global categories of space, time and quality. Currently, machine learning methods are being actively developed, which make it possible to automate data analysis processes with the support of making control decisions. For example, methods for constructing models of decision trees (DT), used in addition to typical problems of classification and regression to describe key ontological information about heterogeneous datasets of large volumes. This provides a more compact form of their presentation with subsequent processing and analysis.

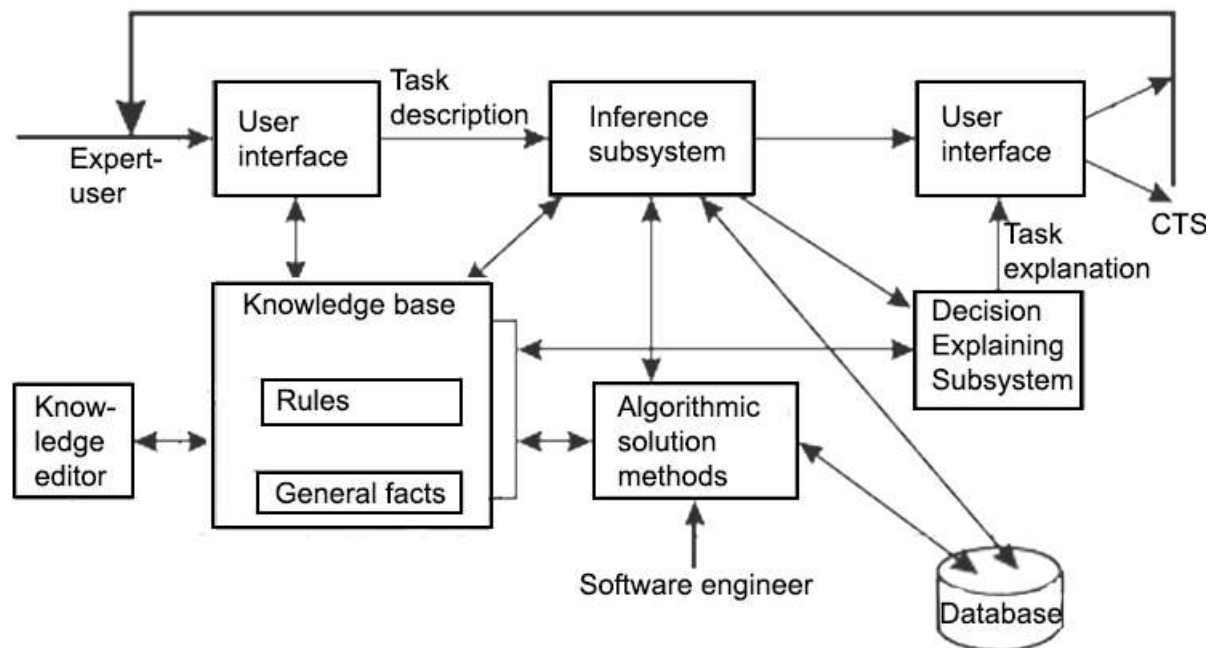


Figure 2: Typical ES scheme

DT visualization can be done by displaying a connected directed acyclic graph. The branches of the graph of the constructed decision tree DT can store the values of attributes, which are functional parameters of the elements of the studied CTS, on which its reliability depends, and the values of the failure risk are displayed on the leaves of DT [23,24]. Consisting in the fact that the structure of the DT partition can be expressed as a set of logical rules in production form, the peculiarities of constructing DT models makes it possible to use such models as the basis for KB. There are known works in which the following are proposed: an approach to constructing a fuzzy KB, which is the basis for constructing fuzzy inference systems for operational control; architectural, multi-level KB model for storing information in a subject-oriented intelligent system [26].

Thus, the further development of methods of cognitive-imitation and fuzzy-probabilistic modeling, the use of DT models as the basis for KB, is relevant when constructing the DMIS of the CTS risk failure of interrelated and interacting elements and EMI.

3. Method for diagnostics and prediction CTS risk failures in an data mining system

The method of diagnostics and forecasting involves the determination of an assessment of the risk of CTS failures based on heterogeneous information of different volume and content. The method is implemented in the DMIS scheme (Fig. 3), which includes cognitive and fuzzy models for assessing the failure risk of CTS, KB and DB elements. In the proposed DMIS, the knowledge base uses the results of fuzzy and cognitive modeling, as well as scenarios for the development of the risk of CTS

failures arising during system designing and portioning. CM is based on Bayesian analysis and Bayesian networks. A Bayesian network is a directed graph. The nodes of the network are many random variables. The vertices are connected in pairs by oriented edges. The complex indicator of the reliability of CTS elements is the risk of failure (R) - a combination of the probabilities of failures (P) and losses from failures (D), Conseq, means consequence.

$$R\left(\frac{Conseq.}{Time}\right) = P\left(\frac{Event}{Time}\right) \cdot D\left(\frac{Conseq.}{Time}\right) \quad (1)$$

$$R = \{ \langle P_{i(j)}, D_{i(j)}, q_{i(j)}, P_{ij}, D_{ij}, q_{ij} \rangle, i, j = 1, 2, \dots, N \}, \quad (2)$$

where i – the number of the vertex from which the edge CM emerges; j - the number of the vertex that contains the edge CM; $P_{i(j)}$ - element failure risk probability (IEC) CTS; $D_{i(j)}$ - damage from the consequences of the risk of failure of an element (IEC) CTS; $q_{i(j)}$ - weight $i(j)$ for each element failure risk (IEC) CTS within 0 ... 1 under the conditions $\sum_{i=1, j=1}^N q_{i(j)} = 1 (i, j = 1, N)$ and $\sum_{ij=1}^M q_{ij} = 1 (ij = 1, M)$; N - amount of elements CTS, M – amount of IEC.

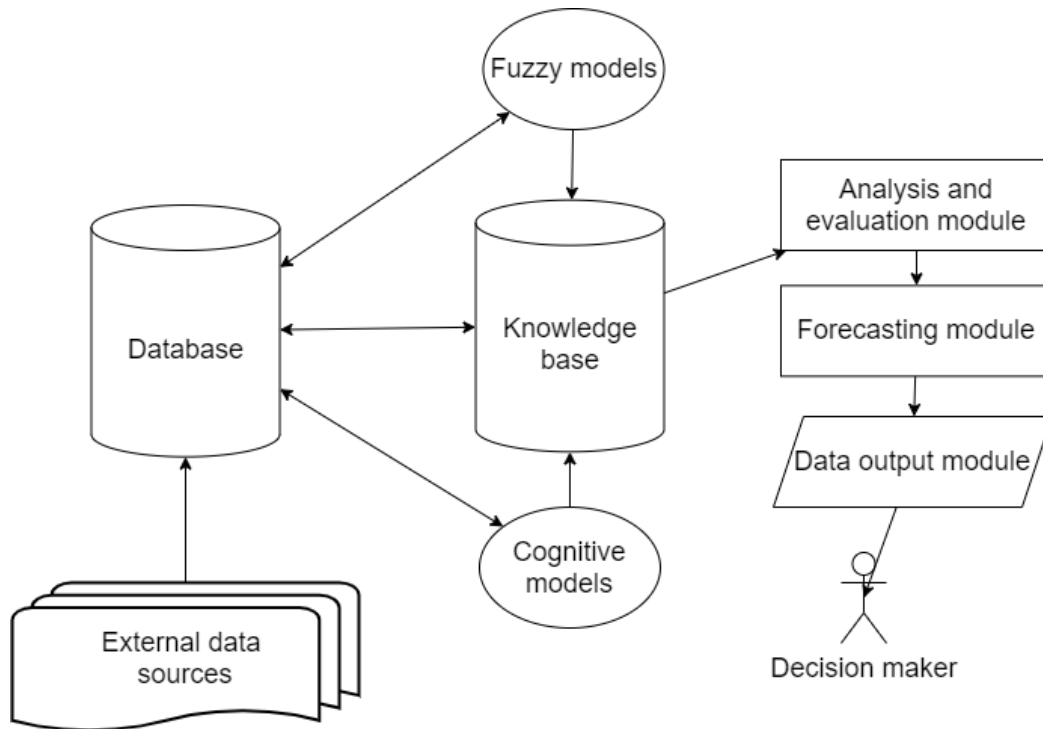


Figure 3: DMIS structural scheme

Cumulative estimates of the risk of failures, taking into account the relationships and interactions of elements and IEC CTS

$$R = \sum_{i=1}^N R_i \cdot q_i + \sum_{j=1}^M R_j \cdot q_j + \sum_{i,j=1}^{N,M} R_{i,j} \cdot q_{i,j}. \quad (3)$$

The model is based on the structural representation of the CTS in the form of a directed graph, which is negatively influenced by external impulses. Each directed graph is reflected using an incidence matrix, the columns of which are the edges of the directed graph CM CTS, and the rows are the vertices. The construction of a model for assessing the structural and functional risk of failures by elements and IEC for CTS is carried out in 3 stages. The first stage of CS is the development of a CM, taking into account the type of the EMI resource in the form of a cognitive map and an oriented graph with N vertices and M edges [2,15]. In the second stage, the structural and functional damages from

the failure of the vertices and edges of the model are researched. At the third stage, structural and functional risk assessments of failures by elements and IEC CTS are made. If it is possible to single out the structural or functional elements of the CTS, then the vertices of the directed graph correspond to the structural elements of the system, the edges of the directed graph correspond to the IEC. If it is difficult to select the elements, the vertices of the directed graph correspond to the parameters of the system, and the edges of the directed graph correspond to the cause-and-effect relationships between the parameters. Several edges can enter (exit) each vertex. Each edge is incident with two vertices located at its ends.

To assess the structural performance of the CTS and the associated damage, CM uses striking simulation impulses (SSI). The proposed SSI propagation procedure is close to the approach considered in [15], but it is simpler, there are no nonlinear operations. The procedure is easily formalized and transformed into a computational algorithm, which is important for CTS with numerous elements.

Uncertainty, incompleteness and fuzziness of information in assessing the structural and functional risk of failure of interrelated and interacting components designed and operated by CTS requires the ISDF to use the apparatus of a fuzzy probabilistic model. The developed structure of the fuzzy-probabilistic model includes input linguistic variables: D - damage, P - probability of failure and the output value R - risk of failure. The Takagi-Sugeno model is used to assess the risk of CTS failures based on a fuzzy probabilistic model. This allows us to obtain the final fuzzy set for the output variable of the failure risk assessment with a membership function of the form

$$\mu_{x_i}(f_{R_i, q_i}) = \max_{Y_R} [\mu_{x_P}(f_R), \mu_{x_D}(f_R), \mu_{x_{q_i}}(f_R)] , \quad (4)$$

where $\mu_{x_P}(f_R)$ - fuzzy set of element failure probabilities CTS; $\mu_{x_D}(f_R)$ - fuzzy set of damage from failures of CTS elements; $\mu_{x_{q_i}}(f_R)$ - fuzzy set of weights of the failure risks of CTS elements.

The production rules were entered on the basis of the Harrington generalized desirability function in accordance with the method of assessing the structural and functional risk of CTS failures. As a membership function, the Gauss distribution function was used, implemented in Matlab in the form of `gaussmf` to set smooth symmetric membership functions. At the fuzzification stage, the input variables of the fuzzy-probabilistic model for assessing the failure risk of designed and operated CTS are set in the form of the probabilities of failures and damages from failures of elements and EMI. To describe the key ontological information about heterogeneous datasets of large volumes, as well as to provide a more compact form of their presentation with subsequent processing and analysis, the structure of partitioning decision trees can be expressed as a set of logical rules in production form. This makes it possible to use such models as the basis for KB in ISDF. It is proposed to formalize the structure of the DT model using the CART algorithm. The use of the algorithm is due to its focus on building binary DTs, each individual node of which, in the process of logical partitioning, forms only two descendants. CART is implemented by splitting a lot of data by gender at each step. On one branch of the tree there are examples where the logical rule is executed successfully, and on the other it is not successful [1]. In the process of increasing the KB structure, at each DT node, an enumeration of all possible attributes is performed, separating only the one that maximizes the target indicator.

4. Experiments and results analysis

The process of conducting experiments was carried out on the basis of several stages:

- collection and extraction of data on the studied CTS structure, which is a transport engine cooling system [2] from the OREDA database [6];
- estimation of the probability distribution of CTS failure components based on normal distribution and failover statistics;
- allocation of individual categories of failures by priority (low, medium, high);
- data structuring, preparing and computer simulation with data visualization.

To study the DMIS of the risk of CTS failures in emergency scenarios, as an example, we studied a system consisting of 12 elements and 26 IEC (Fig. 4), 17 of which are energy carriers (indicated by

solid lines), 9 are carriers of matter (indicated by dashed lines). To determine the damage to the edges and vertices of the directed graph CM DMIS of the failure risk CTS, we set the weight values S_{vi} - for vertices (Table 1) and S_{aj} - for edges (Table 2). Values of the probabilities of failure of vertices and edges of the directed graph CM CTS $p_{vi}(t)$ and $p_{aj}(t)$ are presented in tables 1,2.

The results of the impact of the SSI on the CM CTS were established by research in the developed software package based on the cross-platform Python language. The structure of the CTS in the form of a directed graph is represented using the dot language. The technical state of the CTS elements in CM is reflected in the JSON format. Graphviz is used to visualize graphs. Analysis of the simulation results was carried out in MS Office and Open Office.

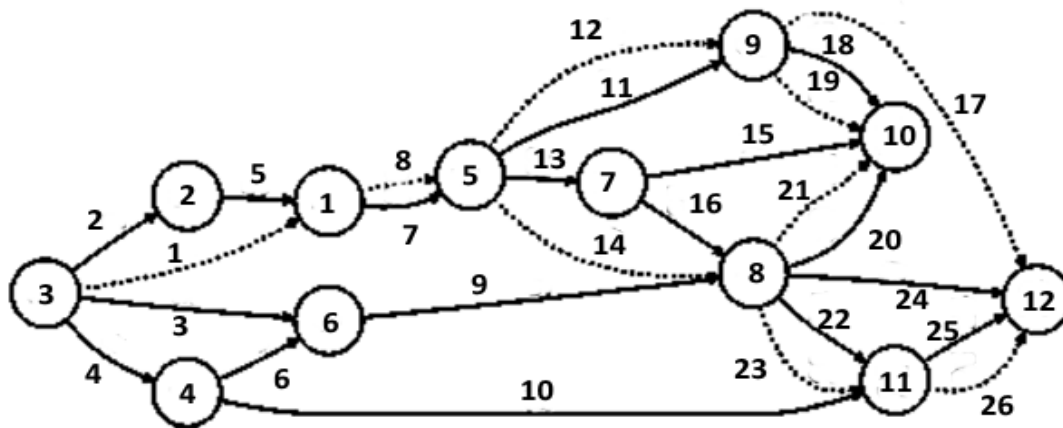


Figure 4: Oriented graph of CM elements and EMI CTS

Table 1

Values of the probabilities of failure of the vertices of the directed graph CM CTS

No. of the i-th verticle	$p_{vi}(t)$	No. of the i-th verticle	$p_{vi}(t)$
1	0,07	7	0,11
2	0,09	8	0,16
3	0,17	9	0,13
4	0,03	10	0,2
5	0,12	11	0,21
6	0,05	12	0,26

Table 2

The values of the probabilities of failure of the edges of the directed graph CM CTS

No. of the j-th edge	$p_{aj}(t)$	No. of the j-th edge	$p_{aj}(t)$	No. of the j-th edge	$p_{aj}(t)$
1	0,12	10	0,18	19	0,21
2	0,2	11	0,25	20	0,16
3	0,15	12	0,35	21	0,17
4	0,17	13	0,14	22	0,11
5	0,11	14	0,17	23	0,06
6	0,41	15	0,09	24	0,23
7	0,32	16	0,29	25	0,38
8	0,21	17	0,23	26	0,26
9	0,29	18	0,19		

From the results of the research it follows that the place of application of the SSI has a significant impact on the process of distribution of the SSI along the directed graph CM CTS. Taking into

account such structural features of a directed graph as connectivity, the presence of contours, the vulnerability of its vertices and the type of EMI resource, the most connected components of a directed graph, which are affected by a number of EMI of various resource types, have been identified.

The values of estimates of the structural risk of failures established by modeling for the vertices of the CM CTS and their ranking are shown in Fig. 5. In the course of research, it has been confirmed that if the place of application of the SSI does not belong to the structural component of strong connectivity, then the passage of the SSI is less associated with the risk of CTS failure. In the studied CTS, the components of strong connectivity are elements 3, 8 and IEC 2, 5, the performance of which determines the risk of failures of adjacent elements and IEC. If IEC 2 is damaged, risks arise during the operation of elements 1, 5, 7-12, and if element 3 is damaged, risks arise for all CTS elements.

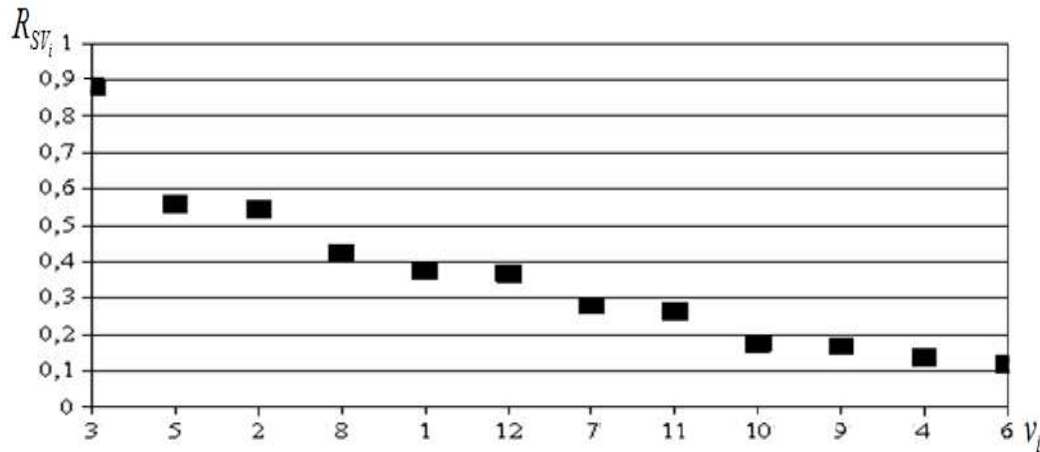


Figure 5: Ranking of the values of the structural risk of failures R_{SV_i} by elements v_i CTS

To effectively store data for different types of system elements, an add-on was created for distributed data warehouse based on the use of Microsoft Azure Data Lake Storage Service, which will allow scaling the system's capabilities.

The developed CM reflects the direct dependence of the risk of failures of elements and IEC on their positions in the system structure, as well as the dependence of the entire CTS on the selected topology at the design stage. The advantage of the method is its simplicity, clarity and applicability for assessing the risk of failures of a wide class of CTS. The procedures of the method implementation are easily formalized and transformed into a computational algorithm and model for assessing the risk of failure, which is important for CTS with a large number of elements and IEC. Considering that the objects of technical condition diagnostics often consist of various CTS, for some of which the use of CM for assessing the risk of element failures and IEC is difficult due to the uncertainty, incompleteness and fuzziness of information received by the diagnostic systems and predicting their technical condition, it is advisable in DMIS additionally in addition to cognitive modeling, use fuzzy probabilistic models. An example of such a CTS can be a ship power plant (SPP), the diagnostic information from which is characterized by uncertainty, incompleteness and indistinctness [30]. As a fuzzy model, the Takagi-Sugeno model is used, the advantages of which include a more compact size of the rules repository compared to the Mamdani fuzzy inference models, as well as lower computational complexity due to a simpler defuzzification procedure [24,37]. Fuzzy rules have the structure “If damage or probability then risk”, where damage and probability are the linguistic value of the antecedent, determined by the fuzzy membership function, which is a consequent stochastic variable equal to one of the values in the range from 0 to 1.

The result of the fuzzification stage is the establishment of a correspondence between the specific value of a separate input variable of the fuzzy inference system and the value of the membership function of the corresponding term of the input linguistic variable.

The aggregation operation is based on the maximization method implemented in Matlab as max. The weight values for all input variables of the fuzzy-probabilistic model were taken equal to one.

During accumulation, the max-disjunction method was used to combine all degrees of truth of the conclusions of the rules. The developed fuzzy-probabilistic model contains 25 logical rules. Defuzzification was carried out using the center of gravity method implemented in Matlab.

The damage to the elements and IEC of the SPP was set using the linguistic variables Not significant (0-0.25), Low (0.12-0.36), Medium (0.3-0.6), High (0.55-0.8), Critical (0.7-1). On the basis of the developed fuzzy-probabilistic model, it is possible to minimize the time of assessing the risk of failures of the designed and operated SPP. This is achieved due to the mechanism of fuzzy rules incorporated into the developed model for assessing the risk of failure of the SPP using the Matlab Rule Viewer module. Such a mechanism makes it possible to identify and quantitatively reflect the degree of influence of specific values of damage and the elements and IEC probabilities failure on the risk of failure of the SPP as a whole. In particular, with the value of damage equal to 0.234 and the value of the probability of failure equal to 0.5, the value of the output variable of the risk of failure was 0.157, i.e. about 16%, which testifies to the low degree of danger of the SPP efficiency loss.

The probability of SPP elements and IEC failure was set using the linguistic variables: Not significant (0-0.15), Low (0.1-0.31), Average (0.22-0.6), High (0.45-0.75), Critical (0.6-1). The terms Minimum risk (0-0.2), Acceptable risk (0.1-0.35), Significant risk (0.3-0.65), Critical Risk (0.63-1). In the course of modeling, a three-dimensional visualization model's surface for assessing the risk of failure of elements and IEC of the SPP was obtained (Fig. 6).

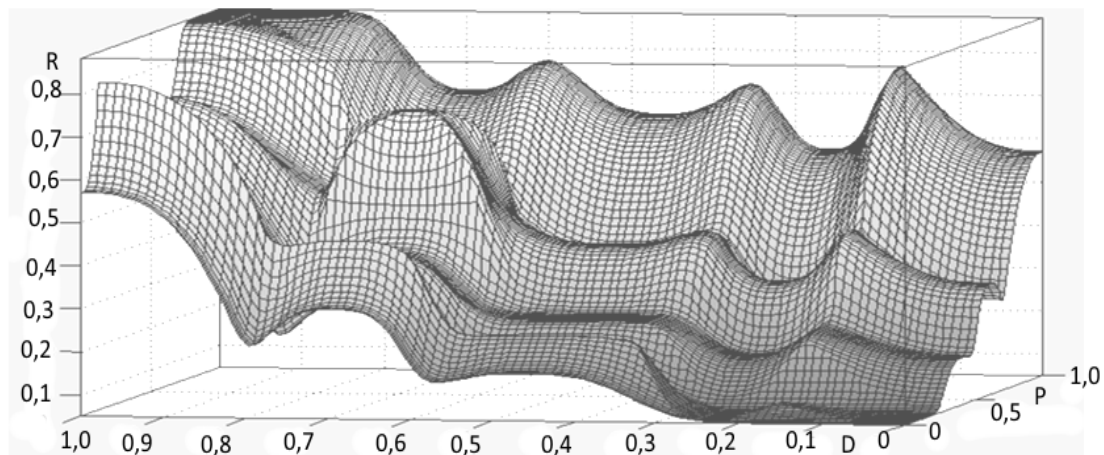


Figure 6: Visualization of the results of the fuzzy-probabilistic model of risk assessments

The results fuzzy-probabilistic model studying make it possible to establish that the probability of failure of elements and EMI has a more significant effect on the assessment of the risk of failure of the EMI.

According to the obtained three-dimensional visualization fuzzy-probabilistic assessment model's surface the greatest increase in the level of risk of EEC is observed with values of the probability of failures in the range from 0.455 to 1. In order to use KB for possible scenarios of actions when making control decisions and forming a visual structure for a data set from CTS, the DT graph is used (Fig. 7).

The composition of the logical rules records for KB based on the interpretation of this model can be expressed as: IF StructRisk > 0.435 AND Average cost ≤ 1600 AND Maintainability = "High" THEN Damage is Middle. Additional clarification of this rule is possible by using the priorities of features or weight values (degree of significance) for each of them. An example of a selection of options for the values of input and output parameters obtained from the results of cognitive and fuzzy-probabilistic modeling is shown in table. 3.

Table 3 is supplemented with a column of logical inference obtained based on the use of KB for possible target alternatives (action scenarios) when making management decisions out of 4 possible options: prevention, repair, replacement, no action). The received recommendations can be assessed by a decision-maker to make changes to the KB structure both in manual and automatic mode, which provides feedback to the model and its adaptability

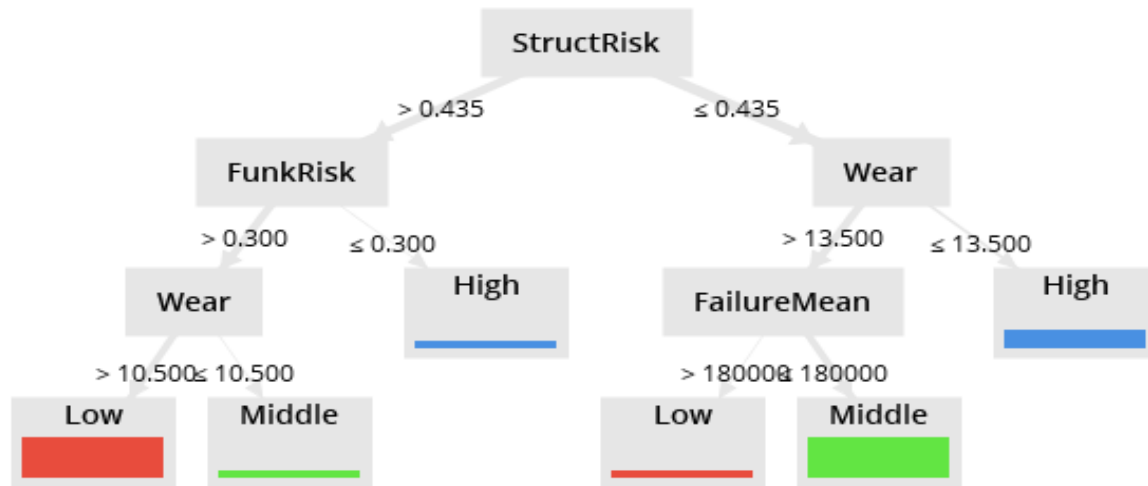


Figure 7: DT model visualization based on logical rules

Table 3

Selections of options for values of input and output parameters

No	Probability of failure	Damage	Risk	Most Priority Action Scenario
1	0,078	0,21	0,016	no action
2	0,194	0,3	0,058	prevention
3	0,29	0,31	0,09	prevention
4	0,354	0,37	0,13	prevention
5	0,398	0,08	0,032	no action
6	0,462	0,21	0,1	prevention
7	0,107	0,64	0,07	prevention
8	0,106	0,23	0,025	no action
9	0,634	0,63	0,4	replacement
10	0,794	0,43	0,34	repair

5. Conclusion

A data mining system has been developed for diagnostics, predicting failure risk assessments of designed and operated interconnected and interacting elements, inter-element connections of complex technical systems, based on the application of the method and cognitive-imitation, fuzzy-probabilistic models of failure risk assessments.

The knowledge base of the proposed system is based on the methodology of cognitive and fuzzy-probabilistic modeling, the decision tree model, which makes it possible to take into account the scenarios of the development of the risk of failure of elements and inter-element connections of systems.

The use of the developed methods and models in the intellectual system makes it possible to assess the reliability of systems based on the results of the effects of the damaging effect of each element in emergency scenarios on the structure of the system, as well as to predict the consequences of failure of elements and inter-element connections of systems. The practical application of the developed system can be performed by automating the tasks of collecting and processing data from various CTS, especially in the transport sector.

The use of machine learning algorithms in the created system makes it possible to form linked and structured metadata sets based on risk assessments of failure of elements and interelement links in various scenarios.

Application of the developed method and models will make it possible to make effective management decisions to ensure the reliability of systems, both in the design and in the operation of complex technical systems.

6. References

- [1] N. Rudnichenko, V. Vychuzhanin, I. Petrov, D. Shibaev, Decision Support System for the Machine Learning Methods Selection in Big Data Mining, in: Proceedings of The Third International Workshop on Computer Modeling and Intelligent Systems (CMIS-2020), CEUR-WS, 2608, 2020, pp. 872-885.
- [2] V. Vychuzhanin, N. Rudnichenko, Z. Sagova, M. Smieszek, V. Cherniavskiy, A. Golovan, Analysis and structuring diagnostic large volume data of technical condition of complex equipment in transport, in: 24th Slovak-Polish International Scientific Conference on Machine Modelling and Simulations - MMS 2019, Liptovský Ján, Slovakia.
- [3] J. Drozd, A. Drozd, S. Antoshchuk, V. Kharchenko, Natural Development of the Resources in Design and Testing of the Computer Systems and their Components, in: 7th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, Berlin, 2013, pp. 233–237. doi: 10.1109/IDAACS.2013.6662656
- [4] Y. Sekretarev, V. Levin, Risk-Based Models of Equipment Repair Management in Power Supply Systems with a Mono Consumer, Journal of Siberian Federal University. Engineering & Technologies, 14 (2021) 17-32. doi: 10.17516/1999-494X-0295
- [5] ISO 13379-1:2015(E) Condition monitoring and diagnostics of machines - Data interpretation and diagnostics techniques (2015).
- [6] M. Rausand, A. Barros, A. Hoyland, System Reliability Theory: Models, Statistical Methods, and Applications, Third Edition, Wiley, 2020
- [7] N. Rudnichenko, V. Vychuzhanin, V. Mateichyk, A. Polyviachuk, Complex Technical System Condition Diagnostics and Prediction Computerization, in: Proceedings of The Third International Workshop on Computer Modeling and Intelligent Systems (CMIS-2020), CEUR-WS, 2608, 2020, pp.1 -15.
- [8] L. Zil'berburg, V. Melochnik, Ye. Yablochnikov, Information technologies in design and production, Politekhnik, Saint-Petersburg, 2018.
- [9] Z. Saleh, Artificial Intelligence Definition, Ethics and Standards, The British University in Egypt, 2019.
- [10] A. Lapina, Intelligent information systems: Textbook. 2017. URL: http://files.lib.sfu-kras.ru/ebibl/umkd/228/u_course.pdf.
- [11] V. Novorusskiy, Expert systems for solving problems of diagnostics, forecasting and management of the functioning of energy systems, etc. (approaches to synthesis). Irkutsk, SEI SO RAN, 2015.
- [12] T. Kravchenko, N. Naumova, Modern information technologies in the development of computer decision support systems (2020). URL: <http://www.mesi.ru/ksit/k4sem24.zip>
- [13] L. Kovtun, N. Sharkov, Methods of simulation modeling and situational analysis of management decisions in ship accidents based on linguistic description of processes, algebraic descriptions and neural networks, in: IMMOD-2019, 2019, pp.139–145.
- [14] R. Shennon, Simulation of systems-art and science, Mir, Moscow, 1978.
- [15] V. Vychuzhanin, N. Rudnichenko, N. Shibaeva, Y. Kondratenko, Cognitive-impulse model for assessing complex technical systems survivability, in: Proceedings of the 9th International Conference Information Control Systems & Technologies (ICST-2020), CEUR-WS, 2711, 2020, pp.571-585
- [16] M. Beetz, M. Buss, D. Wollherr, Cognitive technical systems—what is the role of artificial intelligence, Advances in Artificial Intelligence 4667 (2017) 19–42.
- [17] C. Schwenk, The cognitive perspective on strategic decision making, Journal of Management Studies, 25 1 (2017) 41–55.
- [18] R. Yenikev, Knowledge base for the design of internal combustion engines, SAE Transactions 1 (2017) 15–20.

- [19] A. Zykov, Fundamentals of graph theory. University book, Moscow, 2014.
- [20] S. Shtovba, O. Pankevich, A. Nagorna, Analyzing the criteria for fuzzy classifier learning, Automatic Control and Computer Sciences, 49 (2015) 123–132.
- [21] A. Rotshtein, Selection of Human Working Conditions Based on Fuzzy Perfection, Journal of Computer and Systems Sciences International, 57 (2018) 927–937.
- [22] YU. Kudinov, Fuzzy input models in expert systems, Izv. Theory and Control Systems, 5 (2017) 75–83.
- [23] Zhedunov R.R. The system of identification of pre-emergency situations of the technological process, using the apparatus of fuzzy logic and data of probabilistic deviations, Bulletin of the Astrakhan State Technical University, 38 (2017) 169–173.
- [24] G. Luger, Artificial intelligence: structures and strategies for complex problem solving, Addison-Wesley Publishing Company, United States, 2019.
- [25] D. Tupikov, A. Rezhnikov, An approach to building a fuzzy knowledge base for determining fire-hazardous situations, Mathematical methods in engineering and technologies, 62 (2014) 125-127.
- [26] V. Nechayev, M. Koshkarev, Architecture of a pre-oriented knowledge base of an intelligent system, Educational resources and technologies, 1 (2015) 156-163.
- [27] Yu. Antokhina, V. Balashov, E. Semenova, A. Varzhapetyan, Computer simulation of processes in technical systems, Journal of Physics: Conference Series. 1691 (2020) doi:10.1088/1742-6596/1691/1/012069
- [28] V. Devyatkov, Methodology and technology of simulation studies of complex systems: current state and prospects of development, Vuzovsky textbook, Moscow, 2014.
- [29] I. Nazarenko, M. Delembovskyi, O. Dedov, A. Onyshchenko, I. Rogovskii, M. Nazarenko, I. Zalisko, Study of reliability of technical systems reliability, Dynamic processes in technological technical systems 14 (2021) 110-139. doi: 10.15587/978-617-7319-49-7.ch7

Method for Processing and Assessing the Degree of Digital Image Compression Based on Haar Transformation

Vladimir Vychuzhanin^a, Nickolay Rudnichenko^a, Yurii Bercov^b, Andrii Levchenko^b

^a Odessa National Polytechnic University, Shevchenko Avenue 1, Odessa, 65001, Ukraine

^b I.I.Mechnikov Odessa National University, Nevskogo street 36, Odessa, 65016, Ukraine

Abstract

The article presents the specifics of the developed method for processing and assessing the degree of digital image compression based on Haar transformation. In order to ensure a sufficient level of quality for the presentation of digital images of different formats and classes in the process of their multi-threaded processing in real time, it is proposed to implement a number of orthogonal transformations. The mathematical representation of the created method reflects the stages of restoration of individual image indicators in the process of its reception and transmission. The proposed method is based on the principles of complex adaptive compression using an estimate of the specific global and local sensitivity of the Haar transform functions. The developed method can be used for adaptive compression of digital images of different saturation with support for dynamic scaling of their sizes based on the estimation of the energy index in the analysis of the image transform, providing detailing of the degree of its saturation.

Keywords

Digital image processing, images compression, energy indicator, image saturation, Haar transformation.

1. Introduction

Currently, the flow of various data of various formats from distributed information and technical systems is actively increasing. Graphic data has a high priority in this specificity, in particular, images of various resolutions and formats, which allow us to detail various aspects of the state of the studied or analyzed components of dynamic systems through their automated intelligent recognition using applied methods of pattern recognition theory [1]. Meeting the performance requirement for processing and transmitting digital images in real time is based on reducing the amount of data associated with images by compressing them. However, the efficiency of image compression significantly affects the quality of the reproduced video information. The use of currently known digital image compression methods allows: reduce the time required to transfer images; to increase the number of images transmitted over the communication channel per unit of time; to reduce the requirements for the speed of information transmission devices; reduce the storage capacity of images [2]. Remote sensing of the earth's surface using unmanned aerial systems (UAC) is widely used, for example, to monitor vegetation and environmental parameters, aimed in particular at optimizing activities in agriculture and forestry. In this context, UACs have become useful for assessing crop health by acquiring large amounts of raw data that require processing to further support applications such as moisture, biomass, and others [3]. The analysis of known works indicates a constant increase in the amount of information received from the UAC during remote sensing of the earth's surface, which outstrips the growth rate of the capabilities of technical means for its processing [4].

CMIS-2022: The Fifth International workshop on Computer Modeling and Intelligent Systems, May 12, 2022, Zaporizhzhia, Ukraine.

EMAIL: vint532@yandex.ua (V.Vychuzhanin); nickolay.rud@gmail.com (N.Rudnichenko); bercov@gmail.com (Y. Bercov); KatyaAndreyVel@gmail.com (A. Levchenko)

ORCID: 0000-0002-6302-1832 (V.Vychuzhanin); 0000-0002-7343-8076 (N.Rudnichenko); 0000-0001-9704-1121 (Y. Bercov); 0000-0001-5550-0027 (A. Levchenko)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

In the process of conducting remote sensing of the earth's surface, one of the ways to solve the problem of operational processing and transmission of images in real time is to reduce the amount of data by compressing images.

The image compression method is understood as a set of actions that allows us to unambiguously match a set of compressed data to the original dataset [5].

Currently, many methods have been developed for compressing digital images, which can reduce this or that redundancy. Nevertheless, for this purpose, intensive development of new methods of compression of digital images continues [4].

Analysis of digital image compression methods shows that most of them are complex, consisting of stages: change of color model, image conversion, image encoding [5].

An effective and simple way to increase the compression ratio provided by image compression algorithms is to select the optimal color model in which the image data is represented. Compression algorithms occupy an essential place in the theory of digital image processing.

This is due to the fact that images presented in digital form require a large amount of memory to store, and when transferring images through communication channels, it takes a long time [6].

2. Literature review and description of problem

To transform a digital image in modern practice according to a number of literary sources [7-9], an orthogonal or wavelet transform is used. Orthogonal transformations consist of direct and inverse transformations [4-6]. Discrete forward and backward transformations in general:

$$Y(k) = \left\langle \frac{1}{N} \right\rangle \cdot \sum_{m=0}^{N-1} X(m) \cdot W(k, m), \quad k = \overline{0, N-1} \quad (1)$$

$$X(m) = \sum_{k=0}^{N-1} Y(k) \cdot W(k, m), \quad m = \overline{0, N-1} \quad (2)$$

where:

- $X(m)$ – m -th sample of the original discrete signal;
- $Y(k)$ – k -th conversion coefficient;
- $W(k, m)$ – m -th sample of the k -th basis function of the orthogonal transformation;
- N – the number of samples in the original signal or in the block of the original signal during block processing;
- $\langle 1/N \rangle$ – normalization coefficient, which may be absent in some orthogonal transformations (for example, a discrete-cosine transform).

The entire image or its fragments is subjected to conversion. A typical image fragment consists of 8×8 , 16×16 or 32×32 samples. A fragment size selection is due to the fact that the correlation interval for images does not exceed 8 ... 32 samples.

Analysis of existing varieties of image compression methods shows that most modern compression methods consist of several stages.

The first step in most compression methods involves changing the color model. Currently, various color models are used, which are described in sufficient detail, for example, in publications [7]. The high efficiency of using a color model change as a stage of image compression is confirmed by the experience of operating existing compression methods [8].

However, the use of this stage leads to a significant increase in the computational complexity of the compression algorithm (by $20 \div 40\%$), which is associated with the need to carry out about $9 \times n$ real multiplication operations and $6 \times n$ real addition operations (where n is the number of samples in the image).

The second stage - transformation of the image (for example, orthogonal or wavelet) allows you to translate the original signal from the space-time domain into the spectral-frequency domain. In this

case, the original signal (function of time) can be represented through an orthonormal set of spectral functions in the form of a series.

Representation of signals in the form of a set of spectral functions allows their spectral analysis, convolution of complex signals, processing in the spectral domain with decorrelated spectral signal elements, etc.

In addition, the use of transforming the incoming sequence of image samples allows redistributing the image energy.

Using the transform of the processed image provides further compression procedures with decorrelated transform coefficients.

To assess the effectiveness of transformations in [9], the following particular indicators are used:

- $K_{sum/sub}, K_{sum/sub}^{(-1)}$ – the number of addition / subtraction operations in direct and inverse transformations, respectively;
- $K_{mul/div}, K_{mul/div}^{(-1)}$ – the number of multiplication / division operations for direct and inverse transformations, respectively;
- T_{op} – the type of operations (operands) used in the conversion procedure;
- S_{coef} – sensitivity of conversion factors;
- D_{tr} – dynamic range of transformation ratio values;
- σ – root-mean-square deviation during conversion.

The standard deviation (SD) is a quantitative indicator and characterizes the error of the reconstructed image relative to the original one. SD is determined by the following expression

$$\sigma = \sqrt{\frac{\sum_{i=1}^N \sum_{j=1}^M (x_{i,j}^{(e)} - x_{i,j}^{(u)})^2}{N \cdot M}}, \quad (3)$$

where

- $x_{i,j}^{(e)}, x_{i,j}^{(u)}$ – i, j-th sample of the reconstructed and original images, respectively;
- N, M – the dimension of the image horizontally and vertically.

Conversion coefficient sensitivity S_{coef} divided into local and global [10]. With global sensitivity, the transformation coefficient is a function of all coordinates of the input sequence space, and for local sensitivity, only a certain part of the coordinates. Accordingly, to calculate the conversion factor with global sensitivity, it is necessary to perform more arithmetic operations than for the coefficient with local sensitivity. The sensitivity of the conversion coefficients is determined by the used orthogonal basis.

Currently, the following bases are widely used: discrete cosine transform (DCT), discrete Walsh transform (DWT), Haar transform (HT) [11] and various wavelet transforms (WT).

The complexity of transformations is greatly reduced when using separable orthogonal transformations having fast computational algorithms. These include Fourier transforms, Haar transforms, and cosine transforms, which are most useful for compressing image data.

The main property of the discrete cosine transform is that its basis vectors approximate very well the eigenvectors of Telltz matrices.

In its decorrelating properties, it approaches the Karunen-Loeve transform, while retaining the capabilities of fast Fourier algorithms. This explains the use of discrete cosine transform in JPEG-formats of representation and transmission of images.

The most efficient, from the point of view of computational costs per one transformation ordinate, is the orthonormal system of piecewise constant functions of the Haar basis. The Haar basis possesses the locality property, which underlies the modern theory of wavelet transforms.

When converting an image, the standard deviation (SD) is used with a quantitative indicator (σ) corresponding to the errors of the reconstructed image in relation to the original image. In fig. 1 shows the corresponding values of the transformation efficiency indicators for the orthogonal basis [6,14-16].

Thus, Fig. 1 illustrates the dependence of the root-mean-square deviation σ when using the analyzed transformations for digital images of three classes: $k=1$ - monotonic, $k=2$ - of the "portrait" type; $k=3$ - with a large number of contours.

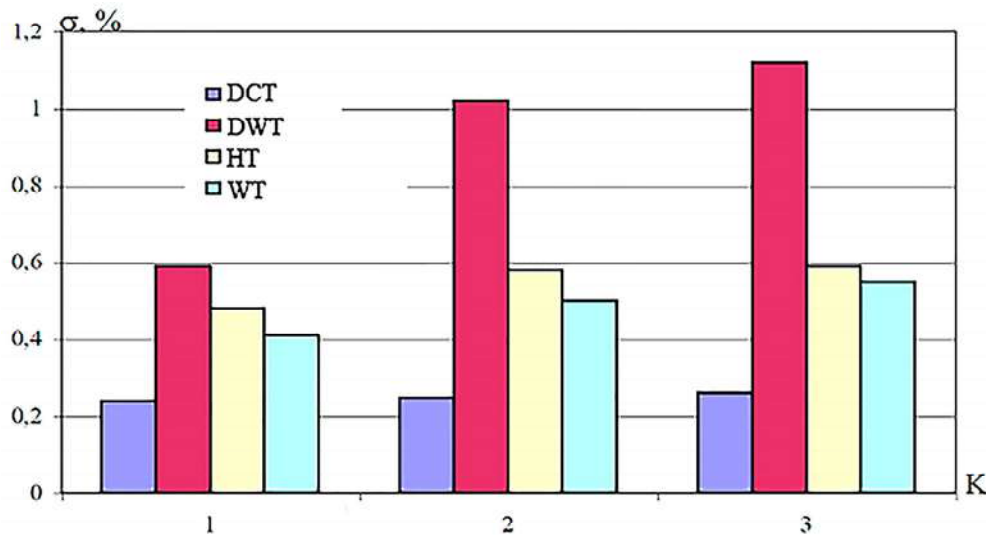


Figure 1: SD dependence when converting images samples with different saturation degrees

As follows from Fig. 1, the Haar transforms have the best values of efficiency indicators. Despite the small SD values provided by DCT and WT, their significant disadvantage is the increased dynamic range.

This leads to the need to perform quantization, as a result of which the SD is increased. One of the stages of digital image compression is image encoding.

As a result of encoding a digital image, a compressed sequence is formed using fixed - length encoding methods.

When compressing digital images, for which the amount of information plays an important role, variable length codes are used.

A distinctive feature of optimal codes is that they do not lead to expansion of the compressible data.

According to the data in Fig. 2, the results of optimal coding (shown by the dotted line) and Huffman coding (shown by the solid line) can be compared. $C(f, S)$ is the source S coding cost f , $p(s)$ is the relative symbol coding frequency.

It can be seen from the graphs that the Huffman coding redundancy approaches optimal only when the relative symbol coding rates are multiples of two.

The advantages of the orthogonal transformation of the original image according to the Haar basis include low computational complexity with a relatively low SD.

Analysis of the efficiency of two-dimensional transformations: DCT, DWT, HT and WT in the Haar basis (Table 1) allows us to single out the HT and DWT transformations as the most promising in terms of the presented indicators, provided that the requirements for the efficiency of image processing are met.

Thus, from the analysis of digital image compression procedures, it follows that in order to meet the requirements for the efficiency of image processing and transmission when implementing methods for compressing them with an acceptable level of distortion, it is necessary to carry out an orthogonal

transformation of the original image. Table 1 also shows the values of the proposed performance indicators for the considered transformations [10,12-14]. The use of WT according to the Haar basis provides the smallest increase in the dynamic range (1.5 times), when using other bases, the increase will be greater.

Despite the small SDs that DCT and WT provide, their significant drawback is the increased dynamic range, which leads to the need to perform quantization and, as a result, to an increase in SD.

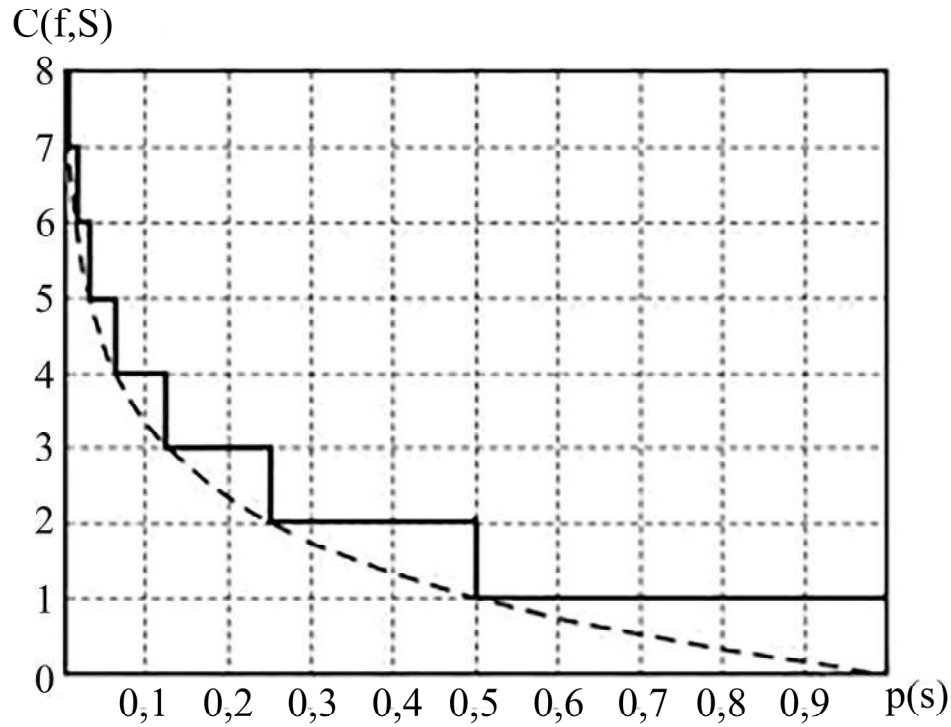


Figure 2: Dependence of coding redundancy for optimal coding - - and Huffman coding -

Table 1. Two-dimensional transform efficiency parameters: DCT, DWT, HT and WT in the Haar basis

Performance indicator		Conversion			
		DCT	DWT	HT	WT (L steps)
	$K_{sum/sub} / K_{sum/sub}^{(-1)}$	$4N^2 \log_2 N$	$2N^2 \log_2 N$	$2(N-1)$	$N^2 \sum_{k=0}^{L-1} (\frac{N}{2^k} - 1)$
	$K_{mul/div} / K_{mul/div}^{(-1)}$	$4N^2 \log_2 N$	N	N	$N^2 \sum_{k=0}^{L-1} (\frac{N}{2^k} - 1)$
	T_{op}	real	integer	partially real	real
	D_{tr}	increases by 2 - 3 times	not increasing	not increasing	increases by 1,5 times
	S_{coef}	global	global	global and local	local (with the Haar basis)

Analysis of the Table 1 shows that the Haar transforms have the best values of efficiency indicators. Existing lossy compression techniques such as JPEG, JPEG-2000, TIFF, WI and others use conversion as a preparatory step for quantization or selection.

At this stage, the search and elimination of transformation coefficients is performed, the contribution of which to the formation of the restored image is minimal. In this case, it is taken into account that to obtain an accurate approximation of the initial data, it is not required to use all the

transformation coefficients. To do this, the JPEG method uses a procedure called quantization. To quantize, simply divide the conversion factors by another number and round to the nearest integer. To determine the quantum value in JPEG, an array of 64 elements is used - a quantization table, in which values for each coefficient of the processed 8×8 image block are defined.

Another approach to finding and eliminating transform coefficients can be the selection of transform coefficients obtained as a result of the transformation transform. Selection is based on the zonal, threshold and zonal-threshold principle [4,15,16]. In zonal selection, only those coefficients are selected that are in a predetermined zone (usually in the lower spatial frequency region). In the case of threshold selection, coefficients are selected that exceed a certain threshold level [16]. Analysis of modern works [17,18] shows that zonal-threshold selection is a combination of zonal and threshold selection and allows you to effectively use the coding procedure in the future.

The final step in most compression methods is encoding, which results in a compressed sequence. It is known that character sets in digital image processing are almost always natively represented by the use of encoding methods like EBCDIC or ASCII, which can't implement the counting characters frequency. Moreover, each symbol is encoded with a constant number of bits, which ensures a high computation speed. In image compression applications, where one of the main criteria is information amount, variable length codes are used.

The analysis of works [4,19] showed that the development and widespread use of existing compression methods received statistical coding algorithms. The processing mode for generating such codes for a list of values, based on their frequency's values, is simple and may include two stages: modeling and coding. If the elements probability distribution $p(s)$ generated by the source is known, then the length of the character codes of the alphabet is proportional $-\log p(s)$. However, in most cases, the statistical source model is not known, therefore, it is necessary to build a source model that would allow us to estimate the probability of each element occurrence at each input sequence position.

There are some algorithms [20], which form the source model as the data stream is processed or use a fixed model based on a priori ideas about the data nature. The encoding procedure effectiveness depends on the following parameters: entropy $H(S)$ coding source S (for example, according to the

Bernoulli scheme $H(S) = \sum_{i=1}^n p(s_i) \log_2 p(s_i)$; cost $C(f, S)$ source S coding f ; coding

redundancy $R(f, S)$. A hallmark of optimal codes f_0 is that they do not expand the compressed data in the worst case and provide such an encoding cost $C(f_0, S)$ at which the inequality $R(f_0, S) \leq R(f, S)$ holds for any coding f [21]. By Shannon's theorem, the best compression when represented in binary form can be obtained by encoding characters with a relative frequency $p(s)$ using the $-\log_2 p(s)$ bit [22]. The redundancy of Huffman coding approaches optimal only when the relative frequencies are multiples of two. It was shown in [19,20,23] that the use of arithmetic coding gives results close to optimal results. Thus, the existing compression methods are a set of procedures, the implementation of which is aimed at reducing various types of redundancy in the representation of images. The analysis of the considered compression methods showed that ensuring the efficiency of processing and transmission of images of the earth's surface with UAC at an acceptable level of their distortion is an urgent task. To solve it, it is necessary to carry out: orthogonal transformation of the original image according to the Haar basis, the advantage of which is low computational complexity with a relatively low SD; carrying out zonal or zonal-threshold selection of conversion coefficients, which allows satisfying the requirements for adaptability and interactivity of the developed image compression method; coding of the selected sequence using optimal codes.

3. Method for assessing the degree of images saturation

To achieve our goal, we will use a direct two-dimensional transformation of images on the Haar basis.

Direct and inverse two-dimensional transformation of images on the Haar basis:

$$y_{k,p} = K_{norm} \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} x_{i,j} h_{k,p}^{(1)}(i, j) \quad (4)$$

$$x_{k,p} = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} y_{i,j} h_{k,p}^{(-1)}(i, j), \quad (5)$$

where:

- $i, j, k, p = \overline{0, (N-1)}$ – the coordinates of the image reference location or its transform coefficient in the block $N \times N$;
- $x_{k,p}$ - image block count $X(n)$;
- $K_{norm} = 1/N^2$ - normalization coefficient;
- $y_{i,j}$ - Haar transform coefficient (transform element $Y(n)$);
- $h_{k,p}^{(1)}(i, j), h_{k,p}^{(-1)}(i, j)$ - i, j -th element of k, p -th submatrix $H_{i,j}^{(1)}$ direct transformation matrix $H_{np}^{(2)}(n)$ and submatrix $H_{i,j}^{(-1)}$ inverse transformation matrix $H_{op}^{(2)}(n)$ respectively.

The performed analysis of the transformation transformants showed that to reduce the dynamic range of the transformant, it is advisable to use subquantization of the coefficients in the direct transformation. Conversion factors should be grouped into selection zones that include factors with the same sensitivity.

According to the selection rule, the configuration of the selection zones of the Haar transformants has the form shown in Fig. 3. (the solid line shows the boundaries of the transformant zones, and the dashed line shows the transformation coefficients in the zones).

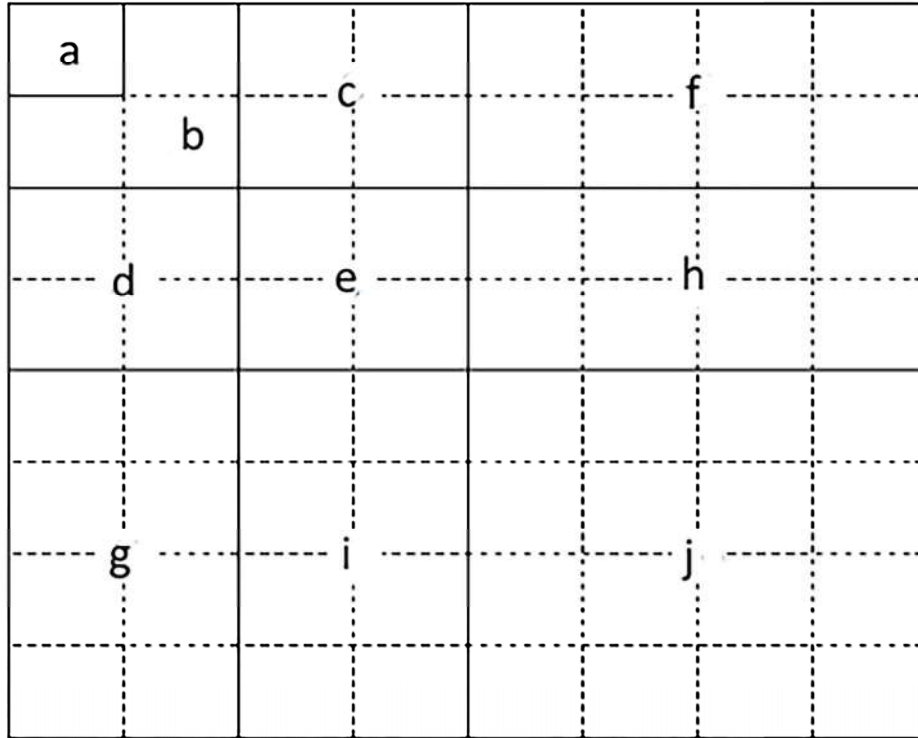


Figure 3: SD dependence when converting images samples with different saturation degrees

The coefficients of the S -th ($s \in a, j$) selection zone of all transformants of N blocks forms the S -th zonal sequence (ZQ). So, for example, f zone ZQ consists of the coefficients f of the selection zones of transformants of all image blocks. Conversion coefficients are grouped into selection zones containing coefficients with the same sensitivity. During subquantization, the value of the transformation coefficients is quantized with less accuracy in relation to the samples of the original digital image.

4. Experiments and results of original image orthogonal transformation research

When conducting a study using the method of adaptive compression of digital images in real time, the distribution of the number of repetitions of the values of the coefficients ZQ $N(y_i)$ was considered using the example of histograms of 3, 5 and 10 zones for test images shown in Fig. 4. The significant differences in the ZQ histograms, which are composed of the all blocks selection zones, for different images suggest the possibility of classifying images according to the energy distributed in them.

$$E = \sum_{s=1}^Z E_s = \sum_{s=1}^Z \sum_{k=1}^{K_s} C_{k,s}^2, \quad (6)$$

where E_s – energy S -th transform selection zones; Z – number of breeding zones; $C_{k,s}$ – k -th coefficient S -th transform zone; K_s – coefficients number $C_{k,s}$ at S -th zone, nonzero. It is proposed to use the weighted energy parameter S -th of the zonal sequence as an indicator that allows us to quantify the image energy distribution over ZQ

$$Q_s = E_s / K_s. \quad (7)$$

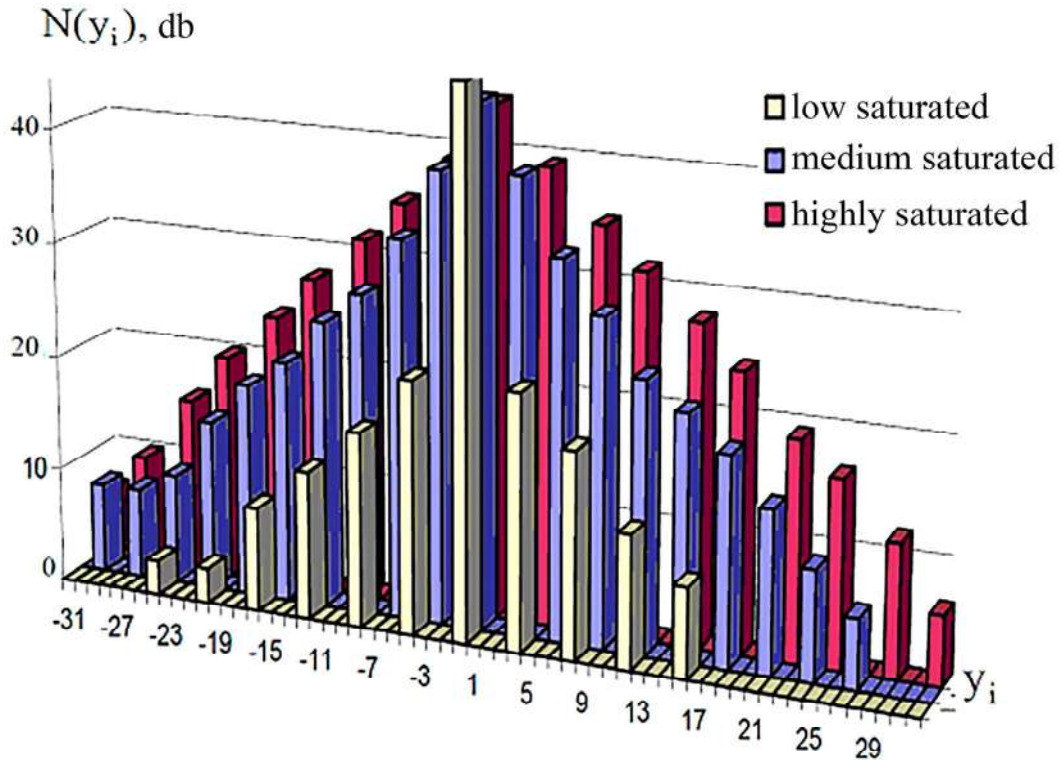


Figure 4: Histograms of the two-dimensional Haar transformation coefficients values distribution ZQ: 10th ZQ

The analysis of ZQ properties showed that their energy depends on the number of nonzero samples of the Haar basis function, which take part in the ZQ coefficients formation, and on the correlation value of the original image samples. Therefore, with an increase in the zone number, zonal energy decreases, and the magnitude of this energy reflects which particular details prevail in the image. An experiment was conducted to accumulate statistics on the energy distribution over zonal sequences for different images.

The original 300 images were used in BMP (bitmap) format with a visualization parameter of 24 bits per pixel.

The results of the zonal sequences average energy estimation $\bar{E} = \frac{\sum_{i=1}^{300} E_{s,i}}{300}$ are presented in fig. 5.

It is proposed to exclude zonal sequences according to the following rules: at $E_s < \bar{E}_s$ S -th ZQ considered priority for exclusion.

For example, when processing images with blocks 8×8 , W ZQ exceptions are proposed (W is the number of ZQ that must be excluded in order to obtain the necessary compression ratio K_{com} , starting from the Z -th) in the following order: 10-th, 9-th and 8-th, 5-th, 7-th and 6-th, 4-th and 3-rd.

The choice of this order is due to the zonal sequence influence, which is excluded, on the change $\Delta\sigma$ restored image; if necessary, eliminate ZQ, whose energy $E_s > \bar{E}_s$, the exception is in the same order.

Because of the values of the coefficients of the selection zone reflect the presence of details in the image of the corresponding size, then, for example, according to the histogram of the 10th ZQ of the third test image, it can be argued that there are more small details in comparison with other test images. According to the histograms of the first image, it can be assumed that there are large areas in the image with a smooth color change. It can be seen that with ZQ number increasing the dynamic range decreases, and the percentage of coefficients with a zero value reaches 40–95%, depending on the processed digital image.

Significant differences in the ZQ histograms, compiled from the selection zones of all blocks, for different images allow for the classification of images. Fig. 6 shows the compression ratio dependence K_{com} from the number of excluded selection zones W .

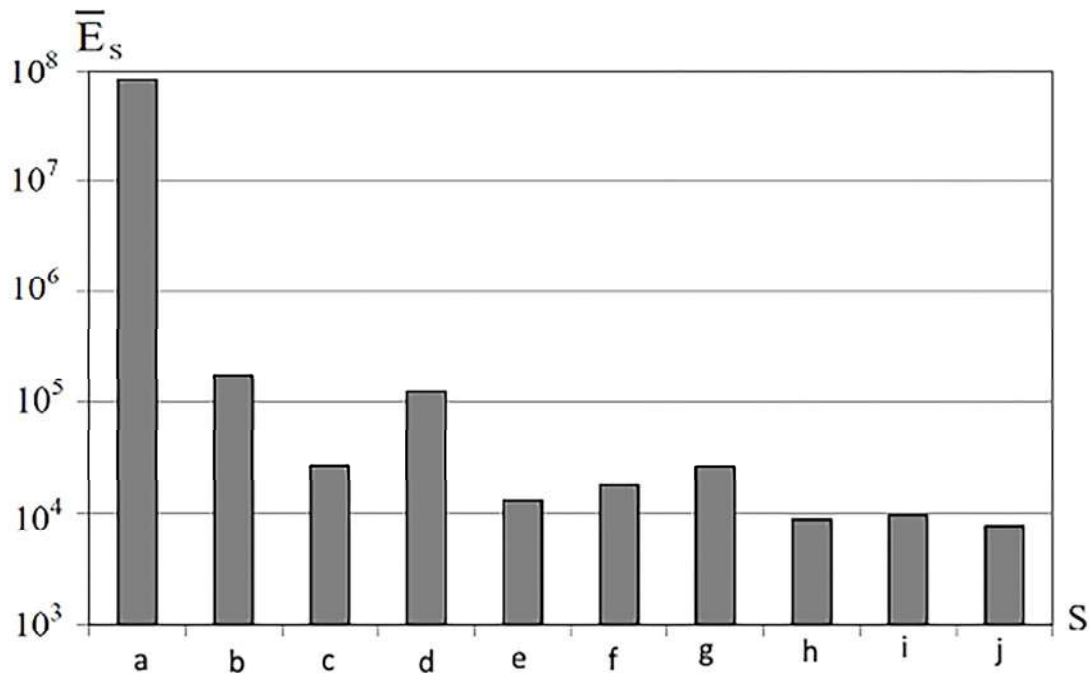


Figure 5: Histogram of arithmetic mean values of energy $\bar{E}_s$ S -th areas of analyzed images

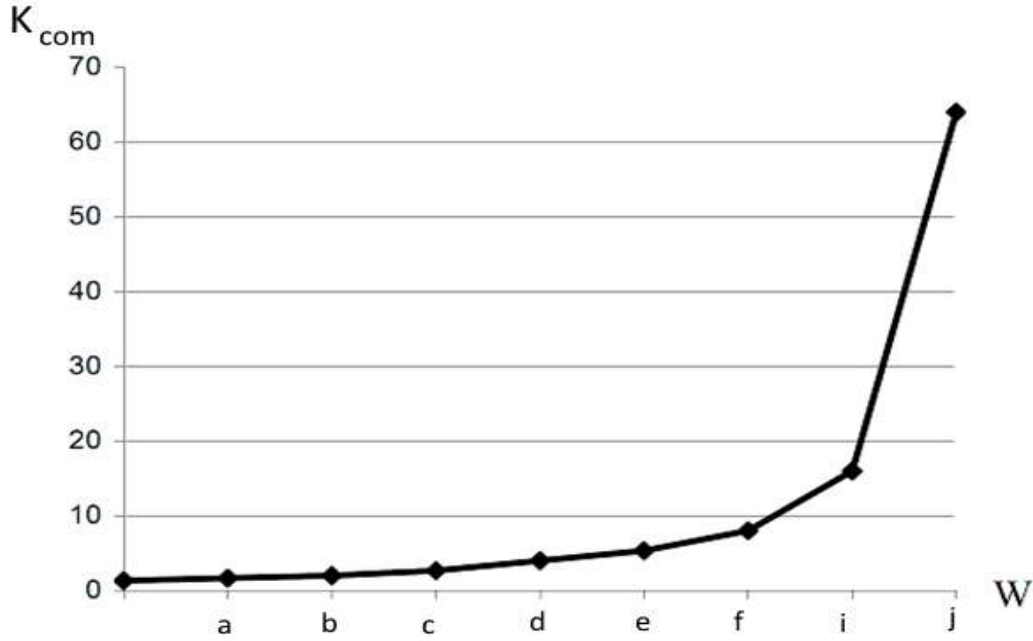


Figure 6: Dependence of the compression coefficients K_{com} on the number of excluded selection zones W , starting from the f zone

It can be seen that when “senior” zones are excluded from processing (at $W = 1$, the tenth selection zone, at $W = 2$, the tenth and ninth selection zones, etc.). the compression ratio increases slightly (from 1.4 to 8.8 for $W = 7$).

5. Results discussion

A significant increase in the compression ratio becomes possible with the “junior” zones (second and third) exclusion. It is not advisable to exclude the first ZQ, because at least 90% of the image energy is distributed in it, and the coefficients have a global sensitivity, therefore, its exclusion will introduce the greatest distortions.

From the analysis of the graph it follows that during the execution of the lossy image compression algorithm: to minimize distortions of the reconstructed image, it should be excluded from the processing of the selection zone with the lowest energy; to obtain significant compression ratios, it is necessary to “discard” ZQ with a large number of conversion ratios.

It can be seen from the figs. 5 - 6 that the most appropriate seems to be an exception from processing j , i , h and e zones, since they have the lowest average energy, and their “dropping” provides a compression ratio $K_{com} = 5$.

However, it should be noted that the root-mean-square deviation for high- and medium-resharpened images exceeds 1.5% already with the exclusion of j , i and f ZQ, which are responsible for fine details.

Wherein K_{com} does not exceed more than 4-5 times, and distortions occur throughout the image, which does not allow us to choose the level of detail that can be neglected in the given processing conditions.

Methods of zonal and zonal-boundary selection of orthogonal transformation coefficients have been obtained, which differ from the known ones in that: in the case of zonal selection, the features of the energy distribution over the processed image zonal sequences are taken into account and only those that introduce minimal distortions into the reconstructed image are excluded, which makes it possible to control the ratio of the compression ratio and the standard deviation; with zonal-boundary selection, the uneven redistribution of energy over zonal sequences of each processed image block is

taken into account, which makes it possible to reduce the entropy relative to the original image by an average of $4 \div 5$, increasing the efficiency of using statistical coding.

6. Conclusion

A method for assessing the degree of saturation of images has been developed, based on taking into account the global and local sensitivity of the basic functions of the Haar transformation using the energy index when analyzing the image transformant, which makes it possible to estimate the degree of saturation of the image with small details from the energy distribution for different images by zonal sequences. A total weighted energy index is proposed, which depends on the energy distribution and the number of orthogonal trans-formation coefficients, forming a zonal sequence.

The scientific novelty of the presented work lies in increasing the accuracy and efficiency of the data transformations carried out during assessing the degree of digital image compression by adapting the Haar transformation for the problem under consideration. The development and application of the method of adaptive compression of digital images will ensure the fulfillment of the requirement for the efficiency of processing and transmission of digital video information in real time at an acceptable level of distortion of the reconstructed image

Using the technique for assessing the degree of saturation of images will allow: more accurately, in comparison with expert judgment, to classify images; select the number of saturation classes required in each specific case; automate the classification process when evaluating an image; use the adaptive compression method as a procedure.

7. References

- [1] N. Rudnichenko, V. Vychuzhanin, V. Mateichyk, A. Polyvianchuk, Complex Technical System Condition Diagnostics and Prediction Computerization, in: Proceedings of The Third International Workshop on Computer Modeling and Intelligent Systems (CMIS-2020), CEUR-WS, 2608, 2020, pp.1-15.
- [2] B. Su, H. Zhang, H. Meng, Development and implementation of a robotic inspection system for power substations, Ind Robot, 44 (2017) 333-342.
- [3] N. Rudnichenko, M. Stepanchuk, D. Shibaev, Concept of physiognomic images recognition and analysis based on the artificial neural networks, in: Proceedings of the International Scientific and Technical Conference "Computer Graphics and Image Recognition", 2, 2018, pp. 65-68.
- [4] L. Guerriero, D. Martire, D. Calcaterra, M. Francioni, Digital Image Correlation of Google Earth Images for Earth's Surface Displacement Estimation, Remote Sens, 21 (2020) 3518. doi: 10.3390/rs12213518
- [5] C. Kuenzer, S. Dech, Remote Sensing and Digital Image Processing, Thermal Infrared Remote Sensing – Sensors, Methods, Applications, 4 (2013) 1-26. doi:10.1007/978-94-007-6639-6_1
- [6] P. Jones, M. Rabbani, Digital Image Compression (2020). doi:10.1201/9781003067054-8.
- [7] P. Chernov, Testing of Image Compression Methods, Vestnik MEI, 44 (2021) 105-113. doi: 10.24160/1993-6982-2021-4-105-113.
- [8] M. Zhou, Z. Liu, H. Hama, A New Digital Image Compression Method Using Chrestenson Transform, Genetic and Evolutionary Computing 387 (2016) 29-34. doi:10.1007/978-3-319-23204-1_4.
- [9] R. Priyanga, M. Kannamma, J. Sathiaselvan, Quality metrics digital image compression, in Proceedings of SPIE, 2016, pp.225-232.
- [10] K. Hashim, S. Baawi, B. Hilal, A New Proposed Method for A Statistical Rules-Based Digital Image Compression, Journal of Physics: Conference Series, 1897 (2021). doi:10.1088/1742-6596/1897/1/012067.
- [11] A. Singh, A Survey on Image Compression Methods, International Journal of Engineering and Computer Science 45 (2017). doi:10.18535/ijecs/v6i5.34.
- [12] M. Bespalov, Wavelet p-analogs of the discrete Haar transform, Izvestiya of Saratov University. New Series. Series: Mathematics. Mechanics. Informatics 21 (2021) 520-531. doi:10.18500/1816-9791-2021-21-4-520-531.

- [13] J. Uri, N.M.R. Radon Haar transforms (2017). doi:10.1117/12.2294846.
- [14] A. Belyaev, O. Yevtushok, N. Ryaboshchuk, Lossless Image Compression Algorithm Based on Haar Transform, in: 2021 IEEE Conference of Russian Young Researchers in Electrical and Electronic Engineering, 2021, pp.1960-1964. doi:10.1109/ElConRus51938.2021.9396588.
- [15] I. Dagher, M. Saliba, R. Farah, Combined DCT-Haar transforms for image compression, International Journal of Imaging Systems and Technology 28 (2018). doi:10.1002/ima.22286.
- [16] A. Dziech, New Orthogonal Transforms for Signal and Image Processing, Applied Science, 11 (2021), 7433. doi:10.3390/app11167433
- [17] R. Gonzalez, R. Woods, Digital Image Processing (2018).
- [18] A. Thompson, The Cascading Haar Wavelet Algorithm for Computing the Walsh–Hadamard Transform, IEEE Signal Process, 24 (2017) 1020–1023.
- [19] A. Andrushia, R. Thangarjan, Saliency-Based Image Compression Using Walsh–Hadamard Transform (WHT), Biologically Rationalized Computing Techniques for Image Processing Applications, 57 (2018) 678-690.
- [20] A. Ashrafi, Walsh–Hadamard Transforms: A Review; Advances in Imaging and Electron (2017).
- [21] K. Marlapalli, R. Bandlamudi, R. Busi, V. Pranav, B. Madhavrao, A Review on Image Compression Techniques. Communication Software and Networks, in Proceedings of INDIA (2019) 271-279. doi:10.1007/978-981-15-5397-4_29.
- [22] C. Yuan-Yuan, Z. He-Xin, S. Hong-Yan, Discrete cosine transform optimization in image compression based on genetic algorithm, in: 8th International Congress on Image and Signal Processing (CISP), 2015, pp. 199-203. doi: 10.1109/CISP.2015.7407875.
- [23] D. Vasin, Description Models, Methods, Algorithms, and Technology for Processing Poorly tructured Raster Graphic Documents, in: Proceedings of the 30th International Conference on Computer Graphics and Machine Vision (GraphiCon 2020) (2020).

Towards the Tokenization of Business Process Models using the Blockchain Technology and Smart Contracts

Andrii Kopp and Dmytro Orlovskyi

National Technical University “Kharkiv Polytechnic Institute”, Kyrpychova str. 2, Kharkiv, 61001, Ukraine

Abstract

Business process modeling helps organizations to capture their workflows visually as diagrams that could be then used to share best practices, identify inefficiencies in ongoing activities, instruct employees, use them as reference solutions, etc. Design and analysis of business process models are essential technologies of the Business Process Management approach, which is successfully adopted and practiced nowadays by many large and medium enterprises. Therefore, business process models should be considered as organizational assets that depict usable and competitive business solutions, which value could be proven by benchmarking. Single reference business process models or even collections of business process models are already accessible on Internet, sometimes for free or, usually, for purchasing because of the value of transferred knowledge. However, peer-to-peer exchange of business process models on a commercial basis is still far from a unification. At the same time, the tremendous growth of the cryptocurrency market and the adoption of Bitcoin by governments (first by El Salvador in June 2021) makes crypto-economics, also referred to as “tokenomics”, usable for organizational knowledge sharing and exchanging without third party authorities, such as banks, or payment systems. Moreover, collaborating parties could reach a consensus when exchanging business process models using smart contracts and crypto-tokens to access shared knowledge artifacts. Therefore, this paper proposes an approach to business process model tokenization using blockchain technology and smart contracts. There was proposed as an Ethereum smart contract that combines features of the business process model collection and the non-fungible token. A prototype of a decentralized application was developed and its usage was demonstrated.

Keywords

Business Process Modeling, Blockchain, Smart Contract, Tokenization

1. Introduction: Related Work and Problem Statement

Nowadays digital transformation is a trend in enterprise management. In the first place, digital transformation is associated with Business Process Management (BPM) and its applications in business process modeling, and automation using BPM suites. However, business process modeling is used not only to draw executable workflows – they are usually simplified enough and describe mostly routine document flows. The main goal of business process modeling includes a visual representation of business activities as graphical diagrams to identify and understand ongoing workflows, find bottlenecks for improvement, and ensure communication between IT (Information Technology) staff and business stakeholders. Widely used reference models of typical enterprise business processes are used by organizations to adopt and tune concerning industry and internal needs. At the highest levels of BPM, maturity organizations have in their possession large collections of business process models that are extremely valuable for them and their competitors. Therefore, successful and time-proven business scenarios captured as process models could be sold to other organizations willing to achieve

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, May 12, 2022, Zaporizhzhia, Ukraine

EMAIL: kopp93@gmail.com (A. Kopp); orlovskyi.dm@gmail.com (D. Orlovskyi)

ORCID: 0000-0002-3189-5623 (A. Kopp); 0000-0002-8261-2988 (D. Orlovskyi)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

positive BPM outcomes. Secure and peer-to-peer exchange of business process models could be organized using blockchain technologies, including cryptocurrencies and smart contracts. The only problem of business process models tokenization, i.e. representation of them as digital tradable assets on a certain crypto-platform [1], must be solved to reach such BPM-driven tokenomics. The object of this research includes sharing and exchange of business process models. The subject of research is the approach to business process model tokenization using blockchain and smart contracts.

This paper is structured as follows: in the introductory section is given an overview of related work (sub-section 1.1) and the problem statement is made (sub-section 1.2); section 2 outlines the approach to business process models tokenization based on smart contracts usage; section 3 contains results and their discussion regarding decentralized application prototype development and validation; conclusion and further research plans are outlined in section 4.

1.1. Related Work

1.1.1. Business Process Modeling

In general business processes are considered as structured collections of manual or automatic (by IT systems, e.g. BPMS, CRM, ERP, and others) executed activities necessary to achieve business goals and satisfy end customers [2]. According to the BPM approach, which focuses on the automation of business processes and support of human interaction with IT applications, there are different phases of business process analysis, modeling, implementation, and deployment to the execution environment of a certain IT system, monitoring, and evaluation [2]. Despite business process models could be defined in two ways: textual or visual [2], graphical diagrams are much more informative in describing how to process activities are triggered by events, which data objects are processed, which organizational units are responsible for process execution, and which outputs are produced [3]. Business process modeling is considered the approach to the depiction of current or future organization activities driven by events and control flow logic [4]. Also in [4], business process models are named as the key tools for process-aware information systems design, business process re-engineering, and service-oriented architectures design [4]. Business process models are also considered graphical knowledge resources and could be used as guidelines to introduce best practices for BPM adoption across multiple enterprises, as it is done by industry reference models that share knowledge of other organizations [4].

An extensive classification of business process modeling perspectives was given by J. Krogstie in [5]. There are object, communication, role, topological, functional, and behavioral perspectives. Some of the most well-known and widely used modeling standards, notations, and languages were covered in this classification: UML (Unified Modeling Language), IDEF0 (Integrated Definition for Functional Modeling), DFD (Data Flow Diagrams), EPC (Event-driven Process Chains), and BPMN (Business Process Modeling and Notation) [5]. A combined behavioral and functional approach was chosen as the most suitable for business process modeling, whereas the BPMN process diagramming notation has been adopted as the standard in BPM in general and, in particular, in business process modeling [5].

Even though EPC notation is still in use in the area of business process modeling, many users have replaced it with BPMN since it is a standard [6], moreover, most modern EPC modeling software tools support BPMN as the second or even alternative modeling notation (e.g. ARIS Express) [6]. As for other languages and standards, such as UML, DFD, and IDEF0, they did not become popular and are rarely used for business-oriented process modeling in practice [6]. In the last decade, BPMN has become a leading business process modeling notation [6], there are over 70 BPMN modeling tools listed on the “BPMN Tool Matrix” [7] and over 50 open-source tools related to BPMN listed on “Source Forge” [8]. BPMN is a complicated notation, there are four core symbols in use [9]:

- Activities. Tasks or units of work that have a certain duration.
- Events. Things that happen instantaneously and indicate when process instances begin (start events), complete (end events), or when something happens inside a process (intermediate events).
- Gateways. Model parallel (AND), exclusive (XOR), and inclusive (OR) branches of a process using combinations of splits and joins.
- Sequence flows. Model logical relations between elements, when one is followed by another.

Pools and lanes are used to model process resources: pools define business parties, i.e. boundaries of a business process, while lanes define roles, i.e. organizational units or persons, that take part in a business process execution [9]. There are also data objects and data stores (containers of data objects that could be databases for electronic objects or some places for physical objects) that could be depicted in BPMN models to represent information or material flows between activities [9]. As for limitations of the BPMN notation, authors of [9] do not recommend using OR-gateways without strong necessity because of their confusing logic, as well as to use data objects and data stores that make diagrams less readable and only useful when communicating to the IT development team for process automation.

Sample BPMN model of a goods purchase business process, which outlines considered symbols, is demonstrated in Fig. 1 below.

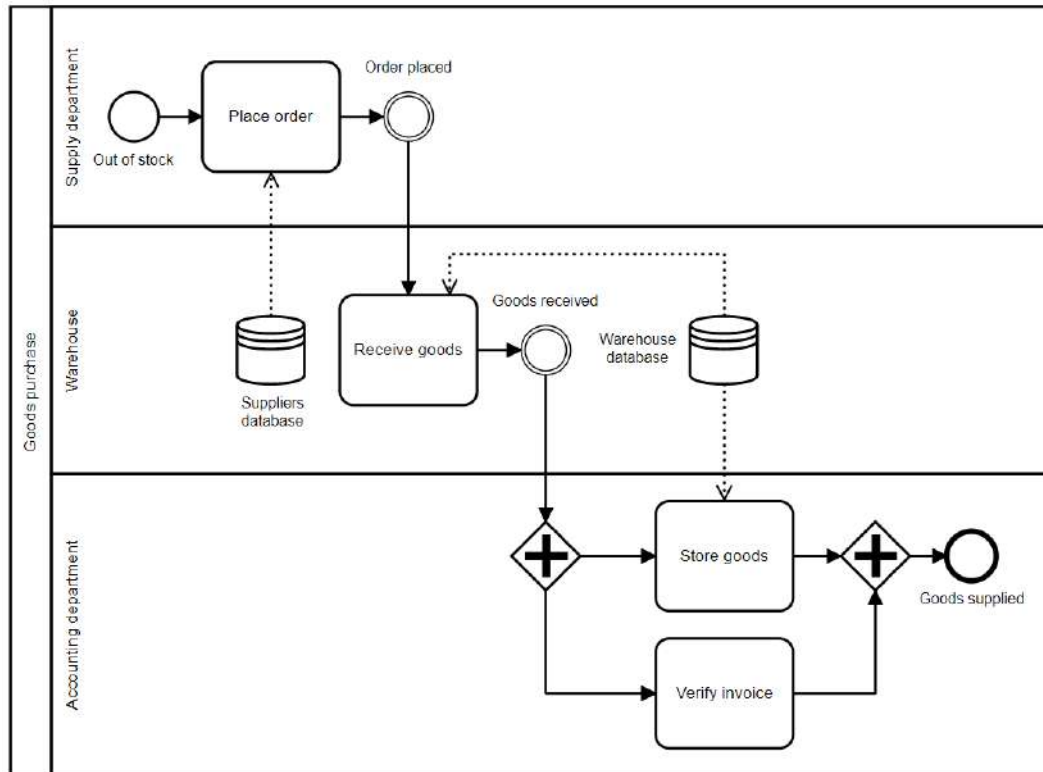


Figure 1: The BPMN model of a goods purchase business process

This diagram contains one pool "Goods purchase" that defined business process boundaries, three lanes "Supply department", "Warehouse", and "Accounting department" that define responsible roles for the process execution. As for core BPMN symbols, there are: start event "Out of stock" that starts process instance, intermediate events (e.g. "Goods received" and "Order placed"), end event "Goods supplied" that completes process instances, tasks (e.g. "Place order", "Receive goods", "Store goods", and "Verify invoice"), split and join AND-gateways to implement parallel execution of "Store goods" and "Verify invoice" tasks (see Fig. 1).

1.1.2. Blockchain Technology

The term "blockchain" means an immutable or read-only data structure – a linked list of blocks, a directed acyclic graph (DAG), or a tree-like data structure, in which new data can be only appended at the end of the blockchain [10]. Blockchain systems are considered trustless distributed networks where a ledger is replicated over several nodes, each of which can participate in the decision-making process (i.e. a consensus protocol used to achieve an agreement between nodes when adding new blocks to the blockchain) in a decentralized and democratic manner [10]. Moreover, blocks contain transactions with certain data about transferred cryptocurrency, assets, tokens, etc. [10].

There are the following principles of blockchain technology that make it adaptable in many industries where security, transparency, and consistency are necessary features: decentralization, peer-to-peer communication, transparency, pseudonymity, irreversibility, and computational logic [11]. In general, blockchain is useful only when several entities are collaborating and sharing data since local copies are maintained on participating nodes to ensure consistency and tamper resistance [11].

The underlying idea of blockchain is using cryptography to link data blocks of referred distributed and decentralized chains. The first block is called the “genesis block”, it does not refer to any previous block in the chain [12]. However, each block contains data (usually as strings), nonce (the unique number usually related to mining – none generated repeatedly until meets consensus requirements of new block adding), and the hash value of all fields of the previous block (i.e. the cryptographic link to a preceding block, also referred as “previous hash”), the hash value of all fields of the current block [12]. The original consensus algorithm Proof-of-Work (PoW) was used to confirm transactions and add new blocks to the Bitcoin blockchain. This algorithm assumes miners (persons who share their computational power to support the network) are competing to confirm transactions and get rewarded in cryptocurrency [11]. Today PoW is used by such most popular cryptocurrencies like Bitcoin and Ethereum. However, Ethereum is planned to be switched to another popular consensus algorithm referred to as Proof-of-Stake (PoS). This algorithm overcomes high energy use since transactions are confirmed by validators who stake coins instead of sharing their computing resources. The higher stake is, the higher the probability to add a new block with pending transactions and getting rewarded [11].

However, in comparison to Bitcoin and other cryptocurrency-oriented blockchain platforms (such as Dogecoin, Bitcoin Cash, Litecoin, etc.), Ethereum was designed as a “world computer” with the ability to host smart contracts – programs to be executed within a blockchain platform [13].

1.1.3. Smart Contracts

Bitcoin as the first cryptocurrency was proposed in 2008 by Satoshi Nakamoto, who is presumably an anonymous person or group of persons hiding by this pseudonym. The difference from traditional payment systems is that electronic currency could be transferred in a peer-to-peer manner without a central party that needs to check the records of ownership. Blockchain technology helps to prevent the double-spending problem and enforce the validity of transaction records [13]. One more crypto-asset “token” is almost the same as the cryptocurrency – a bearer instrument used to transfer value between parties over the blockchain network; whereas tokens are created by a single party to represent a certain value, cryptocurrency is generated by the network as the reward for miners or validators [13].

Token technology was introduced and standardized by Ethereum and its smart contracts, which code describes how each token should work [13]. However, Ethereum was planned as a globally distributed computing network that uses publically stored immutable programs – smart contracts also referred to as decentralized applications [14]. However, decentralized applications usually include a client-side created using markup, style sheets, and JavaScript (JS). Decentralized applications (or DApps) use the Web3 JS library to interact with the smart contracts deployed on the blockchain [12]. DApps are immutable (as any data stored on the blockchain) and perform exactly as they were developed, without the possibility of fraud, downtime, censorship, or interference [14].

1.1.4. Tokenization

In the context of blockchain, tokenization means a representation of real physical or electronic assets digitally on the blockchain [15]. There could be commodities, real estate, ownership rights for arts or other collectibles, currency, or any other kinds of assets. As advantages tokenization offers faster and cheaper transaction processing, flexibility, decentralization, security, and transparency [15], but there are disadvantages, such as regulatory and legality issues, as well as technical barriers caused by the use of DApps [15]. There are two types of crypto-tokens [16]:

- Fungible Tokens (FT), which value is identical among all tokens, and which are exchangeable to each other (i.e. one FT token could be replaced by another FT token similarly to digital cash).

- Non-fungible Tokens (NFT) are unique and not equal to each other in value (i.e. each NFT token is different from others and cannot be replaced by any of them).

Another explanation of FT and NFT given in [16] says that dollar bills are exchangeable to other dollar bills, so they are fungible, whereas baseball cards or other collectibles are unequal in their value and cannot be replaced with another one.

From the technological point of view, both FT and NFT are implemented as smart contracts, which are programmed in a specific way. Several smart contract standards for tokens were developed by the Ethereum community [17]. The most popular contract standards are [17]:

- ERC20 standard for fungible tokens, which is suitable for multiple use cases, such as payment tokens, loyalty coins, gift cards, etc. A great example of NT implemented as the ERC20 token on the Ethereum blockchain is Tether USD, a stable cryptocurrency (also referred to as the “stablecoin”) that digitally represents USD (United States Dollar) [18].
- ERC721 standard for non-fungible tokens, which is suitable for documents, land titles, digital identities, real estate, collectibles, etc. A great example of NFT implemented as the ERC721 token on the Ethereum blockchain is CryptoKitties, an Ethereum-based game in which players buy, sell, and breed collectible digital cats [19].

Nowadays OpenZeppelin organization provides extensive references for the development of ERC20, ERC721, and other smart contract standards [20].

1.2. Problem Statement

The problem of blockchain-driven cross-organizational business process modeling, versioning, and executing was considered by Härer in [21] and Fill in [22], Hull in [23], Viriyasitavat and Hoonsopon in [24], De Sousa and Corentin in [25], Milani and Garcia-Banuelos in [26], and others. Thanks to the considered blockchain advantages, inter-organizational storage of tokenized business process models provides collaborative parties with proof of authorship, censorship resistance, timestamping, and immutability. Smart contracts and token standards provide decentralized financial (DeFi) capabilities, such as peer-to-peer exchange of enterprise knowledge presented as business process models without the need of any third-party authorities, i.e. banks or payment systems. Therefore, the problem of business process model tokenization remains relevant and respective information technologies should consider the latest trends in blockchain technology and digital economics. Since in Spring 2021 NFTs got a focus in global crypto attention, mostly for collectibles and digital art (the entire market exceeds 130 USD by Spring 2021) [27], it seems like a great opportunity to extend the use cases of NFTs with tokenization of enterprise models, in particular – business process models.

The NFT better suites process models that are unique and unequal in terms of their syntactic and semantic properties, which could be used to define the value of shared models. There could be used SEQUAL (Semiotic Quality) framework [5] for the evaluation of syntactic and semantic validity and completeness of BPMN business process models given as graphic diagrams (i.e. images). Hence, the following business process model tokenization workflow could be used (see Fig. 2).

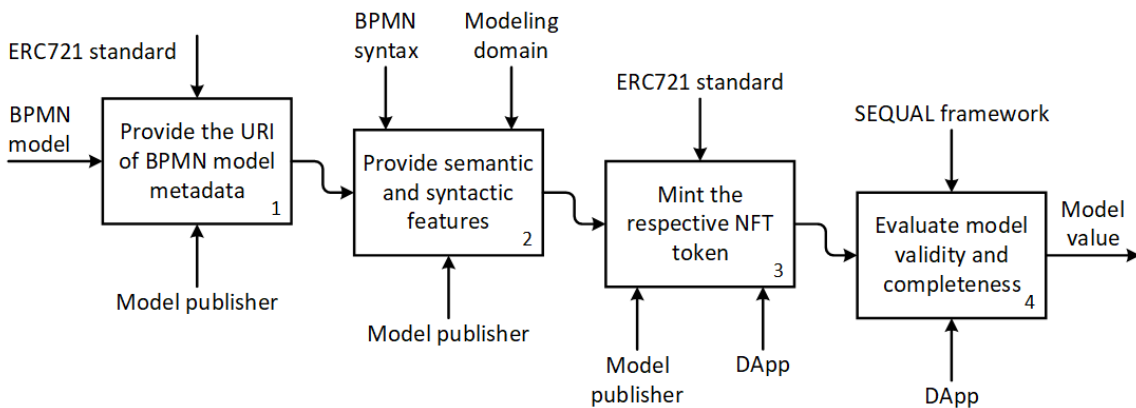


Figure 2: The business process model tokenization workflow

As for the blockchain platform, there could be chosen Ethereum as the pioneering and still leading smart contracts platform despite its competitors, such as Binance Smart Chain, Polkadot, Solana, and others [28]. Thus, the ERC721 standard may be used for business process model tokenization. As it is shown in Fig. 2, the URI (Uniform Resource Identifier) of a business process model should be used as the token URI. Besides the token URI, the ERC721 standard should be extended with syntactic and semantic features of business process models. After NFT is minted (i.e. published on the blockchain), syntactic and semantic features are used to evaluate the validity and completeness of a process model to define its value for collaborating parties (Fig. 2).

2. Tokenization of Business Process Models using Smart Contracts

2.1. Evaluation of Tokenized Business Process Models

According to [5], an extensive contribution to the domain of quality of business process models by J. Krogstie, syntactic and semantic qualities of business process models could be defined as following (according to the proposed specialization of SEQUAL framework for business process models):

- Syntactic quality is formulated as the correspondence of all statements (i.e. symbols, such as activities, etc.) in a business process model to the vocabulary and syntax of modeling notation – in our case BPMN.
- Semantic quality, in its turn, is defined as the correspondence between all statements and the modeling domain (i.e. a real business process).

In terms of formalisms proposed in [5], there could be presented phenomena of syntactic invalidity or syntactic incompleteness:

- Syntactic invalidity – when symbols are not part of BPMN notation (i.e. elements differ from events, activities, gateways, etc.).
- Syntactic incompleteness – when symbols do not obey BPMN syntax (i.e. elements connected improperly).

Formally, syntactic quality is defined as following [5]:

$$\text{Syntactic quality} = 1 - \frac{\#M \setminus L + M_{\text{missing}}}{\#M}, \quad (1)$$

where:

- $\#M \setminus L$ is the number of business process model M statements that do not correspond to the BPMN language L ;
- M_{missing} is the number of missing statements that make a model syntactically incomplete (i.e. missing sequence flows or events);
- $\#M$ is the total number of business process model statements.

As for semantic properties, there are also validity and completeness phenomena [5]:

- Semantic invalidity – when statements are not part of the modeling domain.
- Semantic incompleteness – when statements are not correct or relevant to the modeling domain.

In [5] there are two metrics of semantic quality:

$$\text{Semantic validity} = 1 - \frac{\#M \setminus D}{\#M}, \quad (2)$$

$$\text{Semantic completeness} = 1 - \frac{\#D \setminus M}{\#D}, \quad (3)$$

where:

- $\#M \setminus D$ is the number of business process model M statements that do not belong to the modeling domain D ;
- $\#D \setminus M$ is the number of model M statements that are not correct or relevant to the modeling domain D ;
- $\#D$ is the number of statements in the modeling domain.

Therefore, a quality-oriented smart contract that may be used to assess tokenized business process models should consider formalisms of SEQUAL framework (1), (2), and (3) regarding syntactic and

semantic quality of business process models given in [5]. Thus, each tokenized BPMN model should be described by the following tuples:

$$\langle N_{invalid}^{synt}, N_{incomplete}^{synt} \rangle, \quad (4)$$

$$\langle N_{invalid}^{sem}, N_{incomplete}^{sem} \rangle, \quad (5)$$

where:

- $N_{invalid}^{synt}$ is the number of syntactically invalid statements in a business process model;
- $N_{incomplete}^{synt}$ is the number of syntactically incomplete statements in a business process model;
- $N_{invalid}^{sem}$ is the number of semantically invalid statements in a business process model;
- $N_{incomplete}^{sem}$ is the number of semantically incomplete statements in a business process model.

In addition to (4) and (5), there should be noted the total number of statements N , which could be considered as equal to $\#M$ and $\#D$, since the syntax and semantics of real business process models are completing each other in order to reflect business activities.

Finally, there are following mapping should be noted:

$$Syntactic : tokenId \rightarrow \langle N_{invalid}^{synt}, N_{incomplete}^{synt} \rangle, (N_{invalid}^{synt} + N_{incomplete}^{synt}) \leq N, \quad (6)$$

$$Semantic : tokenId \rightarrow \langle N_{invalid}^{sem}, N_{incomplete}^{sem} \rangle, (N_{invalid}^{sem} + N_{incomplete}^{sem}) \leq N, \quad (7)$$

$$Total : tokenId \rightarrow N, \quad (8)$$

where $tokenId$ is the unique identifier of each tokenized business process model.

Whereas (6), (7), and (8) define the structure of business process model properties, there are following metrics should be calculated then:

$$\text{Syntactic validity} = 1 - \frac{N_{invalid}^{synt}}{N}, \quad (9)$$

$$\text{Syntactic completeness} = 1 - \frac{N_{incomplete}^{synt}}{N}, \quad (10)$$

$$\text{Semantic validity} = 1 - \frac{N_{invalid}^{sem}}{N}, \quad (11)$$

$$\text{Semantic completeness} = 1 - \frac{N_{incomplete}^{sem}}{N}. \quad (12)$$

Interpretation of calculated metrics (9) – (12) may be done using the Harrington scale [29] to transform crisp values of syntactic and semantic validity and completeness into linguistic values. To visualize obtained linguistic values, it is proposed to use “Green”, “Yellow”, and “Red” color codes for “Good”, “Satisfied”, and “Bad” values respectively (see Table 1).

Table 1

Translation of syntactic and semantic quality metrics into linguistic values and color codes

Linguistic value	Threshold	Color code
Good	Syntactic validity ≥ 0.8 , Syntactic completeness ≥ 0.8 , Semantic validity ≥ 0.8 , Semantic completeness ≥ 0.8	Green
Satisfied	Syntactic validity ≥ 0.63 , Syntactic completeness ≥ 0.63 , Semantic validity ≥ 0.63 , Semantic completeness ≥ 0.63	Yellow
Bad	Syntactic validity < 0.63 , Syntactic completeness < 0.63 , Semantic validity < 0.63 , Semantic completeness < 0.63	Red

Obtained evaluation results will help collaborating parties define values and formulate prices of shared business process models for further exchange.

2.2. NFT-Compatible Smart Contract to Store Business Process Models

As it was outlined before, the ERC721 NFT standard [30] will be used for business process model tokenization. According to the problem statement, the ERC721 standard should be extended with the following behavior: model publishing and contacts data publishing (so parties can reach each other for collaboration). Before a model is published to the blockchain, i.e. the respective NFT is minted, special metadata JSON (JavaScript Object Notation) file should be published to the Internet and be accessible by a certain URI (later will be used as the token URI). Such JSON document accessible by the token URI might include the following properties [30]:

- Name (i.e. a brief description of a depicted business process).
- Description (i.e. a more detailed description of a depicted business process).
- Image URI (i.e. an image of a BPMN diagram).

Obviously token metadata stored in JSON files will not be stored on the blockchain. However, such an approach is considered the most efficient from the perspective of transactions' speed and cost [31]. Whereas to save tamper resistance provided by the blockchain technology, only the hash value of an entire document could be found using the secure algorithm (e.g. SHA-256 or others) and stored on the blockchain to verify the identity of the original document [31].

Then, to summarize, we may define two more mappings in addition to formalisms (6), (7), and (8) denoted earlier:

$$Party : address \rightarrow contacts, contacts \subseteq \{\Sigma_{UTF8}, \emptyset\}, \quad (13)$$

$$Hash : tokenId \rightarrow sha256(tokenURI), \quad (14)$$

where:

- *address* is the user address in the Ethereum network;
- *contacts* is the user's contact information outside the blockchain (e.g. email address, phone or messenger number, etc.) represented as the string value of UTF-8 characters set Σ_{UTF8} or even the empty characters set $\emptyset$ (i.e. the empty string) if a party decided to keep pseudonymity;
- *sha256* is the hashing algorithm SHA-256 applied to the JSON document *tokenURI*.

Moreover, besides the (6) – (8) and (13) – (14) mappings, the smart contract should store a number of models *modelsCount* used both as the capacity and the identifier of next minted token. In general, the minting or model publishing sequence diagram may look as following (see Fig.3).

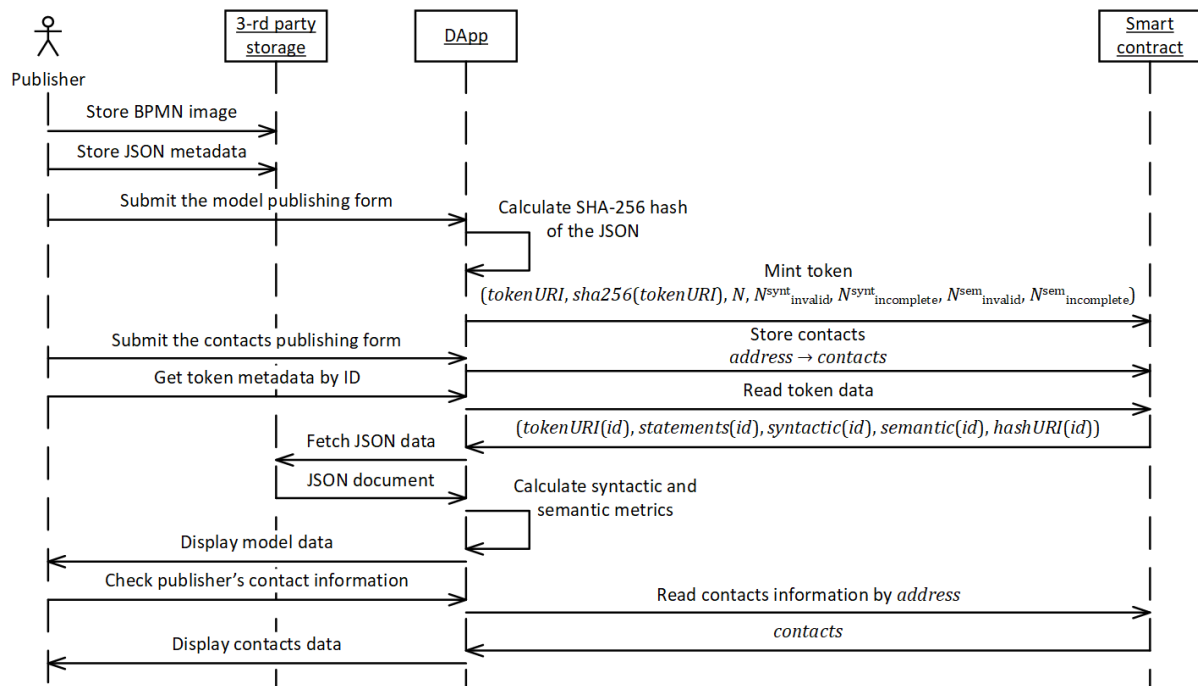


Figure 3: The model publishing sequence diagram

Thus, as depicted in Fig. 3 behavior should be implemented in the form of an ERC721-compatible smart contract that allows for the publication and read tokenized business process models and related data. In terms of UML class diagrams, the smart contract can be depicted as follows (see Fig. 4).

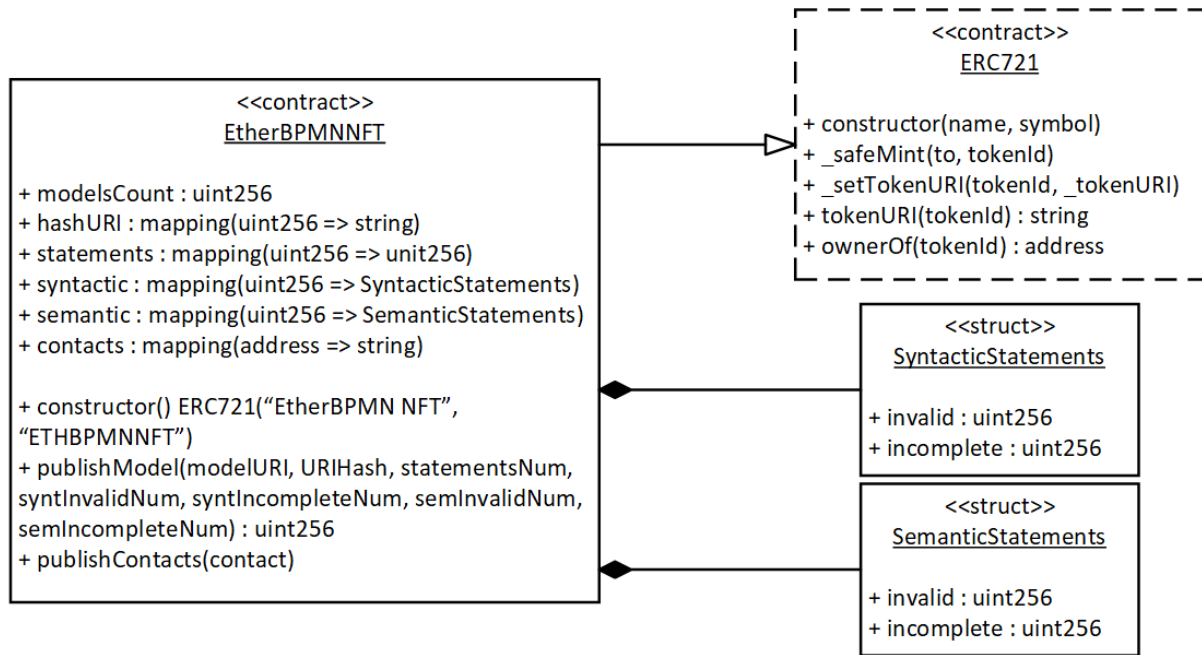


Figure 4: The class diagram of the smart contract

Contract “ERC721” with the dashed border line represents OpenZeppelin ERC721 implementation provided at [30] (see Fig. 4), which is used as the generic for developed contract “EtherBPMNNTF”. The smart contract also contains two structures “SyntacticStatements” and “SemanticStatements” that serve as tuples of semantic and syntactic characteristics of tokenized BPMN models. As shown in Fig. 4 mappings of “EtherBPMNNTF” smart contract implement formalisms (6) – (8) and (13) – (14), while functions “publishModel” and “publishContacts” are used to mint NFT tokens and store collaborator contact information respectively. Some of ERC721 standard functions can be called from the developed contract, these functions are also mentioned in Fig. 4.

3. Results and Discussion

3.1. Design and Development of a Decentralized Application Prototype

Considered smart contract is barely useless without having a special decentralized application that allows ordinary users to interact with the blockchain. The following use cases should be supported by the DApp prototype for BPMN model tokenization (see Fig. 5).

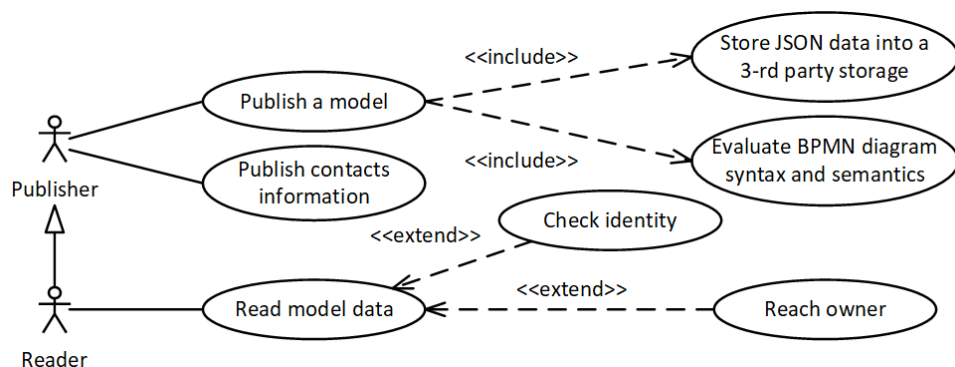


Figure 5: The use cases of DApp prototype

The system architecture of the decentralized application is similar to any client-server web application, whereas instead of an application server and persistent storage the smart contract and blockchain ledger is used respectively [12]. The structure of the DApp prototype is shown in Fig. 6.

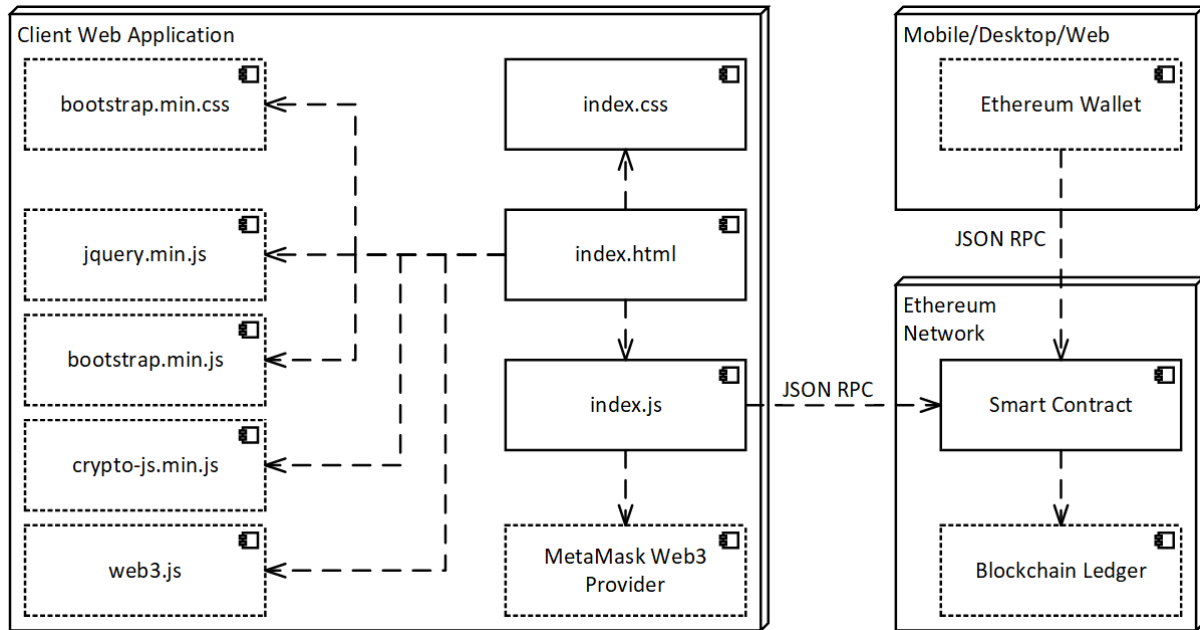


Figure 6: The system architecture of DApp prototype

Components with dashed border lines are third-party libraries, such as Bootstrap for user interface, CryptoJS for SHA-256 calculation, Web3 for interaction with the Ethereum platform, Ethereum blockchain ledger itself, MetaMask Web3 provider (as the Google Chrome extension), and also some Ethereum mobile, desktop, or web wallet that can be used to hold business process models in the form of ERC721-compatible NFTs, which are tradable and exchangeable as any other digital crypto-assets.

3.2. Validation of a Decentralized Application Prototype

Developed prototype of a decentralized application for business process model tokenization that can be used only with MetaMask or another Web3 provider. In the authors' opinion, MetaMask is the easiest and the most popular tool to work with Ethereum networks, either the main network or test network, since it only requires the installation of the Google Chrome extension. According to the sequence diagram (Fig. 3), the tokenization procedure starts with the storing of the JSON metadata and BPMN diagram image in the third-party non-blockchain storage. Let us imagine we are intending to store BPMN shown in Fig. 1. It is image is already given and could be stored somewhere on the Internet with the permanent URI. As for other metadata properties, the following could be given:

- Name: "Goods purchase".
- Description: "The BPMN model of a goods purchase business process".

Therefore, respective JSON (as well as the BPMN diagram image) could be stored in the GitHub repository of this project and may look like the following (see Fig. 7).

```
{
  "name": "Goods purchase",
  "description": "The BPMN model of a goods purchase business process",
  "image": "https://raw.githubusercontent.com/andriikopp/blockchain-repository/main/nft/models/0.png"
}
```

Figure 7: The JSON document representing the NFT metadata

The URI of future NFT is following – "https://raw.githubusercontent.com/andriikopp/blockchain-repository/main/nft/models/0.json". Now, after the non-blockchain data is stored, let us fill and submit the model publishing form in the DApp (see Fig. 8).

Publish a model

Model URI
 https://raw.githubusercontent.com/freebpmnquality/cloud-services/main/storage/nft/models/O.json

URI hash (SHA-256 is calculated automatically, use [this tool](#) to check or re-calculate)
 2ba63e1325c4bcd59a171da0ff3d78bbe381963d0004f3cd16097e8d13798eee

Number of statements
 4

Number of syntactically invalid statements
 0
 Statements are not part of the BPMN notation.

Number of syntactically incomplete statements
 0
 Statements do not obey the BPMN syntax.

Number of semantically invalid statements
 1
 Statements are not part of the modeling domain.

Number of semantically incomplete statements
 2
 Statements are not correct or relevant about the modeling domain

Close Publish

EtherBPMN NFT by freebpmnquality, 2021 CC BY-ND

Figure 8: The model publishing form of the DApp

When the “Publish” button is pressed, the MetaMask appears with the request to sign a transaction as it is shown in Fig. 9. In case of the transaction is successfully mined, a respective token will be minted and can be accessed by the address “0xabc33640b17def441cfb455efa1c3f8f490f4616” and ID “0” (since it is the first NFT in a collection). A respective token could be added to MetaMask for future exchange with other parties (see Fig. 9).

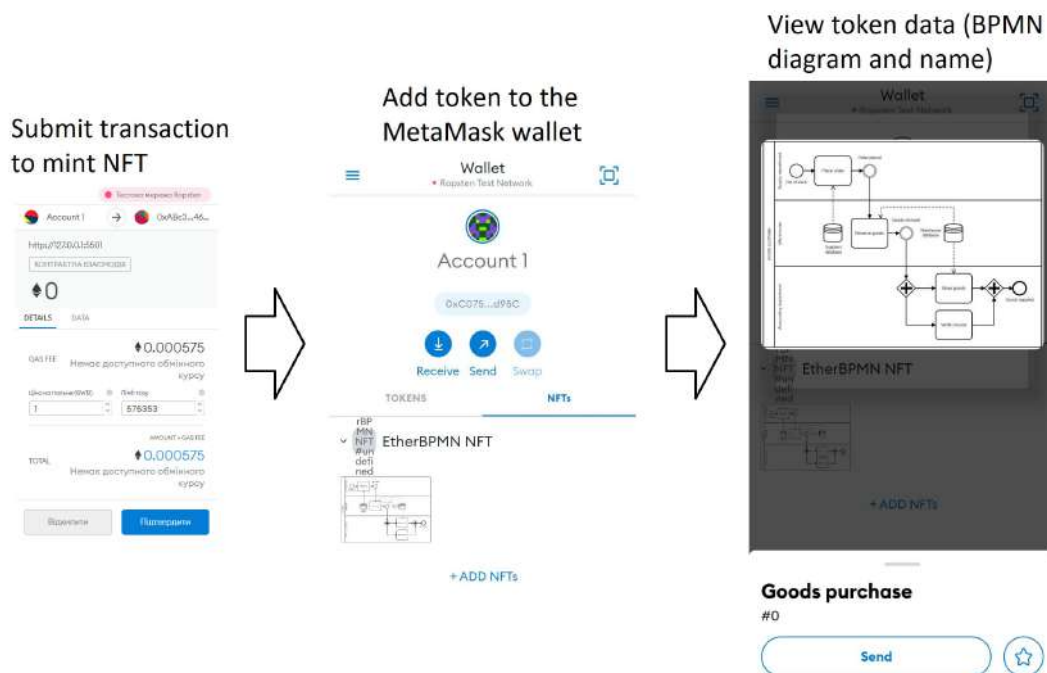


Figure 9: Tokenized business process model in the MetaMask wallet

Now BPMN diagram shown in Fig. 1 can be kept in the MetaMask wallet or any other Ethereum wallet that supports ERC721 tokens, exchanged with other Ethereum users, or even traded using NFT exchanges and marketplaces like other crypto-tokens that represent collectibles or digital art.

The developed DApp prototype also displays tokenized business process models with all given JSON-based data and owner's address, but also with the specific characteristics of BPMN diagrams: quality metrics of syntactic and semantic validity and completeness, a hash value used to ensure the identity of business process model metadata, and owner's contact information (if it was preliminary published to the smart contract, of course) as it is shown in Fig. 10.

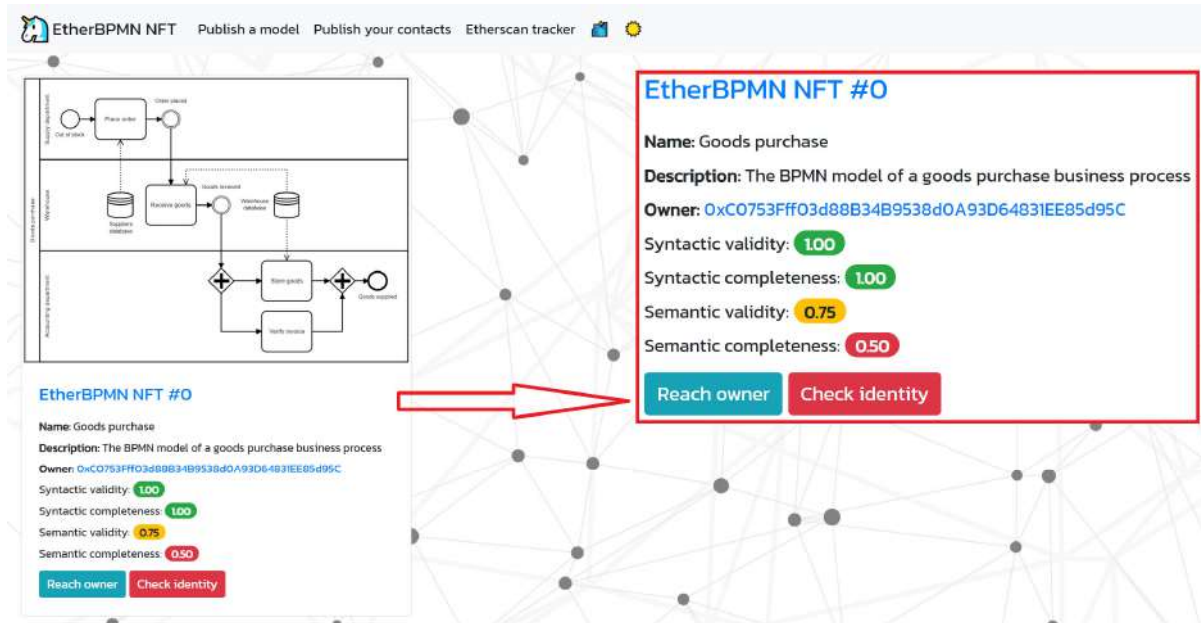


Figure 10: The homepage of developed DApp

According to Fig. 10, each tokenized BPMN model are displayed syntactic and semantic metrics that correspond to (9) – (12) formulas and color codes given in Table 1. Besides metrics, it is possible to reach the owner by retrieving its contact information (in case it was provided) and checking identity by comparing the SHA-256 hash value (stored on the blockchain in the smart contract mapping) to the actual SHA-256 hash value of JSON requested by token URI.

The form of contacts information publishing, as well as examples of requested contact information and SHA-256 hash of the NFT data, are shown in Fig. 11.

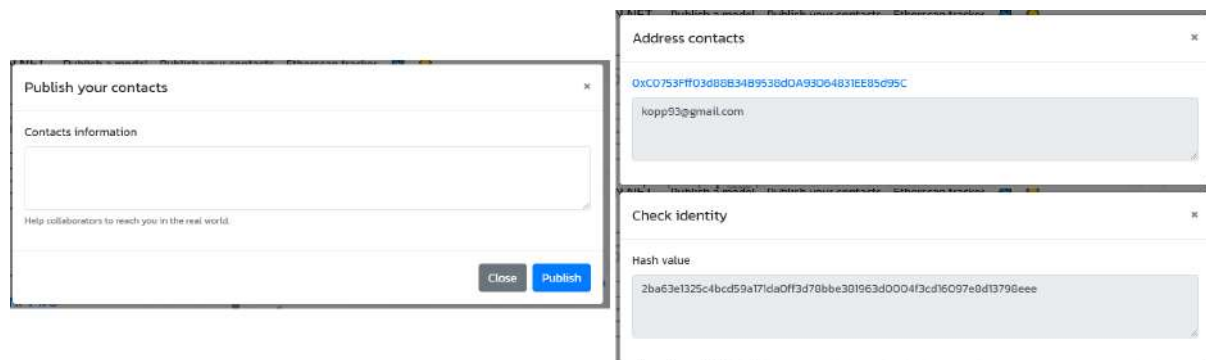


Figure 11: The contacts publishing form, contacts request window, and identity check window

Provided quality metrics may help to define the value of the business process model when exchanged, while the hash value may help to check whether the BPMN has been tampered with by the previous owner. Currently, the DApp is at the software prototype stage being under construction and testing [32]. The MetaMask or another wallet with an in-built Web3 provider should be installed to use the DApp, as well as the Ropsten testnet account should be present for basic usage.

4. Conclusion and Future Work

In this research paper, we proposed the approach to business process model tokenization based on blockchain technology and smart contracts. Based on the performed state-of-the-art overview, the BPMN business process modeling notation was chosen to describe tokenized business process models as the most widely used and considered standard in the BPM industry. Essentials of blockchain technology were considered to prove the relevance of business process model tokenization, and the essentials of smart contracts and decentralized applications were overviewed as well. Two standards of tokens – ERC20 and ERC721, which represent fungible and non-fungible tokens respectively were considered to select the appropriate token standard for BPMN diagrams. Based on features of NFTs, the ERC721 standard has been chosen, as well as Ethereum as the pioneering and still leading smart contracting platform has been chosen for implementation.

To provide tokenized business process models with specific features useful to define their value for exchanging and trading, essential syntactic and semantic quality metrics were used: validity and completeness [5]. Hence, originally provided by OpenZeppelin ERC721 smart contract for NFTs has been extended to keep syntactic and semantic properties of business process models, SHA-256 hash values of token metadata documents to ensure identity, and owner contact information to ensure collaboration of parties.

Developed DApp prototype allows to publish a model as the NFT, review already published NFTs, including model names, descriptions, BPMN diagram images (preview and full size), quality metrics, owner addresses, request owner contact information (if provided), check the identity of tokenized BPMN models, and share own contact information if necessary. The smart contract has been deployed to the Ropsten Ethereum test network, while the DApp is under development and testing.

Future work in this area includes the development of the decentralized marketplace and exchange for BPMN models as NFTs, as well as a more rigorous evaluation of tokenized business process models using special methods and algorithms, rather than human judgment. The DApp should evolve into the full-scale ecosystem of tokenized BPMN diagrams that could be traded and swapped in the same way it could be done with cryptocurrency and NFTs nowadays.

5. References

- [1] B. Hines, What Types of Assets May be Tokenized?, in: *Digital Finance: Security Tokens and Unlocking the Real Potential of Blockchain*, John Wiley & Sons, 2020, pp. 47–60.
- [2] C. Gerth, Business Process Modeling, in: *Business Process Models: Change Management*, Lecture Notes in Computer Science, Springer, 2013, pp. 14–23. doi:10.1007/978-3-642-38604-6.
- [3] A.-W. Scheer, Strategic Business Process Analysis, in: *ARIS – Business Process Modeling*, Springer Science & Business Media, 2013, pp. 7–20. doi:10.1007/978-3-642-57108-4.
- [4] A. Bitkowska, The relationship between Business Process Management and Knowledge Management – selected aspects from a study of companies in Poland, in: R. Gabryelczyk, T. Hernaus (Eds.), *Business Process Management: Current Applications and the Challenges of Adoption*, Cognitione Foundation for Dissemination of Knowledge and Science, 2020, pp 169 – 193. doi:10.7341/20201616.
- [5] J. Krogstie, Quality of Business Process Models, in: *Quality in Business Process Modeling*, Springer, 2016, pp. 53–102. doi:10.1007/978-3-319-42512-2.
- [6] T. Allweyer, BPMN – A Standard from Business Process Modeling, in: *BPMN 2.0: Introduction to the Standard for Business Process Modeling*, BoD – Books on Demand, 2016, pp. 9–14.
- [7] BPMN Tool Matrix. URL: <https://bpmnmatrix.github.io/>.
- [8] SourceForge – Search results for “bpmn”. URL: <https://sourceforge.net/directory/?q=bpmn>.
- [9] M. Dumas, M. La Rosa, J. Mendling, H. A. Reijers, Essential Process Modeling, in: *Fundamentals of Business Process Management*, Springer, 2018, pp. 75–116. doi:10.1007/978-3-662-56509-4.
- [10] M. H. Rehmani, Introduction to Blockchain Systems, in: *Blockchain Systems and Communication Networks: From Concepts to Implementation*, Springer Nature, 2021, pp. 3–14. doi:10.1007/978-3-030-71788-9.

- [11] J. Jayapriya, N. Jeyanthi, Distributed Consensus Protocols and Algorithms, in: K. Saini, P. R. Chelliah, D. K. Saini (Eds.), *Essential Enterprise Blockchain Concepts and Applications*, CRC Press, 2021, pp. 15–38.
- [12] H. L. Gururaj et al., Blockahin: A New Era of Technology, in: G. Shrivastava, D.-N. Le, K. Sharma (Eds.), *Cryptocurrencies and Blockchain Technology Applications*, John Wiley & Sons, 2020, pp. 3–24.
- [13] T. Laurence, Second generation applications of Blockchain technology, in: *Introduction to Blockchain Technology*, Van Haren, 2019, pp. 69–82.
- [14] C. Blackburn, The Which and How's!, in: *Ethereum Blockchain Revolution Explained: Understanding Ethereum Technology For Beginners*, Christopher Blackburn, 2020, pp. 5–13.
- [15] I. Bashir, Tokenization, in: *Mastering Blockchain: A deep dive into distributed ledgers, consensus protocols, smart contracts, DApps, cryptocurrencies, Ethereum, and more*, 3rd Edition, Packt Publishing Ltd, 2020, pp. 567–614.
- [16] B. Ramamurthy, Tokenization of assets, in: *Blockchain in Action*, Simon and Schuster, 2020, pp. 227–248.
- [17] I. Roy, Building a Blockchain Wallet for Fungible and Non-Fungible Assets, in: *Blockchain Development for Finance Projects: Building next-generation financial applications using Ethereum, Hyperledger Fabric, and Stellar*, Packt Publishing Ltd, 2020, pp. 21–57.
- [18] Tether USD (USDT) Token Tracker. URL: <https://etherscan.io/token/0xdac17f958d2ee523a2206206994597c13d831ec7>.
- [19] Cryptopedia Staff, CryptoKitties: A Pioneer in Ethereum Gaming and NFTs, 2021. URL: <https://www.gemini.com/cryptopedia/cryptokitties-nft-crypto-ethereum-token>.
- [20] Tokens – OpenZeppelin Docs. URL: <https://docs.openzeppelin.com/contracts/2.x/tokens>.
- [21] F. Härer, Decentralized business process modeling and instance tracking secured by a blockchain, in: *Proceedings of the 26th European Conference on Information Systems (ECIS2018)*, Portsmouth, UK, 2018, pp. 1–17. doi:10.5281/zenodo.2585717.
- [22] H. G. Fill, F. Härer, Knowledge blockchains: Applying blockchain technologies to enterprise modeling, in: *Proceedings of the 51st Hawaii International Conference on System Sciences*, 2018, pp. 1–10. doi:10.24251/HICSS.2018.509.
- [23] R. Hull, Blockchain: distributed event-based processing in a data-centric world, in: *Proceedings of the 11th ACM International Conference on Distributed and Event-based Systems*, 2017, pp. 2–4. doi:10.1145/3093742.3097982.
- [24] W. Viriyasitavat, D. Hoonsopon, Blockchain characteristics and consensus in modern business processes, *Journal of Industrial Information Integration* 13 (2019) 32–39. doi:10.1016/j.jii.2018.07.004.
- [25] V. A. De Sousa, B. Corentin, Towards an integrated methodology for the development of blockchain-based solutions supporting cross-organizational processes, in: *2019 13th International Conference on Research Challenges in Information Science*, 2019, pp. 1–6. doi:10.1109/RCIS.2019.8877045.
- [26] F. Milani, L. Garcia-Banuelos, Blockchain and principles of business process re-engineering for process innovation, *arXiv* 1806.03054 (2018). doi:10.48550/arXiv.1806.03054.
- [27] S. Jaitly, What are NFTs and Why are They Becoming Popular?, 2021. URL: <https://medium.com/tribalscale/what-are-nfts-and-why-are-they-becoming-popular-c3ca2c84a4b3>.
- [28] The Shrimpy Team, The Best Smart Contract Platforms, 2021. URL: <https://academy.shrimpy.io/post/the-best-smart-contract-platforms>.
- [29] A. Voronin (Ed.), *Multicriteria Problems in Complex Systems Management*, in: *Multi-Criteria Decision Making for the Management of Complex Systems*, IGI Global, 2017, pp. 17–44. doi:10.4018/978-1-5225-2509-7.ch002.
- [30] ERC721. URL: <https://docs.openzeppelin.com/contracts/2.x/erc721>.
- [31] B. Whittle, Storing Documents on the Blockchain: Why, How, and Where, 2018. URL: <https://coincentral.com/storing-documents-on-the-blockchain-why-how-and-where/>.
- [32] The GitHub repository of a software prototype. URL: <https://github.com/andriikopp/blockchain-repository/tree/main/nft>.

The Hierarchical Information System for Management of the Targeted Advertising

Karina Melnyk¹, Natalia Borysova¹ and Viktoriia Melnyk²

¹ National Technical University "Kharkiv Polytechnic Institute", Kirpichova street, 2, Kharkiv, 61002, Ukraine

² Kharkiv general education school of I-III degrees № 145, Amosova street, 24a, Kharkiv, 61171, Ukraine

Abstract

The problems of target audience identification and user segmentation for managing a process of the targeted advertising to customers are considered. It has proposed to combine the solution of these two tasks within the framework of one information system. An overview of the existing methods for identifying and segmenting of the target audience are presented. In addition, the existing applications and tools for solving the given problem are considered. The formalization of these tasks is presented. The functional models of business processes corresponding to the identifying and segmenting of the target audience are developed. The architecture of the information system is proposed. It is presented in the form of a deployment scheme, a database model is developed. The results of numerical studies and evaluation of the effectiveness of the developed information system are presented.

Keywords

Targeted advertising, target audience, buyer persona, customer segmentation, classification methods, K-Means Clustering, similarity measure.

1. Introduction

Profitability of any B2B or B2C company depends on many factors. One of the important one is advertising, which presents a product or service to a potential buyer for the selling purpose. Development of information technology has led to possibility to identify the target audience of customers and to influence on this audience using targeted advertising. Many various methods of Market Research are used to determine the target audience. Depending on the needs of the market, the strategic goals of an enterprise and characteristics of the advertised product or service, all potential buyers are divided into groups. There are several cases of dividing the customers into groups. Manager can create two groups: included in the target audience and not included in the target audience; three groups: primary target audience, secondary target audience, not included in the target audience; four groups: core, primary target audience, secondary target audience, not are included in the target audience. The advantages of determining the target audience are obvious:

- development of a more effective advertising campaign is a reason of saving the advertising budget;
- increasing of a brand awareness and loyalty of clients contribute to a class of regular customers;
- quick return on investment;
- a clearer understanding of the needs of their customers, it means more effective interaction with them, etc.

However, to improve the effectiveness of an advertising campaign, manager can divide the target audience into segments based on the personal characteristics of the clients. The customer segmentation models are used in this case. There are many benefits of this process. The segmentation

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022

EMAIL: karina.v.melnyk@gmail.com (K. Melnyk); borysova.n.v@gmail.com (N. Borysova); v13121423@gmail.com (V. Melnyk)

ORCID: 0000-0001-9642-5414 (K. Melnyk); 0000-0002-8834-2536 (N. Borysova); 0000-0003-2958-3935 (V. Melnyk)



© 2022 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

CEUR Workshop Proceedings (CEUR-WS.org)

process can be used for promotion of a certain product to a group of buyers from the target audience. Other advantages of using the segmentation models are following: drawing up personalized price proposals, improving customer experience, searching for ideal customer, reducing the customer churn, prices optimization, searching for new market opportunities, etc.

Therefore, the purpose of this research is to develop an integrated hierarchical information system (HIS) that would combine the functions of systems for determining the target audience and a system of user segmentation for the distribution of targeted advertising.

2. Formal problem statement

The Hierarchical Information System is a system that solves tasks stage-by-stage. The first task is to identify the target audience. It is a classification task since it divides customers into several predefined groups. Input information is a set of characteristics of potential buyers. To solve the problem, it is necessary to do:

- analyze the domain area for creating a portrait of buyer persona;
- undertake review of approaches and methods for determining the targeted audience;
- formalize a method for the resolving the task.

The second level of the hierarchical management system is designed to solve the problem of segmentation of the target audience. This task is a clustering task and belongs to the class of unsupervised learning tasks. Potential customers should be grouped into clusters in such a way that objects from one cluster are closer to one another than objects from other clusters by any criterion. To solve the clustering problem, it is necessary to do:

- analyze methods and applications to solve the clustering task;
- formalize proposed method;
- perform numerical study;
- evaluate the proposed approach.

3. Overview of existing models, methods and applications to solve the given issue

This study proposes to solve the given issue in two stages. Therefore, it is necessary to conduct review the existing methods and applications separately for each stage: for identifying the target audience and for customer segmentation. There are many generally accepted models in direct marketing, which can be used for both target audience identifying and customer segmentation. The usage of these models depends on user's data (Table 1) [1-4].

Table 1

Models for target audience identifying and customer segmentation

Model name	Characteristics
Demographic based	Age, gender, occupation, income, education, marital status, ethnicity, race, religion, profession or role in the company
Geographic based	Continent, country, region or state, city, district, postal code, timezone
Psychographic based	Social class, lifestyle, personality, values, presence in digital and/or social media space, personal convictions, beliefs, attitude, interests
Technographic based	Usage of devises, applications and software
Behavioral	Habits, spending, consumption, usage and desired benefits, usage, loyalty, awareness, types of payment, demands, quality fanatics, price and/or brand sensitiveness

In addition, it is possible to use multiple models or mixed models.

3.1. Review of target audience identifying methods and applications

Any classification method can be used for resolving of the identifying task of the target audience. For example, Naive Bayes, logistic regression, Support Vector Machine, k-nearest neighbor, decision trees etc. allow dividing objects into different groups [5-9]. The set of buyers is divided into two classes: included in the target audience and not included in the target audience. The classification signs are different characteristics of the portrait of an ideal buyer. The portrait has built before the launch of an advertising campaign for a product or service and can be adjusted. The characteristics correspond to one or more used models or a mixed model (Table 1). The best case when a marketing specialist who understands the intricacies of promoting a company's product or service draws up the portrait. She/he can use various approaches, for example, Mark Sherrington's "5W" approach. 5W means answering such questions: What? (type of product or service we sale); Who? (our customer); Why? (motivation for buying); When? (conditions for buying); Where? (place of buying). There are alternative marketing approaches to the 5W method: the method "from the opposite", the method "from the product", the method "from the market", and the method "from the target" [1]. The result of using any of the marketing methods is a set of values for the characteristics of an ideal buyer. This is the most important and crucial stage for the further success of the advertising campaign. It influences on the size of the company's profit. An interesting tool from HubSpot named Make My Persona [10] can come in handy when drawing up a portrait of an ideal buyer. This tool allows building some virtual image of buyer personas in seven simple steps:

1. Creating Personas' avatar and choosing Personas' name.
2. Identifying Personas' demographic characteristics, such as age and level of education.
3. Identifying Personas' business, such as working industry and size of company.
4. Identifying Personas' careers, such as their job title, their job measuring, their boss.
5. Identifying Personas' job characteristics, such as their goals or objectives, their biggest challenges, their job responsibilities
6. Identifying Personas' lovely tools to do their job and to communicate with vendors and other businesses.
7. Identifying Personas' consumption habits, such as their lovely social networks and their training in their job.

After answering these questions, the user is taken to a page where she/he can see her/his Buyer Personas 'Overview, and can also expand it by adding own fields to this Overview. It can be saved, downloaded or shared on social networks after filling out a special pop-up form.

Existing applications, services, tools and software, which are intended for target audience identifying, for example, such as Google Analytics, Facebook Insights, Twitter Followers Dashboard and others can only be used to analyze the company's existing customer data, creating various analytical reports, visualizing analysis results, making forecasts, tracking the activity of regular customers, inflow and outflow of customers, etc. However, all of these feature-rich applications, services, tools and software have "a cold start" problem. They cannot be used for analysis in the absence of data. For example, when some startup is launched and the portrait of the ideal buyer has already been built, manager have to create advertising, but target audience is unknown. Obviously, the target audience identifying task has not been completely solved, and in some cases it has not been solved at all.

Thus, in order to identify potential buyers, it is necessary to form a portrait of the buyer persona and compare it with the considered objects. There are many ways of comparing objects. For instance, manager of a company can perform calculation of similarities between objects. Input data of buyer persona and potential customers are information of mixed type. Therefore, to calculate a similarity between them, are encouraged to use corresponding metrics. It can be Voronin similarity measure, Zhuravlev metric, Gower coefficient etc. [11].

3.2. Review of customer segmentation methods and applications

After identifying of the target audience, HIS can start customer segmentation. Manager can use rule-based methods or cluster analysis methods for customer segmentation. The use of rule-based

methods implies the creation of rules or the selection of a priori thresholds for strict customer segmentation. Such way of segmentation can lead to situation, when customers from one group have significant differences. It is also quite difficult to perform segmentation in more than two dimensions. In addition, the segmentation results are more consistent with the initial assumptions of the marketer and do not always reveal significant differences between customers. In this sense, cluster analysis methods show higher efficiency in comparison with rule-based methods. Since they are unsupervised machine learning methods, they do not need a training sample; they can work directly with the input data without prior training. These methods are more practical, they divide the training set of customers into more homogeneous groups, within the differences between customers are very small, in addition, cluster analysis methods allow to conduct the dynamic clustering, which is fully reflect the state of the available data at a given time [12].

There are special applications and software for customer segmentation on the market. For example, Segmentor the customer segmentation tool from Optimove [13], CleverTap [14], HubSpot [15], Experian [16], SproutSocial [17], Qualtrics [18], MailChimp [19] etc. In addition, there are applications and software for customer segmentation that only use the behavioral model. For example, Yieldify [20], Amplitude [21], Indicative [22], Mixpanel [23]. All of them certainly have their own advantages and perform the function of customer segmentation using one or several models. However, not all of them are free, but it is important for novice businesspersons or startups. In addition, there is no description of the used algorithms and methods, which is not will allow the user to double-check the obtained results and may lead to erroneous conclusions.

3.3. Review of targeted advertising management applications and services

After segmentation of the target audience, it is necessary to prepare special advertisements for each group of customers and send them out. For this, it is advisable to use special targeted advertising management applications and services. The article [24] provides a brief description of some of these services. Of course, there is a huge number of such applications and services, there are free and paid ones, they have different functionality, and some even solve the problem of finding the target audience in addition. Nevertheless, none of them solves three problems: target audience search, target audience segmentation and targeted advertising delivery.

4. Development of the hierarchical information system

4.1. Formalization of the management process of the targeted advertising

Let's consider designing and using the hierarchical information system to resolve the management task of the Targeted Advertising in a more detailed way. To solve the problem of determining the targeted audience and the segmentation of potential users, it is necessary to develop a functional model of the business process for managing targeted advertising delivery. For this, it is proposed to use the IDEF0 methodology (Figure 1).

The first step is marketing research. It allow to find, to collect and to analyze the received information for reducing the uncertainty in making managerial decisions. For example, it is important for a new startup to find a product or service that can be successfully implemented. For existing companies, it is possible to assess the prospects for demand for a specific product or service based on the study of consumer behavior. There are many methods to perform market research: observation, surveys, focus groups, personal interviews etc. Each method has its own advantages, disadvantages, limitations and is capable of providing information of varying completeness and accuracy.

The goal of the next stage is to highlight informative features that will fully reflect the portrait of a potential buyer. Next, a profile of ideal buyer persona is created. Defining a buyer or audience persona helps to create product or service to better target ideal customer. Depending on the final product is designed for, the appropriate templates are used: B2B or B2C Buyer Persona templates, as well as the results of marketing research.

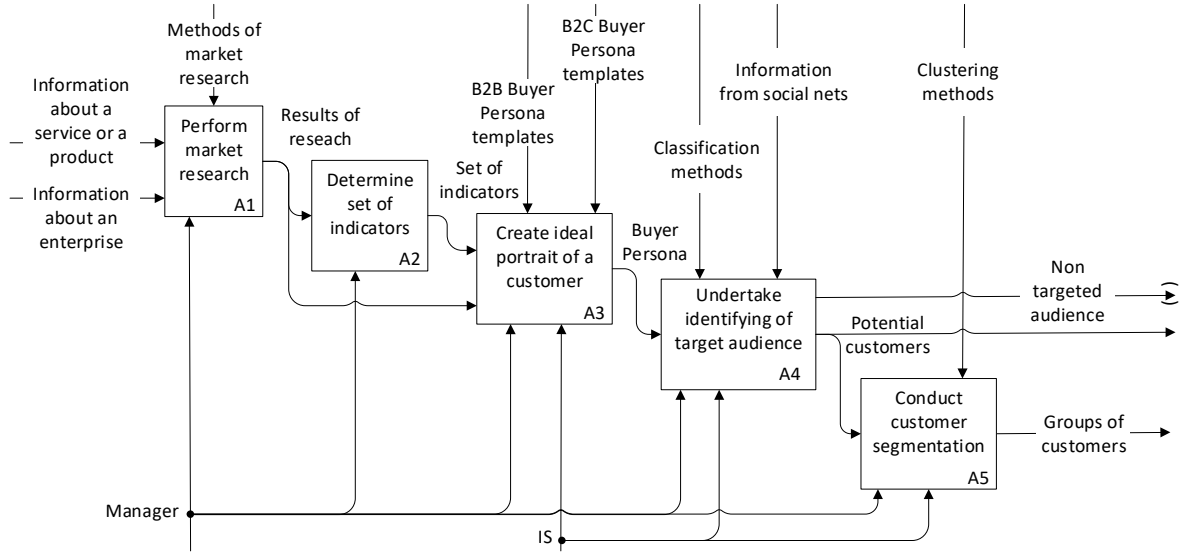


Figure 1: The model of Management of the Targeted Advertising

The next step is to split buyers into two groups: targeted audience and non-targeted audience. Various classification methods can be used for this. According to the above analysis, the paper proposes to use an approach based on calculating a measure of similarity between an ideal customer and a potential customer. To improve the effectiveness of an advertising campaign, the target audience can be segmented depending on the needs of potential customers. It is necessary to solve the clustering task according to the selected indicators. For example, if you advertise certain sport clothes, then the younger generation will choose bright colors, and the middle age will prefer convenience to appearance. In general, the task of segmenting potential customers can be reduced to the task of determining the target audience. The expert for each group sets its own boundary values of the similarity measure. Thus, she/he can create several target groups: core audience, several groups with different values of similarity measures, non-targeted people. This approach can be applied in the case of limited finances. However, the expert can miss significant differences between clients, therefore, for a more efficient segmenting process, it is recommended to use automatic segmentation, namely clustering methods, since these are methods of unsupervised learning.

4.2. The model of resolving the identifying task of the target audience

Consider the task of identifying potential customers. The model for solving the task is described by the following activity diagram in Figure 2.

Let K be the set of clients selected by the manager to determine the value for the advertising company. The input data for the task of identifying the target group is an ideal customer profile. Let's specify F as a set of indicators of a potential client, and $F_i (F_i \in F)$ is a set of values of the i -th indicator. Then x_i^b and $x_{ki} (k \in K, i \in F_i)$ are the value of the i -th indicator of profile of buyer persona and of k -th client accordingly.

Let's designate $s_k (k \in K)$ as a measure of similarity between the profile of an ideal client and k -th client, then s_{ki} is a similarity in the no i -th indicator. Data about a potential customer can be qualitative data, in the form of categories, and in the form of numbers. For the simultaneous processing of such data, there are special proximity measures. Let us consider the calculation of the similarity on the example of using the Gower coefficient [11, 25]. To calculate the measure, it is necessary to turn qualitative data into categorical data, and then use formula (1) for them:

$$s_{ki} = \begin{cases} 1, & x_i^b = x_{ki} \\ 0, & x_i^b \neq x_{ki} \end{cases}. \quad (1)$$

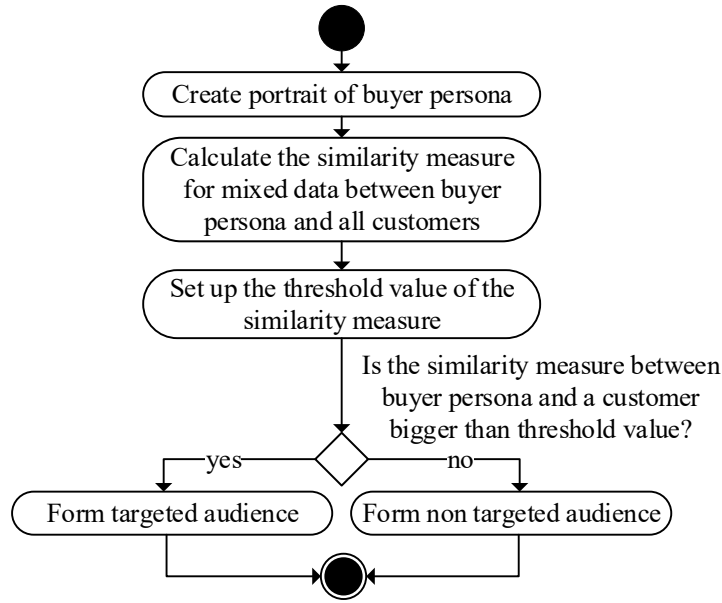


Figure 2: The model of the identifying process of the targeted audience

If the indicators of the portrait of audience persona are quantitative, then formula (2) is applied:

$$s_{ki} = 1 - \frac{|x_i^b - x_{ki}|}{\max\{x_i^b\} - \min\{x_i^b\}} \quad (2)$$

To calculate the Gower coefficient s_k according to formula (3), it is necessary to determine w_i is the coefficient of the presence of the i -th indicator for the client: $w_i = 0$, if information about the indicator is absent in the profile of the ideal client or potential client, it equals to 1 if all the information is available .

$$s_k = \frac{\sum_i w_i s_{ki}}{\sum_i w_i} \quad (3)$$

Next, the manager determines the similarity values that are acceptable for the target group of buyers and forms the corresponding sets. The obtained information allows formulating solutions for the implementation of an advertising campaign.

4.3. The model of resolving the segmenting task

Review of related works according to resolving the segmenting task has showed the feasibility of using the clustering methods. Clustering process is a machine learning technique that groups of objects according to chosen indicators. Modern science knows a lot of clustering methods: Affinity Propagation, Balanced Iterative Reducing, K-Means Clustering, Clustering using Hierarchies, Mean shift clustering, Agglomerative Hierarchical Clustering, Expectation–Maximization Clustering using Gaussian Mixture Models etc. [26, 27]. One of the simplest and most commonly used method is the K-Means Clustering method. Let's consider a model for solving the task of segmenting a targeted audience using the proposed method (Figure 3).

Let the manager has a hypothesis about the number n of clusters or segments. It can be based on theoretical considerations or the results of marketing research. If the assumptions are not obvious, then a series of experiments with different numbers of clusters can be performed to find the optimal partition.

The K-Means algorithm starts with randomly selected clusters and then reassigns objects to them to minimize intra-cluster variability and maximize inter-cluster variability. The main drawback of the K-Means algorithm is that it only works with numeric values. The input customer information is mixed information. Therefore, in this paper, it is proposed to use the Gower coefficient to calculate the distances between the centers of the clusters and the objects, since it works with mixed data.

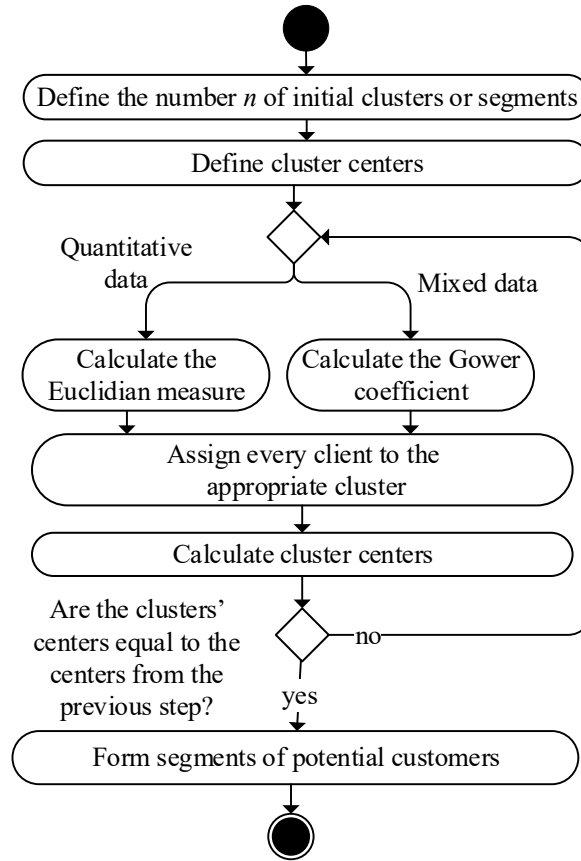


Figure 3: The model of the segmenting process of the targeted audience

The algorithm of K-Means is following. The center of each cluster is randomly determined. Denote $c_q(\{x_i\})$ ($q = \overline{1, n}$) as the center of q -th cluster with the set of values of all indicators $\{x_i\}$. Then $\|c_q\|$ is the cardinality of the set of objects of the q -th cluster. Then it is necessary to calculate the distance between the centers of the clusters and each object s_{qt} ($t \in T$) according to formulas (1)-(3), where T is the set of potential buyers or target audience obtained in the previous step, and t is a specific buyer from this set. The object or client is assigned to the closest cluster

$$t_q = \operatorname{argmax}_{q=\overline{1, n}} \{s_{qt}\}, \quad (4)$$

where t_q is designation of belonging of t -th customer to q -th cluster.

After calculating the distances and assigning objects to the new cluster, it is necessary to find the coordinates of the new center for each cluster. The new values of each indicator of the center of the cluster are found based on the use of the formula for the arithmetic mean of all values of the i -th indicator of objects that belong to the current cluster:

$$c_q(\{x_i\}) = \frac{1}{\|c_q\|} \sum_{t \in c_q} x_{ti}, \quad i \in F_t. \quad (5)$$

Next, the algorithm again calculates the distance from each object to the cluster centers using formulas (1)-(3) and assigns the objects to the nearest cluster using formula (4). The centers of gravity of the clusters are calculated again according to (5). This process is repeated until the centers of gravity stop “migrating” in space.

Thus, the model for target audience segmentation has been proposed.

4.4. Architectural solution of the Information System

The HIS design process is revised from the development of the software requirement specification (SRS). All requirements from SRS are divided into functional and non-functional ones. Functional

requirements are subdivided into business requirements, user requirements, and system requirements. Non-functional requirements describe how the HIS should work and which properties and quality attributes it should have. The main requirements described in the SRS are presented in the form of Requirement Diagram (Figure 4).

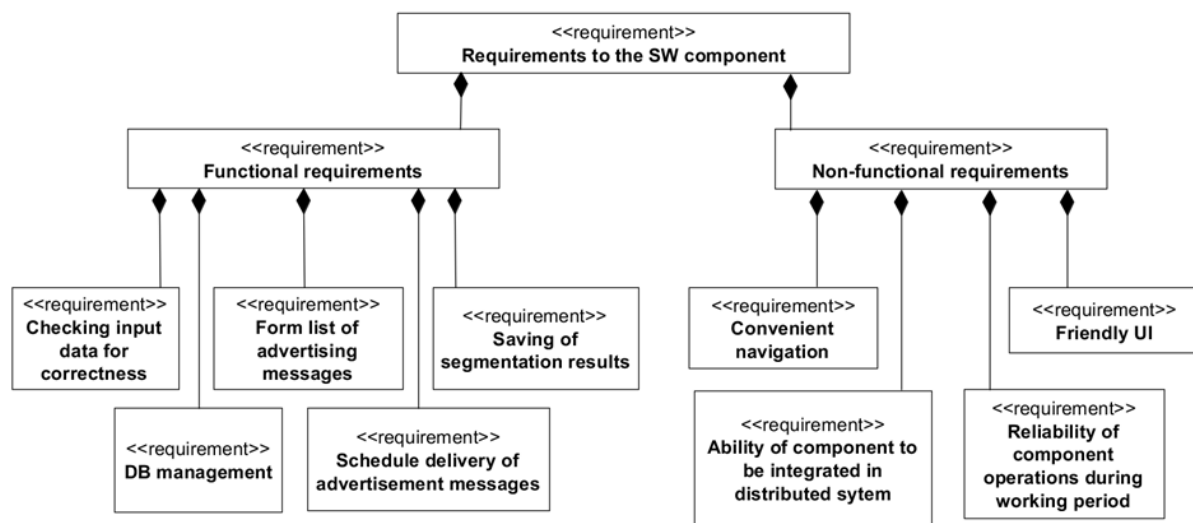


Figure 4: The Requirements Diagram for the HIS

The list of functional requirements for the system is following:

- The HIS has to verify the correctness of the entered data.
- The HIS has to generate a list of advertising messages and advertising objects or generate an error in the process of creating advertising messages.
- The HIS has to create the advertising messages and customer groups after the customer segmentation process is completed.
- The HIS has to provide access to all necessary data and documents (list of product categories, information on discounts, rules of message formation, advertising budget, segmentation rules, list of customer wishes, description of the segmentation method, list of customer preferences) to the advertiser at the stage of formation of advertising messages.
- The HIS has to provide access to the database (Products_DB, Customers_DB etc.) at any time.
- The HIS has to provide opportunities to work with files that have been created in other systems.
- The HIS has to allow the advertiser to suspend his/her work and save to a file, then to resume his/her work from the saved file.
- The HIS has to send advertising messages at the scheduled time or allow the system user to do so manually.

The result of turning out the requirements into a structured solution that meets both technical and business requirements is an architectural solution for HIS. Deployment Diagram is used to visualize the topology of the physical components of the system. The proposed architectural solution for The Hierarchical Information System for the Management of Targeted Advertising in the form of a Deployment Diagram is shown in Figure 5.

This research proposes to use the “client-server” architectural pattern for the development of the HIS architecture. Such architecture allows sharing the data processing function to several separate servers. It separates the functions of storing, processing and presenting data for more efficient use. The presentation component or a “client” is responsible for the user interface. Application logic is executed at the middle level of the architecture, namely the application server layer. It provides data exchange between users and databases. The middle layer is split into two separate components to improve the performance of the HIS:

- IIS Application Server is responsible for the application’s logic;
- IIS Mail Server is devoted to process advertising messages.

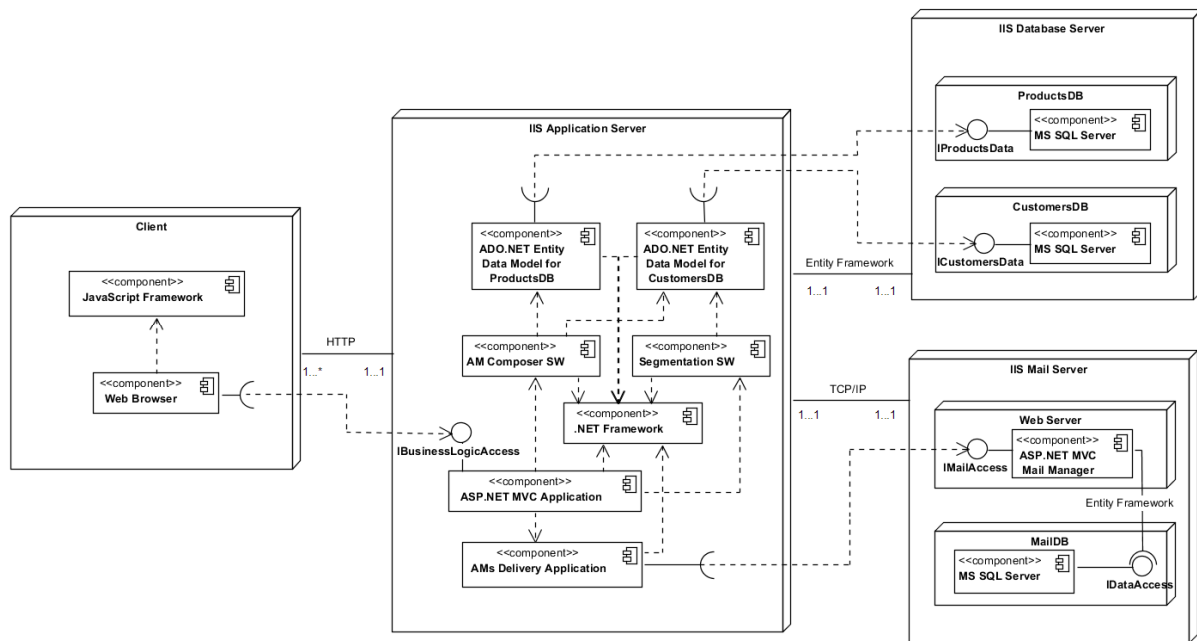


Figure 5: The Architecture of the HIS

The data layer is designed to store and manage the information processed by the HIS. Here is a database that allows to implement the interpretation of information about the domain area in the form of formalized data in accordance with certain requirements (Figure 6).

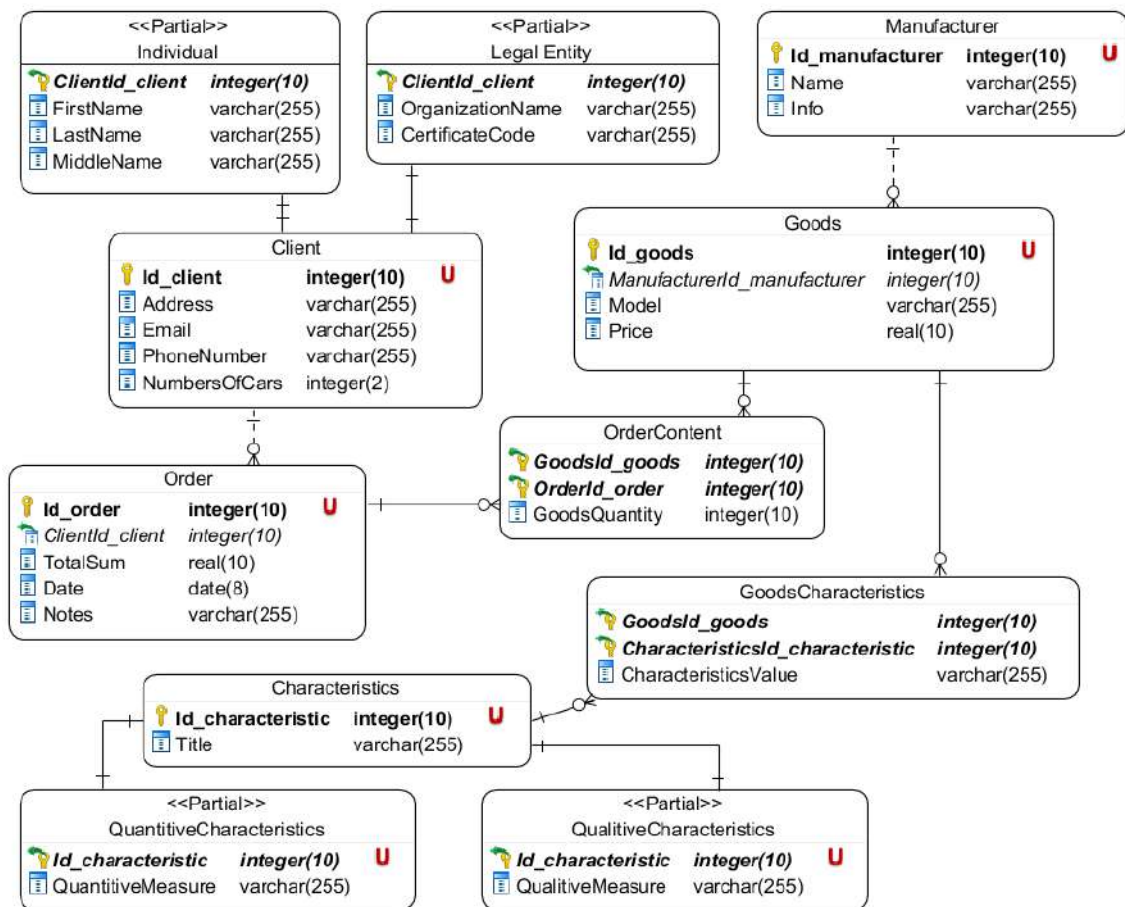


Figure 6: The Database structure

The structure of the developed database consists of 11 entities:

- The entity of “Clients” is the company’s customers.
- The categorical entity “Individual” describes common clients or people.
- The categorical entity of “LegalEntity” describes a legal entity.
- The entity of “Order” is a table with orders.
- The entity of “Manufacturer” is a list of producers of goods.
- The entity of “Goods” is a list of goods.
- The entity of “OrderContent” is the content of the order.
- The entity of “Characteristics” is a list of product characteristics.
- The categorical entity of “QuantitiveCharacteristics” is a list of quantitative characteristics of the product.
- The categorical entity of “QualitiveCharacteristics” is a list of quality characteristics of the product.
- The entity of “GoodsCharacteristics” shows the value of a specific characteristic of a particular product.

Thus, the architectural solution of the hierarchical information system was proposed. It allows solving the task of determining and segmentation of the target audience. In addition, the component, which is responsible for sending messages to clients from the target audience, was developed.

5. Experiments

To check the performance of the developed HIS, three tasks have been solved to promote a new sports club service: target audience identifying task, customer segmentation task and management task of targeted advertising. The target audience has been determined based on the database of the club’s clients and among users of the social networks Instagram and Facebook.

The marketing research of the fitness services market, the communication with clients and analysis of their requests made it possible to see the increased interest of the club’s clients in losing weight, quick recovery after heavy physical exertion, developing muscle endurance and strength, increasing their elasticity, strengthening the cardiovascular system, etc. In this regard, management of the fitness club has been decided to create a new type of service in this fitness center as aqua aerobics. This type of fitness contributes to the normalization of weight, hardening of the body and strengthening its immunity, smoothing out the manifestations of cellulite, increasing skin tone, relieving muscle and emotional stress, neutralizing the negative effects of stress, strengthening the nervous system, normalizing sleep. The aforementioned marketing research would determine the set of indicators for identifying the target audience:

- x_1 – income of clients: x_{11} – low, x_{12} – average, x_{13} – high;
- x_2 – field of activity: x_{21} – office worker, x_{22} – student, x_{23} – business manager, x_{24} – housewife, x_{25} – creative activity;
- x_3 – social media interests in Instagram and/or Facebook: x_{31} – presence of likes, followers and subscriptions of pages with information about sport clubs, diet, weight loss, child goods, mam’s publics, x_{32} – absence of such records, x_{33} – private account or absence of account in social networks;
- x_4 – geographic location of a potential client: x_{41} – residents of the area where the fitness club is located, x_{42} – residents of other areas;
- x_5 – age: x_{51} – less, than 21 years, x_{52} – 21-35 years, x_{53} – 35-45 years, x_{54} – more, than 45 years.

The management of the fitness club has been created the profile of the target audience: x_{12} or x_{13} ; x_{21} or x_{23} or x_{24} ; x_{31} or x_{33} ; x_{41} or x_{42} ; x_{52} or x_{53} or x_{54} . Let’s consider several potential clients of aqua aerobics service. Information from club’s databases, social networks or questionnaires is presented as a set of indicators in Table 2. Analysis of input data allow seeing that some indicators for the portrait of a potential consumer have several values. It, in turn, adds uncertainty during the deciding whether to include a particular potential client in the target group. The usage of the proposed

technology could clearly define the future client based on the varying process of the threshold value, which helps to increase or filter the customer base.

Table 2

Result of the resolving the target audience identifying task

Potential client	Values of the input indicators					Gower coefficient S_k	Threshold value	Belonging to the target audience
	x_1	x_2	x_3	x_4	x_5			
1	x_{12}	x_{23}	x_{33}	x_{41}	x_{53}	1	0,55	+
2	x_{11}	x_{21}	x_{31}	x_{42}	x_{52}	1		+
3	x_{13}	x_{24}	x_{31}	x_{42}	x_{53}	1		+
4	x_{11}	x_{25}	x_{32}	x_{41}	x_{54}	0,4		-
5	x_{12}	x_{22}	x_{31}	x_{41}	x_{51}	0,6		+

The customers 1-3 and 5 fell into the target group with the current threshold as can be seen from Table 2. If the threshold value increases to 0.61 to refine the set of potential customers, then we will consider only customers 1-3. Despite the fact that two of them live in a different area, the scope of their employment and interests allows to offer them the considered service.

Second stage of proposed technology is resolving the customer segmentation task for the potential clients of aqua aerobics service. Analysis of customer preferences and the impact of aqua aerobics on the human body shows that potential customers could be divided into 3 segments. The first segment is responsible for improving the muscles and visible changes of the body of clients, that is, this service can be used as an additional training. The aim of the second group of clients is psychological and physical rehabilitation, that is, to relieve muscle and emotional tension and problems. The purpose of the third segment of clients is the organization of leisure, that is, more entertaining than strengthening the body. The result in this case is the acquisition of new acquaintances, an interesting pastime and an increase in self-esteem. To solve the segmentation task, the following indicators have been proposed:

- x_1 – health status: x_{11} – poor health, x_{12} – average state, x_{13} – good health;
- x_2 – social identification: x_{21} – business people, x_{22} – student, x_{23} – athlete, x_{24} – former athlete, x_{25} – housewife;
- x_3 – social media interests in Instagram and/or Facebook: x_{31} – sport clubs, x_{32} – diet, x_{33} – weight loss, x_{34} – child goods, x_{35} – pools, x_{36} – private account or absence of account in social networks;
- x_4 – free time: x_{41} – all day, x_{42} – weekend, x_{43} – weekday evenings;
- x_5 – age: x_{51} – less, than 21 years, x_{52} – 21-35 years, x_{53} – 35-45 years, x_{54} – more, than 45 years;
- x_6 – marital status: x_{61} – single, x_{62} – married, x_{63} – have kids, x_{64} – childless;
- x_7 – regularity of training: x_{71} – 0-1 per week, x_{72} – 2-3 per week, x_{73} – more than 3 times per week;
- x_8 – possible problems of the target audience: x_{81} – a weight problem, x_{82} – problems with communication, x_{83} – problems with appearance, x_{84} – bad mood, x_{85} – narrow social circle, x_{86} – low immunity, x_{87} – absence of complaints.

The snippet of input data for ten potential clients and result of resolving the segmentation task is presented in Table 3. The clients 1-3 were chosen as initial centers for the three clusters. Result of using the K-Means Clustering method at the first stage is the set of values of similarity measure between clients 1-3 and 4-10. It is a base for defining the belonging of each client to one of the clusters. Next step is to calculate new cluster centers. After the sixth iteration, the cluster centers have stabilized.

An analysis of the content of each segment shows that the first cluster includes middle-aged housewives with overweight problems and a former athlete. Based on the values of other indicators, it can be understood that this segment of clients is aimed at rehabilitation, support and restoration of health. The second cluster of potential customers includes those people who are seriously into sports. For this group, it is necessary to create special training programs. The third segment is a group of

young people, mostly unmarried, who are focused not only on sports, but also on a pleasant and interesting pastime.

Table 3

The result of the segmenting process of the targeted audience

Client	Values of the input indicators								Number of iteration and content of each cluster					
	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	1	2	3	4	5	6
1	x_{13}	x_{21}	x_{31}	x_{43}	x_{52}	x_{61}	x_{72}	x_{87}	I 1,	I 4,	I 3,	I 5,	I 2,	I 2,
2	x_{12}	x_{25}	x_{33}	x_{41}	x_{53}	x_{62}	x_{73}	x_{81}	4, 8,	5, 8	5, 10	8, 10	5, 10	5, 10
3	x_{11}	x_{22}	x_{31}	x_{43}	x_{51}	x_{61}	x_{71}	x_{82}	9					
4	x_{12}	x_{21}	x_{31}	x_{43}	x_{53}	x_{63}	x_{71}	x_{86}		II 2,	II 2,	II 2,	II 4,	II 4,
5	x_{13}	x_{25}	x_{32}	x_{41}	x_{52}	x_{64}	x_{73}	x_{83}	II 2,	6, 7,	4, 6,	4, 6,	II 4,	6, 7
6	x_{12}	x_{23}	x_{31}	x_{42}	x_{53}	x_{63}	x_{73}	x_{87}	5, 6,	10	7, 8	7	6, 7	
7	x_{11}	x_{24}	x_{34}	x_{42}	x_{53}	x_{63}	x_{72}	x_{81}	7, 10					
8	x_{11}	x_{21}	x_{35}	x_{43}	x_{52}	x_{62}	x_{71}	x_{84}		III 1,	III 1,	III 1,	III 1,	III 1,
9	x_{13}	x_{23}	x_{35}	x_{43}	x_{52}	x_{64}	x_{73}	x_{82}	III 3	3, 9	9	3, 9	3, 8,	3, 8,
10	x_{13}	x_{24}	x_{36}	x_{42}	x_{54}	x_{63}	x_{73}	x_{85}					9	9

Thus, the practical obtained result shows the feasibility of using the proposed technology in the article in real conditions to create a target audience.

6. Discussion

Efficiency estimation of the HIS was carried out separately for each task: target audience identifying task, customer segmentation task and management task of targeted advertising. The classification method based on the calculation of the similarity measure has been used for target audience identifying task. Therefore the standard metrics for assessing the effectiveness of the Precision and Recall classification were chosen to assess the effectiveness of HIS. The Confusion matrix with the corresponding indicators for calculating these metrics is presented in Table 4.

Table 4

Confusion matrix for Precision and Recall metrics evaluation

		Predicted	
		Positive	Negative
Actual	Positive	TP = 1123	FN = 77
	Negative	FP = 41	TN = 429

Data for the matrix is the result of comparison of the results of the work of HIS and the opinion of an expert. The classifier assigned 1164 clients from 1670 to the target audience. It incorrectly attributed 41 clients to the target audience and did not attribute 77 clients from the target audience to it. The formulas for calculating and Precision and Recall metrics values are presented below:

$$Precision = \frac{TP}{TP + FP} = \frac{1123}{1123 + 41} = 0.96,$$

$$Recall = \frac{TP}{TP + FN} = \frac{1123}{1123 + 77} = 0.94$$

The values of the Precision and Recall metrics are quite high and very close to one, which indicates a high classification efficiency.

K-means clustering algorithm was used for resolving the customer segmentation task. The Rand index (RI) has been chosen for evaluating the quality of clustering algorithm. The RI calculates a proximity measure between two clusters based on the comparing results of pairs that are assigned in the same or different clusters in the predicted and true clustering process. The input data for clustering were 1200 clients of the target audience. The contingency table with the indicators values for calculating the RI is presented in Table 5.

Table 5

Contingency table for Rand index evaluation

	Same cluster	Different clusters
Same class	SS = 337	DS = 61
Different classes	SD = 53	DD = 749

An expert compared the clustering results of the HIS with the own results of the distribution of clients into groups. The formula for calculating and the value of the Rand index are presented below:

$$Rand = \frac{SS + DD}{SS + DS + SD + SS} = \frac{337 + 749}{337 + 61 + 53 + 749} = 0.91$$

The obtained value of the RI is rather close to one. It indicates an almost complete coincidence of clusters and classes, which showed a high efficiency of clustering.

Conversion rate (CR) was used to evaluate the effectiveness of targeted advertising mailing. The target audience of 1200 clients was split into two clusters. The first cluster included 418 clients, the second one 782. Different advertisements about the new service of the sports club were prepared for the clients of each cluster. The set of data was prepared according to the customers' actions within a month from the date of sending advertisements. The formula for CR calculating and CR values for the two clusters are presented below:

$$CR = \frac{\text{purchase}}{\text{total number of ads}} 100\%,$$

$$CR_{1cluster} = \frac{57}{418} 100\% = 13.64\%,$$

$$CR_{2cluster} = \frac{108}{782} 100\% = 13.81\%.$$

In SMM, an advertising company is considered as successful if the CR values have reached 3-5%. In our case, the rather high CR values can be explained by the fact that the target audience included both real customers from the sports club's database and new customers found in social networks.

7. Conclusion

In this study, the approach was proposed for solving the task of dividing potential customers into non-targeted audience and targeted audience with additional segmentation for a more effective advertising campaign. The methods and existing applications for solving the given task were considered. The model of resolving the identifying task of the targeted audience and the model of the segmenting process of the targeted audience were developed. The architectural solution for the HIS has been developed on the base of the chosen architectural pattern "client-server". The conducted experiment and the assessment of the performance of the HIS have showed the feasibility of usage of the developed HIS in real conditions for managers that conduct an advertising campaign in order to attract new customers and improve the financial condition of the enterprise.

8. References

- [1] Matan Naveh, How To Identify a Target Audience for Your Business, 2022. URL: <https://elementor.com/blog/how-to-identify-target-audience/>.
- [2] Customer Segmentation Models, 2017. URL: <https://medium.com/think-with-startupflux/customer-segmentation-models-52ef7738823a>.
- [3] What is Customer Segmentation? – Types, Techniques, Models. URL: <https://survicate.com/customer-segmentation/what-is-customer-segmentation/>.
- [4] K. Baker, The Ultimate Guide to Customer Segmentation: How to Organize Your Customers to Grow Better, 2020. URL: <https://blog.hubspot.com/service/customer-segmentation>.
- [5] E. F. Ayetiran, A. B. Adeyemo, A Data Mining-Based Response Model for Target Selection in Direct Marketing, *International Journal of Information Technology and Computer Science* 4(1) (2012). doi:10.5815/ijits.2012.01.02.
- [6] E. W. Maibach, A. Leiserowitz, C. Roser-Renouf, C. K. Mertz, Identifying like-minded audiences for global warming public engagement campaigns: an audience segmentation analysis and tool development, *PLoS ONE* 6(3) (2011). URL: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0017571> doi:10.1371/journal.pone.0017571
- [7] G. Tirenni, C. Kaiser, A. Herrmann, Applying decision trees for value-based customer relations management: Predicting airline customers' future values, *J Database Mark Cust Strategy Manag* 14 (2007) 130–142. doi:10.1057/palgrave.dbm.3250044.
- [8] K. Melnyk, S. Kirkin, Intelligent Data Processing in Creating Targeted Advertising, in: *Proceedings of the 1st International Conference Computational Linguistics And Intelligent Systems*, volume 1 of COLINS 2017, NTU «KhPI», Kharkiv Ukraine, 2017, pp. 131–132.
- [9] M. Karim, R. M. Rahman, Decision Tree and Naïve Bayes Algorithm for Classification and Generation of Actionable Knowledge for Direct Marketing, *Journal of Software Engineering and Applications* 6 (4) (2013). URL: https://www.scirp.org/html/6-9301587_30463.htm. doi:10.4236/jsea.2013.64025.
- [10] HubSpot tools. Make My Persona. A Buyer Persona Generator from HubSpot. URL: <https://www.hubspot.com/make-my-persona>
- [11] K. Melnyk, N. Borysova, Integrated Technology for Personnel Assessment Based on the Competencies Model, in: T. Hovorushchenko, A. Pakštas, V. Vychuzhanin, H. Yin and N. Rudnichenko (Eds.), *Proceedings of the 9th International Conference “Information Control Systems & Technologies”, ICST-2020, Odessa, 2020*, pp. 343–357. URL: <http://ceur-ws.org/Vol-2711/paper27.pdf>. doi:10.13140/RG.2.2.26024.60169.
- [12] Customer Segmentation via Cluster Analysis, 2020. URL: <https://www.optimove.com/resources/learning-center/customer-segmentation-via-cluster-analysis>.
- [13] Segmentor. Customer segmentation tool, 2022. URL: <https://segmentor.optimove.com/#/>.
- [14] CleverTap. Audience segmentation. Build actionable user segments with ease, 2022. URL: <https://clevertap.com/segmentation/>.
- [15] HubSpot's Product and Services Catalog, 2022. URL: https://legal.hubspot.com/hubspot-product-and-services-catalog?_ga=2.129780293.818963249.1623571090-2102544527.1617356624.
- [16] Experian. Marketing solutions, 2022. URL: <https://www.experian.com/business/solutions/marketing-solutions>.
- [17] SproutSocial. Listening Tools. Inform your business strategy with social listening, 2022. URL: <https://sproutsocial.com/features/social-media-listening/>.
- [18] Qualtrics. Platforma dlya segmentacii rynka. Izuchajte svoyu celevuyu auditoriyu s pomoshch'yu resheniya dlya segmentacii rynka [Qualtrics. Market segmentation platform. Research your target audience with a market segmentation solution], 2022. URL: <https://www.qualtrics.com/ru/product-experience/po-dlya-segmentirovaniya-rynka/?rid=langMatch&prevsite=en&newsite=ru&geo=UA&geomatch=>.

- [19] MailChimp. Put your audience at the heart of your marketing, 2022. URL: <https://mailchimp.com/audience/>.
- [20] Yieldify. Behavioral segmentation, 2022. URL: <https://www.yieldify.com/platform/behavioral-segmentation/>.
- [21] Amplitude Analytics, 2022. URL: <https://help.amplitude.com/hc/en-us/categories/360006505092-Amplitude-Analytics>.
- [22] Indicative. Segmentation, 2022. URL: <https://www.indicative.com/feature/segmentation/>.
- [23] Mixpanel. Limitless segmentation. Analyze why metrics change, 2022. URL: <https://mixpanel.com/segmentation/>.
- [24] D.O. Maidebura, K.V. Melnyk, N.V. Borysova, Analiz isnuichykh servisiv nalashtuvannia tarhetovanoi reklamy [Analysis of existing services of setting up targeted advertising], in: E. I. Sokol (Eds.), Proceedings of XXIX International scientific-practical conference in Information technologies: science, engineering, technology, education, health, part 1 of MicroCAD-2021, NTU “KhPI”, Kharkiv Ukraine, 2021. p. 32.
- [25] B. S. Everitt, S. Landau, D. Stahl, Cluster Analysis, Wiley, New York, NY, 2011.
- [26] G. Seif. The 5 Clustering Algorithms Data Scientists Need to Know, 2018. URL: <https://towardsdatascience.com/the-5-clustering-algorithms-data-scientists-need-to-know-a36d136ef68>.
- [27] M. McGregor. 8 Clustering Algorithms in Machine Learning that All Data Scientists Should Know, 2020. URL: <https://www.freecodecamp.org/news/8-clustering-algorithms-in-machine-learning-that-all-data-scientists-should-know/>.

Virtual Environment for Internet of Things Technologies Studying

Yevhenii Yaremchenko^a, Anzhelika Parkhomenko^a, Artem Tulenkov^a, Andriy Parkhomenko^a, Yaroslav Zalyubovskiy^a, Aleksandr Sokolyanskii^a and Olga Gladkova^a

^a National University Zaporizhzhia Polytechnic, 64, Zhukovskogo str., Zaporizhzhia, 69063, Ukraine

Abstract

For the development of Industry 4.0 educational course on Internet of Things we propose to implement the practical part of the course as a cycle of laboratory works with the usage of modern online engineering technologies, in particular developed online Virtual Learning Environment. The application of this environment will provide students with knowledge and skills necessary for selection of components and protocols for IoT systems realization as well as for organization of users' interaction with IoT applications. The scientific novelty of the work lies in fact that the method of organizing user interaction with the Virtual Learning Environment has been improved through the introduction of remote user interface and remote programming method to provide new ways of obtaining data from the user, increasing the convenience of their processing and visibility of presentation.

Keywords

Internet of Things, online engineering, cloud platform, virtual learning environment, remote user interface, MQTT protocol

1. Introduction

The introduction of the Internet of Things (IoT) technologies in industrial production allows enterprises and companies to reduce operating costs, increase production uptime and, accordingly, increase production volumes and reduce its costs. For example, constant monitoring of production processes can ensure timely maintenance of equipment and improve the level of product quality control.

However, despite the growing popularity, the IoT technologies still have certain drawbacks and unresolved issues that hinder their comprehensive integration into industrial production. They are the problems of introducing standardization, ensuring an appropriate level of security and confidentiality, reducing the cost of the IoT system implementation, combining it with cloud technologies, etc. [1]. Therefore, the further development of the IoT technologies is an urgent task for both scientists and industrialists.

In the context of distance learning, mastering the IoT technologies is especially problematic, due to the complex organization of remote practical experiments with components, protocols, hardware and software platforms used during the implementation of the IoT systems. Therefore, the development and application of online engineering technologies and tools for the IoT systems studying are relevant [2].

CMIS-2022: The Fifth International Workshop on Computer Modeling and Intelligent Systems, Zaporizhzhia, Ukraine, May 12, 2022
EMAIL: yaremchenkoyev@gmail.com (Y. Yaremchenko); parkhomenko.anzhelika@gmail.com (A. Parkhomenko); aetulenkov@gmail.com (A. Tulenkov); andriy2872073@gmail.com (A. Parkhomenko); zalubovskiy.yar@gmail.com (Y. Zalyubovskiy); jorapobeditel@gmail.com (A. Sokolyanskii); gladolechka@gmail.com (O. Gladkova)
ORCID: 0000-0001-9971-5599 (Y. Yaremchenko); 0000-0002-6008-1610 (A. Parkhomenko); 0000-0003-4863-4144 (A. Tulenkov); 0000-0001-8265-0530 (A. Parkhomenko); 0000-0002-6847-8778 (Y. Zalyubovskiy); 0000-0002-3623-9019 (A. 6); 0000-0002-6834-2854 (O. Gladkova)



© 2020 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).
CEUR Workshop Proceedings (CEUR-WS.org)

The goals of the work are research and practical implementation of online virtual environment for the IoT technologies studying.

2. Related works

Online engineering is a generic definition for a wide range of related technologies: remote engineering, virtual instrumentation, online modeling, cloud technologies, etc. [3].

The developed classification of online engineering technologies is presented in Fig. 1. As the studies have shown, these technologies can be successfully used for training in the field of the IoT by sharing engineering equipment and resources of remote laboratories (RL) [4], specialized software for virtual laboratories (VL) and virtual environments (VE) [5]-[6] as well as cloud platforms and services [7].

For example, IoTIFY [8] is a cloud platform designed to support the creation, validation and continuous monitoring of IoT systems. This platform allows to create realistic models of IoT devices based on various components and parameters that can be configured to simulate the interaction of devices with the environment and control this process dynamically through Application Programming Interface (API). IoTIFY supports the most popular data transfer protocols such as MQTT, HTTP, CoAP, AMQP, and LWM2M. With Raw UDP / DTLS and TCP / TLS support, it is also possible to simulate any other protocol [8]. IoTIFY offers two solutions: VL IoTIFY and IoT Simulator. VL IoTIFY is an online IoT learning platform that provides an interactive online workspace for creating electronic prototypes using a web browser. This laboratory allows to develop IoT applications in the browser based on the Arduino Uno, STM32 Nucleo F411RE and Raspberry Pi microcontroller platforms as well as various peripherals [9]. Using VL IoTIFY, students can design prototypes of IoT devices and learn embedded systems programming. VL IoTIFY also enables team collaboration on a project, as all components are available online via a web browser. VL IoTIFY is available under a university license for a minimum of 10 student licenses per year. For individual usage, first 2 credits are available (1 credit - 1 project for 1 hour of use), after that it is necessary to buy additional credits [8].

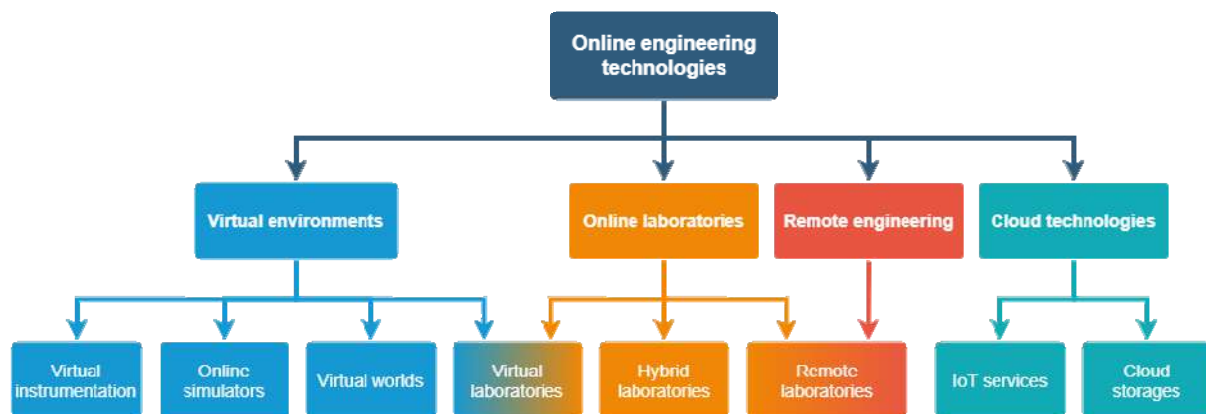


Figure 1: Classification of online engineering technologies

The IoT Simulator from IoTIFY solves the challenge of testing scalability and performance of the IoT system, uses virtualization and simulates thousands of IoT endpoints in the cloud. Thus, it is possible to completely simulate the system before the product reaches the customer [10].

Virtual IoT is a VE for collaborative learning in IoT technologies [11]. In order to start using the VE Virtual IoT, it is necessary to register on the project portal as a teacher. After completing the registration process, the teacher can add students to the class. Students need to download the VE to their computers and log in using the mailbox that the teacher has used to invite the student [12].

The training course consists of three levels with three tasks for each of them. The student's goal is to pass the level, getting as many points as possible for completing the tasks. The transition to the next level is possible only after completing all the tasks at the current level. The training material is presented in the form of movies, which can be viewed in the room located before the task room.

Currently, Virtual IoT is in the process of development, so, there are still certain problems with the displaying of training movies, chat and with the virtual surroundings [12].

Wokwi [13] is an online VL for simulating Arduino and other electronics. The goal of the project is to provide a convenient and powerful tool for learning programming based on the Arduino UNO, Mega, Nano and other platforms for creating prototypes of IoT devices. Wokwi allows to simulate the operation of various devices, in particular, LEDs, shift registers, sensors (ultrasonic, temperature, humidity, etc.), SSD display, a keyboard, etc.

An important idea of the project is to enable users to share their ideas with other developers, because the results of the discussions, suggestions and comments may help to improve the projects. In addition, VL Wokwi offers many ready-made examples of the implemented devices based on various hardware platforms and measuring instruments. The development of documentation for this VL is still in its early stages, so there is a lot of missing information. If necessary, it is possible to contact the project server for help via the Discord communication platform [14].

The work [15] describes the usage of 3D virtual world Second Life with elements of social network for the development of VE for teaching in the field of engineering. Second Life is a powerful system for VE development with the ability to render high-quality 3D graphics. It gives the possibility to create collaborative VE, where many users can interact with each other. There is a disadvantage of this system: it can be noted that it is impossible to create an HTML5-compatible program for usage in a regular browser, because a special browser must be used to access the developed virtual worlds.

Thus, as the studies have shown, teaching methods and learning resources based on online engineering technologies can be used successfully for training in the field of the IoT. Each of them has its own functionality, advantages and disadvantages. The choosing of appropriate learning resources is quite complicated and requires consideration of many criteria, so it is important to use recommendation systems, in particular for selection the appropriate RL and VL [16] as well as the IoT platforms [7].

Existing VLs offer experiments focused on specific programming problems for IoT devices based on different hardware platforms. These VLs are mostly presented as web clients, as such systems are fairly easy to implement for usage in a web browser. Existing VEs offer more possibilities to the user based on the free exploration of the virtual world and performing various tasks in it. Most of them are presented as desktop clients, so they require installation on users' computers.

Thus, the task of further development of online learning environment for organizing virtual experiments in the area of the IoT technologies is relevant.

3. Main and anticipated findings

3.1. Method of user interaction with online Virtual Learning Environment

The developed Virtual Learning Environment (VLE) for the IoT technologies studying can be implemented as a non-immersive online system in the form of a web client, i.e. HTML5-compatible program for using in a regular browser. This approach does not require the application of additional control devices, high computing power and installation of software on users' computers.

To organize the interaction of users with online VLE, it is advisable to implement GUI and WIMP paradigm methods: point and click, select and highlight [17] - [19]. This will reduce the number of steps required to take action and obtain information about the VLE.

The selection of items is carried out using the "point and click" method. The selected items are highlighted and icons and widgets with additional information are visualized. In addition, it is necessary to use specific for VLE interaction methods:

- Heads-up display (HUD), on which the actual data, error messages, information widgets and widgets for performing certain actions are visualized. HUD is used in all VLEs, because it is the most user-friendly ways of displaying important information to the user [20] - [21].
- Perspective of observing an agent is top-down. Although the usage of the first-person perspective makes the user more immersed in the VLE, it is more convenient to use the top-down perspective for performing experiments and complex research of VLE. In this case, the user can see what is happening in all or almost all areas of VLE and the agent surveillance camera is behind him [21] - [22].

- Computer mouse to organize agent movement inside VLE. The user is able to indicate the place to which the agent should reach using the "point and click" method. The agent is able to perform only simple movements in order to get to the target, in particular, he cannot jump. Such set of available movements was defined based on investigated works (for example, Virtual IoT). This approach to the organization of agent movement is since to move the game character the user needs to use only a computer mouse without the need to use the keyboard, which simplifies the process of interacting with VLE [23].

However, a set of basic methods for organizing user interaction with the VLE should be improved by introducing a remote user interface and a remote programming method (Fig. 2).

A remote user interface is an interface that allows a user to interact with a remote environment using a variety of communication methods, one of which is the remote programming method.

The remote programming method is a method of user interaction with the system, which consists in using a programming language to send commands to the system based on online technologies. The application of the remote programming method for interaction with the system allows the user to develop external programs that can control the full range of the IoT system functions.

A typical approach to this method implementation is the usage of TCP / IP sockets and any programming language with built-in capabilities or methods for working with them. In the system, that implements this method of interaction, the target device can be remotely controlled and it provides API for external applications or devices with which they can send commands and receive results.

Based on this method, various interfaces for the developed VLE can be supported using one API. For example, interactions between VLE and the user based on ready or user-created hardware or software with sensors and actuators, PCs, web clients, smartphones, virtual instruments, etc. [24].

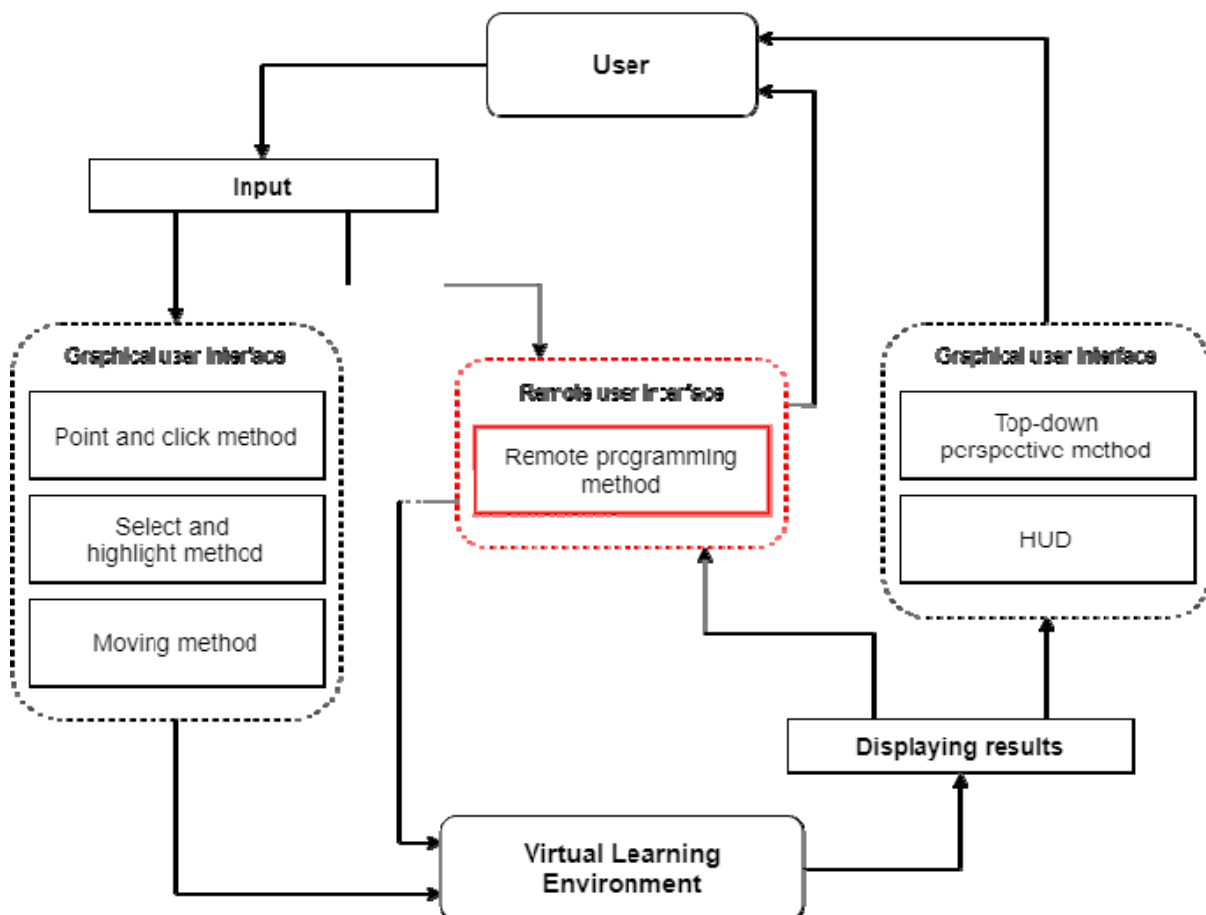


Figure 2: Improved method of organizing user interaction with VLE

3.2. Implementation of online Virtual Learning Environment

The developed architecture of VLE is presented in Fig. 3. To ensure the possibility of remote interaction with the developed system, the data transmission protocols used in the IoT systems were analysed. The investigations [25] - [26] have shown that protocols are an integral part of the technology stack used in such systems.

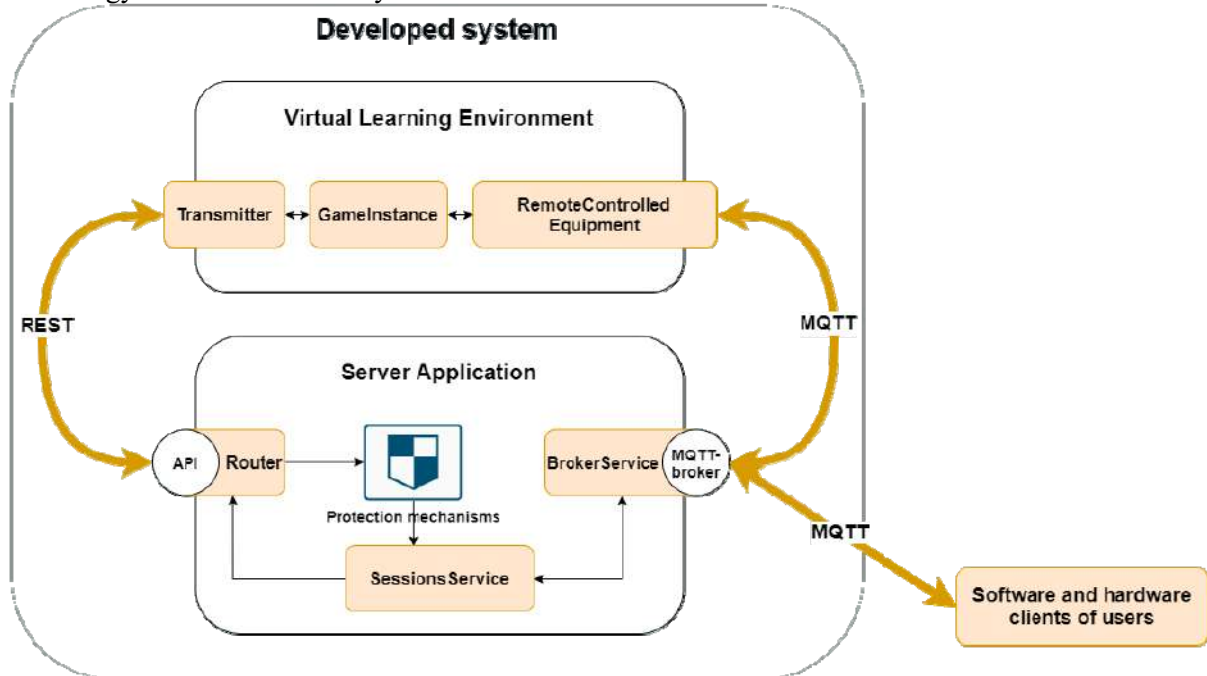


Figure 3: Architecture of developed system

The stack of protocols within the OSI model, which are most often considered in the context of IoT technologies, is the next:

- Application layer protocols (IETF CoAP, IBM MQTT, XMPP and AMQP) are used to connect Smart Things as well as HSC of the end users to the Internet.
- Transport layer protocols (UDP and TCP) are designed to transfer data between HSC.
- Network layer protocols (IPv6, ROLL RPL and 6LoWPAN) provide network routing and encapsulation capabilities.
- BLE, Z-Wave, ZigBee, HomePlug GP and Dash7 are five standard low-power channel layer protocols for short-distance wireless communication within a single LAN.
- At the physical level, the main protocols are IEEE 802.11 (Wi-Fi), IEEE 802.15.1 (Bluetooth) and IEEE 802.15.4 (LR-WPAN).

In the context of the development of VLE, which should provide the opportunity to gain practical skills in developing hardware and software for working with IoT systems, it was necessary to analyse application layer data transmission protocols, because users and HSC directly interact with them.

Comparisons of application layer data protocols used in IoT systems are presented in Table 1. The investigations [25] - [26] have shown that the MQTT data transfer protocol can provide reliable and secure routing for small and cheap devices with low power consumption and memory in unstable and low bandwidth networks.

Therefore, the MQTT application-level data transfer protocol can be recommended for organizing remote interaction between VLE and users' HSC, because this popular and reliable protocol is often used in the IoT systems.

The system is built on the basis of a client-server architecture that combines a server and VLE. The developed server class diagram is presented in Fig. 4. The Node.js platform, the Express.js framework, the JavaScript and TypeScript programming languages and the Visual Studio Code IDE were chosen as the system server development toolkit. To implement online VLE the game engine Unreal Engine 4 was chosen, which provides a high level of graphic capabilities, ready-made

templates and models, textures and materials for them. It is important that with the help of this game engine it is possible to create applications for web browsers that is HTML5-compatible programs [27].

For the server, the structure of entry points for communication between the server and VLE was developed. The naming rules and the structure of MQTT equipment topics as well as the rules for publishing in them were determined for VLE.

Table 1

Comparison of application layer data protocols used in IoT systems

Criterion \ Protocol	CoAP	MQTT	XMPP	AMQP
Goal	Reduce the overhead of messages, their size and the need for their fragmentation	Ensure reliable communication in networks with limited bandwidth or unstable communication	Provide users with the ability to connect to multiple instant messaging systems	To transfer business messages between HSC
Messaging mode	Synchronous and asynchronous	Asynchronous	Synchronous and asynchronous	Asynchronous
Works in real time	+	+	+	-
Reliability	Built-in mechanism for ensuring the reliability of messaging	Reliability of message delivery through different levels of QoS	TCP guarantees reliability	The intermediate storage function ensures reliability
Security	DTLS	SSL / TLS	SSL / TLS	SSL / TLS
Quality of Service (QoS) option support	+	+	-	+
Architecture	Request / response	Publication / subscription	Publication / subscription, request / response	Publication / subscription
Transport layer protocol	UDP	TCP	TCP	TCP
Energy consumption	Reduces energy consumption	Reduces energy consumption	Reduces energy consumption	Consumption is higher than in MQTT
Network bandwidth	Increases network bandwidth	Optimizes network bandwidth	Significantly reduces network bandwidth	Reduces network bandwidth

Implemented VLE is the non-immersion system in the form of a web client to run in Google Chrome and Mozilla Firefox web browsers. The developed software provides the functions of simulating VLE and its equipment (using the example of home automation system) as well as the possibilities for user interaction with them. This VLE initiates actions on the session by sending requests about this to the server using its REST API. It also displays the parameters for connecting the users' HSC, the message exchange log and other information to the user (Fig. 5). The main areas of HUD are:

- A - Log of events;
- B - Session timer;
- C - Data for connecting users' software and hardware clients;
- D - Log of MQTT messages published by the equipment.

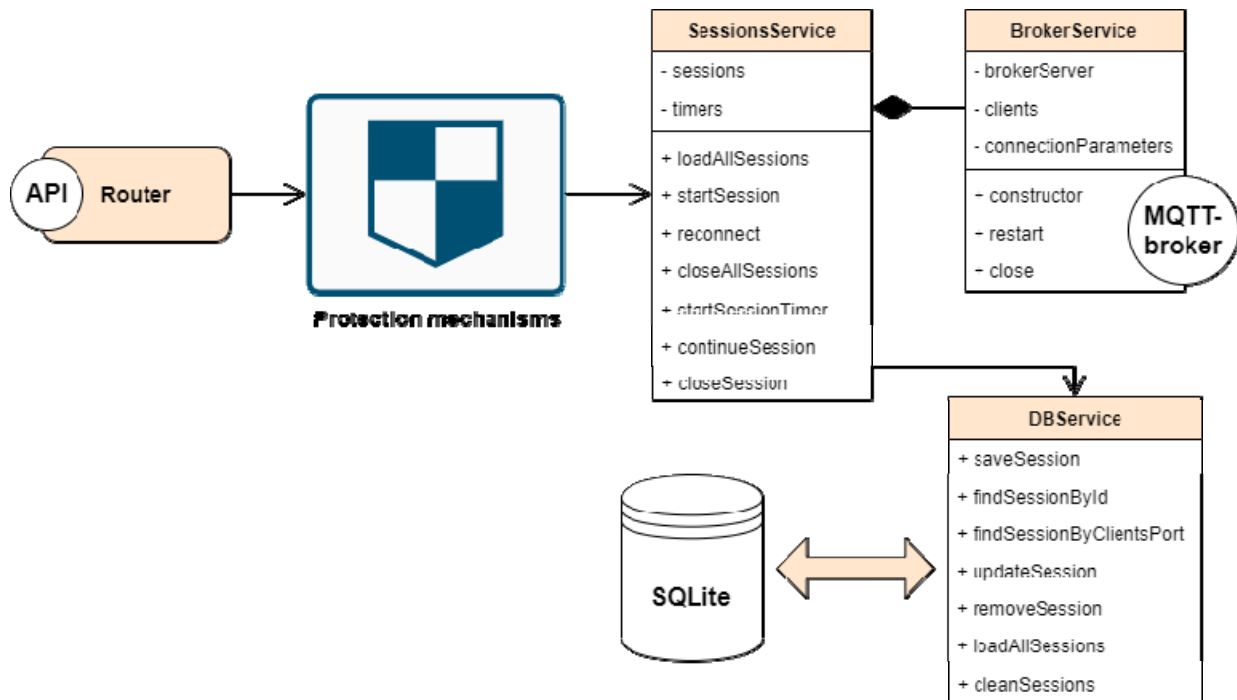


Figure 4: Server class diagram

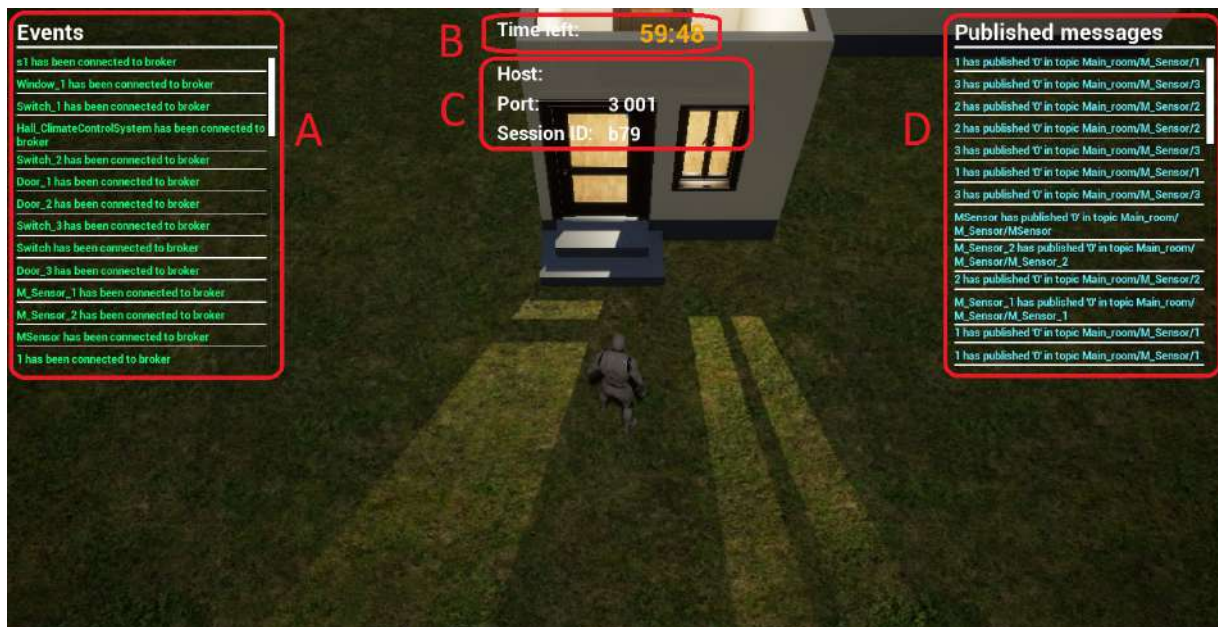


Figure 5: HUD of VLE

3.3. User interaction with the developed Virtual Learning Environment

In order to be able to control the events that take place in the home automation system reproduced in VLE as well as to create them, the structure of MQTT-topics was implemented (Table 2).

The server adheres to the following rules for naming and publishing MQTT-topics of equipment:

- MQTT case-sensitive topics;
- MQTT-topics are structured in a hierarchy similar to folders and files in the file system using a slash (/) as a delimiter;
- only UTF-8 is used in topic names;
- \$SYS is a reserved topic and is used to publish broker information;
- # (hash symbol) - multilevel substitution symbol;
- the client can publish only in a separate topic. That is, the use of wildcards during publication is prohibited. For example, to publish messages in two topics, the client needs to publish messages separately in each of them.

Table 2
The structure of MQTT-topics of VLE equipment

Topic	Structure
'Hello'-topic	session_id/hello
\$SYS-topic	session_id/\$SYS
Equipment	session_id/room_id/equipment_type/equipment_id/...
Subscription on all non-\$SYS topics	session_id/#
Subscription on all equipment topics in specific room	session_id/room_id/#
Subscription on all topics of specific equipment type in specific room	session_id/room_id/equipment_type/#

UML diagram of the algorithm of user interaction with the developed system is presented in Fig. 6.

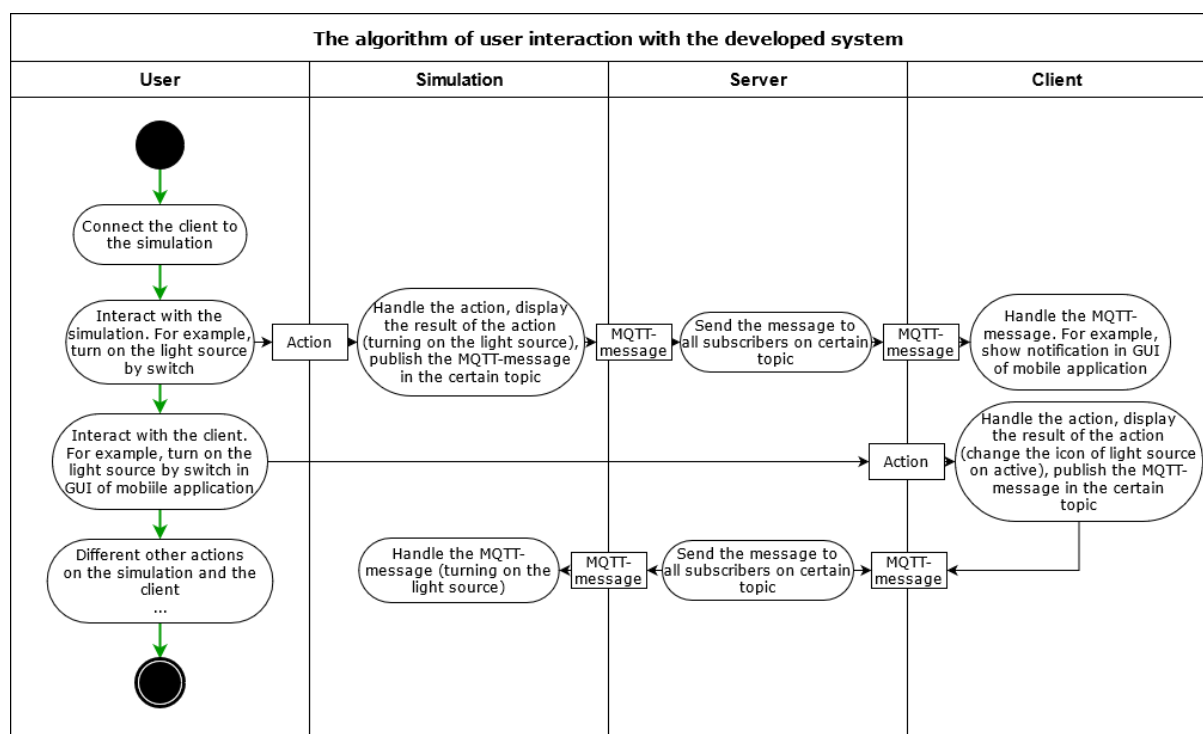


Figure 6: The algorithm of user interaction with the developed system

After connection of HSC to VLE user can interact with both and the result of such interaction affects on another part. For example, user turns on the light source by switch in VLE. VLE handles this action and, in this example, display the turning on of the desired light source as well as publish the MQTT-message about this change in the certain topic. The server sends the received message to

all subscribers on the certain topic. HSC should handle the received message. For example, show notification in GUI of mobile application.

In case of user interaction with HSC, which is developed by himself, HSC should handle the action. For example, user turns on the light source by switch in GUI of mobile application. This application should display the result of the action (in this case, change the icon of the light source on active) as well as publish the MQTT-message about this change in the certain topic. The server sends the received message to all subscribers on the certain topic. VLE handles this action and, in this example, display the turning on of the desired light source.

4. Discussion

Analysis of the possibilities of using the product showed that the user can interact with the developed VLE using hardware as well as self-developed or ready-made software applications. To ensure compatibility between clients and the server of the system, students' clients must be able to work with TCP/IP sockets using the MQTT data transfer protocol.

A software client is software that allows the user to perform actions in a VLE by passing commands to it. Due to the active development of smartphones and other mobile devices, the most common type of software clients are mobile applications. They have opportunities for informative visualization of the VLE state as well as for active interaction with it.

A hardware client is a device that provides various physical tools to interact with VLE. Such tools include simple buttons as well as various sensors (lighting, temperature, etc.), based on which the software component of such clients can interact with VLE. For example, a hardware client can implement behaviour, when based on the measurement of the surrounding room temperature in case of exceeding the normal temperature, will transmit a VLE command to turn on the climate control system.

In addition to developing their own applications, the user can use existing solutions, such as a free MQTT client, which is presented in Fig. 7 [28].



Figure 7: Free MQTT client to work with the developed VLE

The experimental representation of the system data is presented in Table 3 [29]. For today, the number of implemented standard Smart Houses subsystems according to the Table 3 is limited. This is due to the complexity and duration of the development of virtual models of equipment, which often requires the implementation of a mathematical model of physical processes (for example, a

mathematical model of room temperature measurement result changing by the influence of the climate control system taking into account room configuration).

5. Conclusion

As a result of the performed research, the features of existing methods and tools of online engineering for studying IoT technologies were identified.

The developed online VLE is intended for teaching and learning the development of software and hardware for interacting with the IoT systems, conducting remote experiments on developing clients for working with the IoT systems as well as creating and testing the operability and efficiency of control scenarios.

Table 3

Experimental representation of the system data

Topic	Payload	VLE action / reaction
f84d592/Room1/Door/Door1	opened /	Open / close the door Door1 / the window
f84d592/Room1/Window/Window1	closed	Window1 in the room Room1
f84d592/Room1/LightSource/LS1	turn_on /	Turn on / off the light source LS1 in the
	turn_off	room Room1
f84d592/Room1/LightSource/LS1/color	58,49,91	Set the color of the light source LS1 to
		yellow (HSV color format)
f84d592/Room1/CCS/CCS1	turn_on /	Turn on / off the climate control system
	turn_off	CCS1 in the room Room1
f84d592/Room1/CCS/CCS1/humidity	60	Set the humidity for the climate control
		system CCS1 to 60%
f84d592/Room1/CCS/CCS1/temperature	20	Set the temperature for the climate
		control system CCS1 to 20°C
f84d592/Room1/L_Sensor/LSens	value	Send to the broker information about
f84d592/Room1/M_Sensor/MSens		measurement results of the specific sensor
f84d592/Room1/T_Sensor/TSens		
f84d592/Room1/H_Sensor/HSens		

The practical value of the proposed solutions is determined by the fact that implemented VLE which is based on improved method, provides a tool for developing practical skills for programmable interaction with the IoT systems using hardware and software clients (HSC). User develops these clients independently and he can connect them to VLE for its control via API. Such system offers a high level of interactivity, because users' actions in VLE affect the state of the HSC developed by him and vice versa.

In the future, it is planned to increase the number of available components of VLE, in particular for industrial IoT systems as well as to implement new possibilities for user interaction with VLE.

6. Acknowledgements

This work is partly carried out with the support of Erasmus + KA2 project WORK4CE "Cross-domain competences for healthy and safe work in the 21st century" (619034-EPP-1-2020-1-UA-EPPKA2-CBHE-JP).

7. References

- [1] What are the main advantages and disadvantages of the Internet of Things? URL: <https://www.imd.org/iot/iot-reflections/pros-and-cons-of-iot>.

- [2] D. Pirrone, C. Fornaro, D. Assante, Open-source multi-purpose remote laboratory for IoT education, in: Proceedings of the 2021 IEEE Global engineering education conference, EDUCON, IEEE, Vienna, Austria, 2021, pp. 1462-1468. doi:10.1109/EDUCON46332.2021.9454034.
- [3] International journal of online and biomedical engineering. About the journal. URL: <https://online-journals.org/index.php/i-joe/about>.
- [4] A. Parkhomenko, A. Tulenkov, A. Sokolyanskii, Y. Zalyubovskiy, A. Parkhomenko, Integrated complex for IoT technologies study, in: M. Auer, D. Zutin (Eds.), Online engineering & Internet of Things, volume 22(31) of Lecture notes in networks and systems, Springer, Cham, 2017, pp. 322–330. doi:10.1007/978-3-319-64352-6_31.
- [5] B. S. Wästberg, T. Eriksson, G. Karlsson, M. Sunnerstam, M. Axelsson, M. Billger, Design considerations for virtual laboratories: A comparative study of two virtual laboratories for learning about gas solubility and colour appearance, education and information technologies 24 (2019) 2059–2080. doi:10.1007/s10639-018-09857-0.
- [6] Y. Francillette, E. Boucher, A. Bouzouane, S. Gaboury, The virtual environment for rapid prototyping of the intelligent environment, Sensors (Basel) 17.11 (2017). doi:10.3390/s17112562.
- [7] A. Tulenkov, A. Parkhomenko, A. Sokolyanskii, Evaluation and selection of IoT service for Smart House system big data processing, in: Proceedings of the IEEE 14th International scientific and technical conference on Computer sciences and information technologies, CSIT, IEEE, Lviv, Ukraine, 2019, pp. 124-129. doi:10.1109/STC-CSIT.2019.8929810.
- [8] IoTIFY. URL: <https://iotify.io/>.
- [9] Virtual lab. URL: <https://docs.vlab.iotify.io/>.
- [10] IoTIFY. IoT simulator. URL: <https://iotify.io/iot-network-simulator/>.
- [11] Virtual IOT. URL: <http://www.virtualiot.net/>.
- [12] Virtual IOT. How to use. URL: <http://www.virtualiot.net/how-to-use/>.
- [13] Wokwi. URL: <https://wokwi.com/>.
- [14] Wokwi docs. URL: <https://docs.wokwi.com/>.
- [15] M. J. Callaghan, K. McCusker, J. Lopez Losada, J.G. Harkin, S. Wilson, Engineering education island: teaching engineering in virtual worlds, Innovation in teaching and learning in Information and computer sciences 8.3 (2009) 2–18. doi:10.11120/ital.2009.08030002.
- [16] A. Parkhomenko, O. Gladkova, A. Parkhomenko, Recommendation system as a user-oriented service for the remote and virtual labs selecting, in: M. Auer, T. Tsiatsos (Eds.), The challenges of the digital transformation in education, volume 917 of Advances in intelligent systems and computing, Springer, Cham, 2019, pp. 600-610. doi:10.1007/978-3-030-11935-5_57.
- [17] K. Hinckley, D. Wigdor, Input Technologies and Techniques, in: J. A. Jacko (Ed.), The human-computer interaction handbook, 3rd. ed., CRC Press, Boca Raton, FL, 2012, pp. 95–132. doi:10.1201/b10368-12.
- [18] K. M. Inkpen, Drag-and-drop versus point-and-click mouse interaction styles for children, ACM Transactions on Computer-Human Interaction 8.1 (2001) 1–33. doi:10.1145/371127.371146
- [19] Selection. URL: <https://material.io/design/interaction/selection.html>.
- [20] L. Caroux, K. Isbister, Influence of head-up displays' characteristics on user experience in video games, International journal of human-computer studies 87 (2016) 65–79. doi: 10.1016/j.ijhcs.2015.11.001.
- [21] L. Caroux, K. Isbister, L. Le Bigot, N. Vibert, Player–video game interaction: A systematic review of current concepts, Computers in human behavior 48 (2015) 366–381. doi: 10.1016/j.chb.2015.01.066.
- [22] J. Garcia, Perspectives and points of view, 2020. URL: https://www.ign.com/wikis/video-game-dictionary/Perspectives_and_Points_of_View.
- [23] C. Compton, Run, jump and climb: designing fun movement in games, 2019. URL: <https://www.gamedeveloper.com/audio/run-jump-and-climb-designing-fun-movement-in-games>.
- [24] R. M. Prades, P. J. Sanz, P. Nebot, R. Wirz, A multimodal interface to control a robot arm via the web: a case study on remote programming, IEEE Transactions on industrial electronics 52.6 (2006) 1506–1520. doi:10.1109/TIE.2005.858733.

- [25] C. Sharma, N. K. Gondhi, Communication protocol stack for constrained IoT systems, in: Proceedings of the 3rd. International conference on Internet of Things: smart innovation and usages, IoT-SIU, IEEE, Bhimtal, India, 2018, pp. 1–6. doi:10.1109/IoT-SIU.2018.8519904.
- [26] T. M. Tukade, R. Banakar, Data transfer protocols in IoT an overview, International journal of pure and applied mathematics 118 (2018) 121–138.
- [27] P. Gestwicki, Unreal Engine 4 for Computer Scientists, Journal of computing sciences in colleges 35.5 (2019) 109–110.
- [28] BP_Mqtt. URL: https://github.com/damody/BP_MQTT_Demo/releases/download/1.6/BP_Mqtt.zip.
- [29] Y. Yaremchenko, J. Nau, D. Streitferdt, K. Henke, A. Parkhomenko, Virtual environment Smart House for hybrid laboratory GOLD, in: M. E. Auer, H. Hortsch, O. Michler, T. Köhler (Eds.), Mobility for Smart cities and regional development - challenges for higher education, volume 389 of Lecture Notes in Networks and Systems, Springer, Cham, 2022, pp. 250–257. https://doi.org/10.1007/978-3-030-93904-5_25