

Lecture Notes in Networks and Systems 2048

Shahnaz N. Shahbazova
Vladik Kreinovich
Janusz Kacprzyk
Bernard De Baets *Editors*

Soft Computing Applications

Selected Papers of the 9th World
Conference on Soft Computing,
September 25–27, 2024, Baku,
Azerbaijan. Volume 1

 Springer

Lecture Notes in Networks and Systems

Volume 2048

Series Editor

Janusz Kacprzyk , Systems Research Institute, Polish Academy of Sciences,
Warsaw, Poland

Advisory Editors

Fernando Gomide, Department of Computer Engineering and Automation—DCA,
School of Electrical and Computer Engineering—FEEC, University of
Campinas—UNICAMP, São Paulo, Brazil

Okyay Kaynak, Department of Electrical and Electronics Engineering,
Bogazici University, Istanbul, Türkiye

Derong Liu, Department of Electrical and Computer Engineering, University of
Illinois at Chicago, Chicago, USA

Institute of Automation, Chinese Academy of Sciences, Beijing, China

Witold Pedrycz, Department of Electrical and Computer Engineering, University of
Alberta, Edmonton, Alberta, Canada

Systems Research Institute, Polish Academy of Sciences, Warsaw, Poland

Marios M. Polycarpou, Department of Electrical and Computer Engineering,
KIOS Research Center for Intelligent Systems and Networks, University of Cyprus,
Nicosia, Cyprus

Imre J. Rudas, Óbuda University, Budapest, Hungary

Jun Wang, Department of Computer Science, City University of Hong Kong,
Kowloon, Hong Kong

The series “Lecture Notes in Networks and Systems” publishes the latest developments in Networks and Systems—quickly, informally and with high quality. Original research reported in proceedings and post-proceedings represents the core of LNNS.

Volumes published in LNNS embrace all aspects and subfields of, as well as new challenges in, Networks and Systems.

The series contains proceedings and edited volumes in systems and networks, spanning the areas of Cyber-Physical Systems, Autonomous Systems, Sensor Networks, Control Systems, Energy Systems, Automotive Systems, Biological Systems, Vehicular Networking and Connected Vehicles, Aerospace Systems, Automation, Manufacturing, Smart Grids, Nonlinear Systems, Power Systems, Robotics, Social Systems, Economic Systems and other. Of particular value to both the contributors and the readership are the short publication timeframe and the world-wide distribution and exposure which enable both a wide and rapid dissemination of research output.

The series covers the theory, applications, and perspectives on the state of the art and future developments relevant to systems and networks, decision making, control, complex processes and related areas, as embedded in the fields of interdisciplinary and applied sciences, engineering, computer science, physics, economics, social, and life sciences, as well as the paradigms and methodologies behind them.

Indexed by SCOPUS, EI Compendex, INSPEC, WTI Frankfurt eG, zbMATH, SCImago.

All books published in the series are submitted for consideration in Web of Science.

For proposals from Asia please contact Aninda Bose (aninda.bose@springer.com).

Shahnaz N. Shahbazova · Vladik Kreinovich ·
Janusz Kacprzyk · Bernard De Baets
Editors

Soft Computing Applications

Selected Papers of the 9th World Conference
on Soft Computing, September 25–27, 2024,
Baku, Azerbaijan. Volume 1

Editors

Shahnaz N. Shahbazova
Department of Digital Technologies
and Applied Informatics,
Azerbaijan State University of Economics
(UNEC)
Baku, Azerbaijan

Janusz Kacprzyk
Systems Research Institute
Polish Academy of Sciences
Warsaw, Poland

Vladik Kreinovich
Department of Computer Science
University of Texas at El Paso
El Paso, TX, USA

Bernard De Baets
Faculty of Bioscience Engineering
Ghent University
Ghent, Oost-Vlaanderen, Belgium

ISSN 2367-3370

ISSN 2367-3389 (electronic)

Lecture Notes in Networks and Systems

ISBN 978-3-032-04802-8

ISBN 978-3-032-04803-5 (eBook)

<https://doi.org/10.1007/978-3-032-04803-5>

© The Editor(s) (if applicable) and The Author(s), under exclusive license to Springer Nature Switzerland AG 2026

This work is subject to copyright. All rights are solely and exclusively licensed by the Publisher, whether the whole or part of the material is concerned, specifically the rights of translation, reprinting, reuse of illustrations, recitation, broadcasting, reproduction on microfilms or in any other physical way, and transmission or information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed.

The use of general descriptive names, registered names, trademarks, service marks, etc. in this publication does not imply, even in the absence of a specific statement, that such names are exempt from the relevant protective laws and regulations and therefore free for general use.

The publisher, the authors and the editors are safe to assume that the advice and information in this book are believed to be true and accurate at the date of publication. Neither the publisher nor the authors or the editors give a warranty, expressed or implied, with respect to the material contained herein or for any errors or omissions that may have been made. The publisher remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

This Springer imprint is published by the registered company Springer Nature Switzerland AG
The registered company address is: Gewerbestrasse 11, 6330 Cham, Switzerland

If disposing of this product, please recycle the paper.

*The 9th World Conference on Soft Computing
is dedicated to the Research Heritage of
Prof. Lotfi A. Zadeh.*

Advisory Editors

Vladik Kreinovich, University of Texas at El Paso, El Paso, USA

Ronald R. Yager, New York, USA

Oscar Castillo, Ph.D., D.Sc., Tijuana, Institute of Technology, Tijuana, Mexico

Laszlo T. Koczy, Szechenyi Istvan University, Gyor, Hungary

Jerry Mendel, Los Angeles, USA

Valentina E. Balas, Aurel Vlaicu University of Arad, Romania

Imre Rudas, Obuda University, Hungary

Yingxu Wang, University of Calgary, Canada

Fernando Gomide, University of Campinas, Brazil

Bernard De Baets, Ghent University, Belgium

Shahnaz N. Shahbazova, Azerbaijan State University of Economics, Baku,
Azerbaijan

Preface

These two volumes constitute the selected papers from of the 9th World Conference on Soft Computing or WConSC 2024, held on September 25–27, 2024, in Baku, Azerbaijan. The conference was organized by the University of California, Berkeley, California, USA, and Azerbaijan State Economic University, Azerbaijan.

In 1991, Zadeh introduced the concept of Soft Computing—a consortium of methodologies that collectively provide a foundation for the conception, design, and utilization of intelligent systems. One of the principal components of Soft Computing is Fuzzy Logic. Today, the concept of Soft Computing is growing rapidly in visibility and importance. Many papers written in the eighties and early nineties were concerned, for the most part, with applications of fuzzy logic to knowledge representation and commonsense reasoning. Soft Computing facilitates the use of fuzzy logic, neurocomputing, evolutionary computing, and probabilistic computing in combination, leading to the concept of hybrid intelligent systems.

Combining these intelligent systems tools has already led to a large number of applications, which shows the great potential of Soft Computing in many application domains.

The volumes cover a broad spectrum of Soft Computing techniques, including theoretical results and practical applications that help find solutions for industrial, economic, medical, and other problems.

The conference papers included in these proceedings were grouped into the following research areas:

- Fuzzy Simulation and Data mining.
- Intelligent and Multi-agent Systems.
- Fuzzy Recognition.
- Artificial Intelligence.
- Fuzzy sets and Quantum Computing.
- Fuzzy Control Systems.
- Fuzzy Methods and Modern Approaches.
- Fuzzy control and Fuzzy Approach.
- Fuzzy sets and Decision on making.

- Soft Computing and Fuzzy Logic.
- Fuzzy Applications and Fuzzy Control.
- Fuzzy Logic and Its Applications.
- Fuzzy Neural Networks.
- Intelligent Methods and Fuzzy Approach.
- Fuzzy Methods and Applications in Medicine.
- Fuzzy Applications and Analysis System.

During the WConSC 2024 conference, we had two panels. The first panel titled Lotfi Zadeh legacy: Fuzzy Logic, Fuzzy sets and decision-making featured presentations by Prof. Janusz Kacprzyk, Prof. Bernard De Baets, Prof. Nadejda G. Yarushkina, Prof. Annamaria R. Varkonyi-Koczy, Prof. Leonid Lyubchik and Prof. Shahnaz N. Shahbazova. The second panel titled Fuzzy Logic and Soft Computing, Computational Intelligence and Intelligent Technologies featured Prof. Vladik Kreinovich, Prof. Janusz Kacprzyk, Prof. Bernard De Baets, Leonid Lyubchik, Koksai Erenturk and Prof. Shahnaz N. Shahbazova. Many thanks to all the panel speakers for their very valuable memories of Prof. Lotfi A. Zadeh and for their interesting talks.

At the WConSC 2024 conference, we had ten keynote speakers: Prof. Janusz Kacprzyk (Poland), Prof. Vladik Kreinovich (Texas, USA), Prof. Bernard De Baets (Belgium), Prof. Nadezhda Yarushkina (Russia), Prof. Valentina E. Balas (Romania), Prof. Marius Balas (Romania), Prof. Koksai Erenturk (Erzerum, Turkiye), Prof. Leonid Lyubchik (Ukraine), Prof. Annamaria R. Varkonyi-Koczy (Hungary), Fernando Gomide (Brazil) and Prof. Shahnaz N. Shahbazova (Azerbaijan). Summaries of their talks are included in this book.

The editors of this book would like to acknowledge all the authors for their contributions that kept the quality of the WConSC 2024 conference at a high level. We have received invaluable help from the members of the International Program Committee and from the Program Committee of Conference. We would like to specially thank Prof. Vladik Kreinovich (Texas, USA), for his great organization and support from the review of all articles of the accepted conference.

Thanks to all of them we had remarkably interesting presentations and stimulating discussions.

Our special thanks go to Janusz Kacprzyk, Editor-in-Chief of the Springer book series *Advances in Intelligent Systems and Computing* for the opportunity to organize this guest-edited volume.

We are grateful to the Springer staff, especially to Dr. Thomas Ditzinger (Senior Editor, Applied Sciences and Engineering, Springer-Verlag), for their excellent collaboration, patience, and help during the preparation of these volumes.

We hope that the volumes will provide useful information to professors, researchers, and graduated students in the area of Soft Computing techniques and

applications and that all of them will find this collection of papers inspiring, informative and useful. We also hope to see you all at future World Conferences on Soft Computing.

Baku, Azerbaijan
El Paso, USA
Warsaw, Poland
Ghent, Belgium

Prof. Shahnaz N. Shahbazova
Prof. Vladik Kreinovich
Prof. Janusz Kacprzyk
Prof. Bernard De Baets

About Professor Lotfi A. Zadeh



Lotfi A. Zadeh
Department of Electrical Engineering and Computer Sciences
College of Engineering
University of California
Berkeley, USA

Lotfi A. Zadeh joined the Department of Electrical Engineering at the University of California, Berkeley, in 1959, and served as its chair from 1963 to 1968. Earlier, he was a member of the Electrical Engineering faculty at Columbia University. In 1956, he was a visiting member of the Institute for Advanced Study in Princeton, New Jersey. He held a number of visiting appointments, including a visiting professorship in Electrical Engineering at MIT in 1962 and 1968; visiting scientist at IBM Research Laboratory, San Jose, CA, in 1968, 1973, and 1977; visiting scholar appointments at the AI Center, SRI International, in 1981, and at the Center for the Study of Language and Information, Stanford University, in 1987–1988. He was a Professor in the UC Berkeley Graduate School and served as the Director of BISC (Berkeley Initiative in Soft Computing).

Until 1965, Dr. Zadeh's work centered on system theory and decision analysis. His 1965 paper on fuzzy sets has received over 26,000 Google Scholar citations and is by far the highest cited paper in Information and Control. In 1965, his research interests shifted to the theory of fuzzy sets and its applications to artificial intelligence, linguistics, logic, decision analysis, control theory, expert systems and neural networks. His later research focused on fuzzy logic, soft computing, computing with words, and the newly developed computational theory of perceptions and precisiated natural language.

An alumnus of the University of Tehran, MIT, and Columbia University, Dr. Zadeh was a fellow of the IEEE, AAAS, ACM, AAAI, and IFSA and a member of the National Academy of Engineering. He held NSF Senior Postdoctoral Fellowships in 1956–57 and 1962–63 and was a Guggenheim Foundation Fellow in 1968.

He received the IEEE Education Medal in 1973 and the IEEE Centennial Medal in 1984. In 1989, Dr. Zadeh was awarded the Honda Prize by the Honda Foundation and in 1991 received the Berkeley Citation, University of California.

Dr. Zadeh was accorded the following honors:

1992

- IEEE Richard W. Hamming Medal “For seminal contributions to information science and systems, including the conceptualization of fuzzy sets.”
- Foreign Member of the Russian Academy of Natural Sciences (Computer Sciences and Cybernetics Section).
- Certificate of Commendation for AI Special Contributions Award from the International Foundation for Artificial Intelligence.
- Kampe de Fériet Prize.
- Honorary Member of the Austrian Society of Cybernetic Studies.

1993

- Rufus Oldenburger Medal from the American Society of Mechanical Engineers “For seminal contributions in system theory, decision analysis, and theory of fuzzy sets and its applications to AI, linguistics, logic, expert systems, and neural networks.”
- Grigore Moisil Prize for Fundamental Researches.
- Premier Best Paper Award by the Second International Conference on Fuzzy Theory and Technology.

1995

- IEEE Medal of Honor “For pioneering development of fuzzy logic and its many diverse applications.”

1996

- Okawa Prize “For outstanding contribution to information science through the development of fuzzy logic and its applications.”

1997

- B. Bolzano Medal by the Academy of Sciences of the Czech Republic “For outstanding achievements in fuzzy mathematics.”
- J. P. Wohl Career Achievement Award of the IEEE Systems, Science and Cybernetics Society.
- Served as a Lee Kuan Yew Distinguished Visitor, lecturing at the National University of Singapore and the Nanyang Technological University in Singapore, and as the Gulbenkian Foundation Visiting Professor at the New University of Lisbon in Portugal.

1998

- Edward Feigenbaum Medal by the International Society for Intelligent Systems.
- Richard E. Bellman Control Heritage Award by the American Council on Automatic Control. The Information Science Award from the Association for Intelligent Machinery.
- SOFT Scientific Contribution Memorial Award from the Society for Fuzzy Theory in Japan.

1999

- Elected to membership in Berkeley Fellows.
- Certificate of Merit from IFSA (International Fuzzy Systems Association).

2000

- IEEE Millennium Medal.
- IEEE Pioneer Award in Fuzzy Systems.
- ASPIH 2000 Lifetime Distinguished Achievement Award.
- ACIDCA 2000 Award for the paper, “From Computing with Numbers to Computing with Words—From Manipulation of Measurements to Manipulation of Perceptions.”
- Chaos Award from the Center of Hyper incursion and Anticipation in Ordered Systems for his outstanding scientific work on foundations of fuzzy logic, soft computing, computing with words and the computational theory of perceptions.

2001

- ACM 2000 Allen Newell Award for seminal contributions to AI through his development of fuzzy logic.
- A Special Award from the Committee for Automation and Robotics of the Polish Academy of Sciences for his significant contributions to systems and information science, development of fuzzy sets theory, fuzzy logic control, possibility theory, soft computing, computing with words, and computational theory of perceptions.

2003

- Elected as a foreign member of the Finnish Academy of Sciences.
- Norbert Wiener Award of the IEEE Society of Systems, Man and Cybernetics “For pioneering contributions to the development of system theory, fuzzy logic and soft computing.”

2004

- Civitate Honoris Causa by Budapest Tech (BT) Polytechnical Institution, Budapest, Hungary.
- V. Kaufmann Prize by the International Association for Fuzzy-Set Management and Economy (SIGEF).

2005

- Elected as a foreign member of the Polish Academy of Sciences, Korea Academy of Science and Technology and Bulgarian Academy of Sciences.
- Nicolaus Copernicus Medal of the Polish Academy of Sciences.
- J. Keith Brimacombe IPMM Award.

2006

- Elected as a foreign member of the National Academy of Sciences of Azerbaijan.
- Pioneer Award for Outstanding Contributions to Soft Computing, Georgia State University, Atlanta, Georgia.
- Silicon Valley Engineering Hall of Fame.

2007

- Egleston Medal, Columbia University, New York.
- Member of the International Academy of Systems Studies (IASS).

2009

- Franklin Institute Medal, Philadelphia.

2011

- Medal of the Foundation by the Trust of the Foundation for the Advancement of Soft Computing, Spain.
- High State Award 'Friendship Order', from the President of the Republic of Azerbaijan.
- Trans disciplinary Award and Medal of the Society for Design and Process Sciences, Korea.

Dr. Zadeh was a recipient of twenty-four honorary doctorates from the following Universities:

- Paul-Sabatier University, Toulouse, France.
- State University of New York, Binghamton, USA.

- University of Dortmund, Dortmund, Germany.
- University of Oviedo, Oviedo, Spain.
- University of Granada, Granada, Spain.
- Lakehead University, Canada.
- University of Louisville, KY, USA.
- State Oil Academy of Azerbaijan.
- Baku State University, Azerbaijan.
- Silesian Technical University, Gliwice, Poland.
- University of Toronto, Toronto, Canada.
- University of Ostrava, the Czech Republic.
- University of Central Florida, Orlando, FL, USA.
- University of Hamburg, Hamburg, Germany.
- University of Paris(6), Paris, France.
- Johannes Kepler University, Linz, Austria.
- University of Waterloo, Canada.
- University of “Aurel Vlaicu”, Arad, Romania.
- Lappeenranta University of Technology, Finland.
- Muroran Institute of Technology, Japan.
- Hong Kong Baptist University, China.
- Indian Statistical Institute, Kolkata, India.
- University of Saskatchewan, Canada.
- Polytechnic University of Madrid, Spain.
- Ryerson University, Canada.

Keynote Speakers

Lotfi A. Zadeh's Biography and Scientific Legacy

Professor Dr. Shahnaz N. Shahbazova

Department of Digital Technologies and Applied Informatics
Azerbaijan State University of Economics (UNEC)

Baku, Azerbaijan

shahbazova@gmail.com

shahbazova.shahnaz@unec.edu.az

Abstract This paper describes Prof. Lotfi A. Zadeh's life and his scientific legacy. The work begins with his birth where he was born in Baku in Azerbaijan, as well as the period of residence in Baku from birth to 10 years. After 10 years, he moves to Iran, where he lives from 10 to 23 years. After 23 years, he moves to the States for the purpose of getting an education.

The paper also describes Studying and working in the States (Doctoral studies) and some information about his Family.

Zadeh's articles and the first results of scientific work as well as Scientific discoveries (how the theory was born, the history of this theory and features) are described in detail in this work too.

His collaboration with other scientists (his teacher Norbert Wiener, Rudolf Kalman, Charles A. Desoer and others) is written in detail in this article.

The work describes scientific life by years to the last stage of life.

The very important conferences and symposiums that he participated in and he make presentation as a keynote speaker are described in detail in this work.

The paper also describes in detail his visit to his native country, where he was born, which he visited Azerbaijan International Exhibition Telecommunications and Information Technologies

—BakuTel 2008 and reception of the President of the home country and his close ties with Azerbaijan, scientific friends from Azerbaijan.



Dr. Shahnaz N. Shahbazova received the academic degree of Ph.D. in 1995, the academic title of associate professor in 1996, the academic degree of Doctor of Sciences in Engineering 2015. Since 2002, she has been an academician, and since 2014, Vice-President of the Lotfi A. Zadeh International Academy of Sciences. Since 2011, she has held the position of General Chairman and organizer of the World Conference on Soft Computing (WConSC), dedicated to the preservation and development of the scientific heritage of Prof. Lotfi A. Zadeh. She is an honorary professor at the University of Obuda in Hungary and the “Aurel Vlaicu” University of Arad in Romania, an Honorary Doctor of Philosophy of Technical Sciences of the UNESCO International Personnel Academy.

She is a member of the Board of Directors of the North American Society for Fuzzy Information Processing (NAFIPS); moderator of the Berkeley Soft Computing Initiative (BISC) group; member of the council 3338.01 “System Analysis, Management and Information Processing” for the defense of Ph.D. and doctoral dissertations, as well as chairman of the seminar of the same council at the Azerbaijan Technical University.

Dr. Shahbazova participated in many international conferences as Organizer, Honorary Chair, Session Chair, member in Steering, Advisory or International Program Committees and Keynote Speaker. She is a president of Fuzzy System Association in Baku.

She’s awards: India—3 months (1998), Germany (DAAD)—3 months (1999, 2003, 2010), USA, California, Berkeley University (Fulbright)—(2007–2008, 2012, 2015–2017). She is the author of more than 242 scientific articles, 6 methodological manuals, 8 textbooks, 2 monographs and 13 Springer publications. Her research interests include Artificial Intelligence, Soft Computing, Intelligent Systems, Machine Learning Methods for Decision-Making and Fuzzy Neural Network.

A New Extension of Bellman and Zadeh's Decision Making in a Fuzzy Environment Via an Extended and/or Aggregation

Professor Janusz Kacprzyk

Intelligent System Laboratory

Systems Research Institute

Polish Academy of Sciences

Newelska, Warsaw, Poland

kacprzyk@ibspan.waw.pl

Abstract We propose an extension of the basic Bellman and Zadeh's (1970) model of decision-making in a fuzzy environment which is a point of departure of virtually all fuzzy models of decision-making, optimization, control, etc. The model involves in its original, fundamental version the fuzzy constraints and fuzzy goals, both represented as fuzzy sets in the space of options. The fuzzy constraints and fuzzy goals are aggregated using the “and” connective, corresponding to the intersection of fuzzy sets, and usually represented by the minimum or other triangular norm, e.g., the algebraic product (Cf. Kacprzyk, Yager, Merigo, 2019). This aggregation yields a performance function given as a fuzzy decision, and an optimal option is sought that maximizes its membership function.

This powerful model has been very successful, and over the years, some extension has been proposed. Basically, the satisfaction of the fuzzy goals and constraints driven by the traditional “and” connective, even with weights, has suggested that one should not take into account the particular fuzzy goal and fuzzy constraint if its/their satisfaction could interfere, in some sense, with the satisfaction with other fuzzy goal(s)/constraint(s).

We consider first some historical solution to this problem, notably by Yager (1992, 1996), and Bordogna and Pasi (1994), which can be viewed as related to the situation when an “interference” can occur.

In the second part we will present an application of Zadrożny and Kacprzyk's (2012) approach to the so-called bipolarity, originally proposed in the database querying context, when the satisfaction of some goal(s)/constraint(s) is mandatory and the satisfaction of some other goal(s)/constraint(s) is optional which can be verbally represented by “satisfy some specific goal(s)/constraint(s)” and, if possible “satisfy some other specific other goal(s)/constraint(s)”. We show how such a higher order structure can be modeled using multivalued logic and show examples of its use in real estate decision-making.



Janusz Kacprzyk is Professor of Computer Science at the Systems Research Institute, Polish Academy of Sciences, WIT—Warsaw School of Information Technology, and Chongqing Three Gorges University, Wanzhou, Chongqing, China, and Professor of Automatic Control at PIAP—Industrial Institute of Automation and Measurements in Warsaw, Poland. He is Honorary Foreign Professor at the Department of Mathematics, Yli Normal University, Xinjiang, China. He is Full Member of the Polish Academy of Sciences, Member of Academia Europaea, European Academy of Sciences and Arts, European Academy of Sciences, Foreign Member of the: Bulgarian Academy of Sciences, Spanish Royal Academy of Economic and Financial Sciences (RACEF), Finnish Society of Sciences and Letters, Flemish Royal Academy of Belgium of Sciences and the Arts (KVAB), National Academy of Sciences of Ukraine and Lithuanian Academy of Sciences. He was awarded with 6 honorary doctorates. He is Fellow of IEEE, IET, IFSA, EurAI, IFIP, AAIA, I2CICC, and SMIA.

His main research interests include the use of modern computation computational and artificial intelligence tools, notably fuzzy logic, in systems science, decision-making, optimization, control, data analysis and data mining, with applications in mobile robotics, systems modeling, ICT, etc.

He authored 7 books, (co)edited more than 150 volumes, (co)authored more than 650 papers, including ca. 150 in journals indexed by the WoS. He is listed in 2020 and 2021 “World’s 2% Top Scientists” by Stanford University, Elsevier (Scopus) and SciTech Strategies and published in *PLOS Biology Journal*.

He is the editor in chief of 8 book series at Springer, and of 2 journals, and is on the editorial boards of ca. 40 journals. He is President of the Polish Operational and Systems Research Society and Past President of International Fuzzy Systems Association.

How to Make a Decision Under Interval Uncertainty If We Do Not Know the Utility Function

Professor Vladik Kreinovich

University of Texas at El Paso

El Paso, USA

Abstract Decision theory describes how to make decisions, in particular, how to make decisions under interval uncertainty. However, this theory's recommendations assume that we know the utility function—that describes the decision-maker's preferences.

Sometimes, we can make a recommendation even when we do not know the utility function. In this paper, we provide a complete description of all such cases.



Vladik Kreinovich received his MS in Mathematics and Computer Science from St. Petersburg University, Russia, in 1974, and Ph.D. from the Institute of Mathematics, Soviet Academy of Sciences, Novosibirsk, in 1979. From 1975 to 1980, he worked with the Soviet Academy of Sciences; during this time, he worked with the Special Astrophysical Observatory (focusing on the representation and processing of uncertainty in radioastronomy). For most of the 1980s, he worked on error estimation and intelligent information processing for the National Institute for Electrical Measuring Instruments, Russia. In 1989, he was a visiting scholar at Stanford University.

Since 1990, he has worked in the Department of Computer Science at the University of Texas at El Paso. In addition, he has served as an invited professor in Paris (University of Paris VI), France; Hannover, Germany; Hong Kong; St. Petersburg and Kazan, Russia; and Brazil.

His main interests are the representation and processing of uncertainty, especially interval computations and intelligent control. He has published 12 books, 39 edited books, and more than 1800 papers. Vladik is a member of the editorial board of the international journal *Reliable Computing* (formerly *Interval Computations*) and several other journals. In addition, he is the co-maintainer of the international Web site on interval computations <http://www.cs.utep.edu/interval-comp>.

Vladik is Vice President of the International Fuzzy Systems Association (IFSA), Vice President of the European Society for Fuzzy Logic and Technology (EUSFLAT), Fellow of International Fuzzy Systems Association (IFSA), Fellow of Mexican Society for Artificial Intelligence (SMIA), Fellow of the Russian Association for Fuzzy Systems and Soft Computing, Treasurer of the IEEE Systems, Man, and Cybernetics Society; he served as Vice President for Publications of IEEE Systems, Man, and Cybernetics Society 2015–18 and as President of the North American Fuzzy Information Processing Society 2012–14; he is a foreign member of the Russian Academy of Metrological Sciences; he was the recipient of the 2003 El Paso Energy Foundation Faculty Achievement Award for Research awarded by the University of Texas at El Paso; and he was a co-recipient of the 2005 Star Award from the University of Texas System.

Aggregation 2.024

Bernard De Baets

KERMIT

Department of Data Analysis and Mathematical Modelling

Ghent University

Gent, Belgium

Abstract The study of aggregation functions is undeniably one of the most important spin-offs of the fuzzy set community, mainly driven by the need for appropriate logical connectives. The generated body of knowledge has also contributed to the fields of multi-criteria decision-making (e.g. OWA operators) and dependence modelling (e.g. quasi-copulas and copulas). However, the focus has increasingly narrowed to the aggregation of real numbers in the unit interval (letting aside a small group of researchers focusing on posets and lattices). In this setting of real numbers, Yager's penalty functions traditionally play an important role.

In this 'era of aggregation', data aggregation has become a central problem in many fields of application, with needs extending largely beyond the setting of real numbers, such as the aggregation of multidimensional data, compositional data, rankings, relations and strings. In this talk, we draw attention to the old notion of a betweenness relation and propose to replace the currently required property of quasi-convexity of a penalty function by the compatibility with a betweenness relation. Several methods for constructing penalty functions using a monometric, which carries a more appropriate semantics than the more commonly used distance function (metric), are provided. In this way, we are able to model aggregation processes on any set of objects equipped with a betweenness relation, thus considerably expanding the scope of the theory of aggregation.



Bernard De Baets is a senior full professor in applied mathematics at Ghent University, Belgium, where he is leading the research unit KERMIT. As a trained mathematician, computer scientist and knowledge engineer, he has developed a passion for multi- and interdisciplinary research. Over the past 25 years, more than 100 Ph.D. students have graduated under his (co-)supervision. Bernard is a prolific writer, with a bibliography comprising over 650 peer-reviewed journal papers, accumulating more than 32,000 Google Scholar citations (h-index 88). Several of his works have been bestowed upon with a best paper award.

Moreover, he is a much-invited speaker, having delivered over 250 lectures worldwide. Bernard also actively serves the research community, in particular as co-editor-in-chief of *Fuzzy Sets and Systems*.

Bernard has been an affiliated professor at the Anton de Kom Universiteit (Suriname); he is an Honorary Professor of Budapest Tech (Hungary), a Doctor Honoris Causa of the University of Turku (Finland), a Professor Invitado of the Universidad Central “Marta Abreu” de las Villas (Cuba), a Professor Extraordinarius of the University of South Africa and an Honorary Chair Professor (2023) at Bennett University, India. In 2011, he was elected Fellow of International Fuzzy Systems Association (IFSA) and, in 2012, he was a nominee for the Ghent University Prometheus Award for Research. In 2019, he received the EUSFLAT Scientific Excellence Award, and in 2021 he was declared Honorary Member of the EUSFLAT Society.

Fuzzy and Type-2 Fuzzy Models of Time Series to Building Decision Support Systems Based on an Ontology Service

Professor Yarushkina Nadezhda

Rector of Ulyanovsk State Technical University

Ulyanovsk, Russia

Abstract The development of the economy creates new challenges for artificial intelligence methods. Such challenges include the processing of large volumes of data, the analysis of various dynamic indicators, the discovery of complex dependencies in the accumulated data, and the forecasting of the state of processes. The main point of presented study is the development of a set of analytical and prognostic methods.

Data analysis in the context of the features of the problem domain and the dynamics of processes are significant in various industries. Uncertainty modeling based on fuzzy logic allows building approximators for solving a large class of problems. In some cases, type-2 fuzzy sets in the model are used.

We propose the approach to solving the problem of controlling the complex organizational and technical systems based on hybrid models and a new component of intelligent decision support systems (DSS) that is integrated with control systems. The proposed component is based on fuzzy logic and knowledge engineering. We present a model of ontology to form the context of data analysis and time series modeling. The ontological context allows us to represent trends of the analyzed object indicators. An expert can add a set of fuzzy rules to the ontology for systems control based on the fuzzy inference.

Ontologies are designed to represent knowledge of complex structures and to perform inference tasks. Developers must use the OWL API and SWRL API libraries to use ontology features. They are impossible to use in DSSs written in programming languages not for Java Virtual Machines. The FuzzyOWL library and the FuzzyDL inference engine are required to work with fuzzy ontologies.

Integration of Ontologies and Neural Networks is an effective approach to solving the problem of detecting time series anomalies, taking into account the specifics of the subject area. We propose a method based on the integration of a neural network with long short-term memory (LSTM) and Fuzzy Web Ontology Language (Fuzzy OWL) ontology. A LSTM network is used for the mathematical search for anomalies in the first stage. The fuzzy ontology filters the detection results and draws an inference for decision-making in the second stage. The ontology contains a formalized representation of objects in the subject area and inference rules that select only those anomaly values that correspond to this subject area. We propose the architecture of a software system that implements this approach.

Proposed hybrid models were used to create a predictive approach to control software development based on a time series extracted from repositories and hosting platforms. We described this domain as example of using proposed hybrid models. The method analyses based on fuzzy logic principles were used for extracting parameters from repositories, the approach to creating time series models and forecasting their behavior of software project.



Professor Yarushkina Nadezhda Doctor of Sciences in Engineering, Professor; graduated from Ulyanovsk Polytechnic Institute, Russia; rector of Ulyanovsk State Technical University; member of the Russian Association of Artificial Intelligence an author of more than 390 scientific papers in the field of soft computing, fuzzy logic, and hybrid systems.

Author ID (RSCI): 10358; Author ID (Scopus): 6602353202; Researcher ID (WoS): B-4438-2014; ORCID: 0000-0002-5718-8732. e-mail: jng@ulstu.ru

Anytime Information Processing, Fuzzy Logic, and Neural Networks—A Challenging Combination to Cope with the Uncertainty and Finiteness of Real-World

Annamaria R. Varkonyi-Koczy

Obuda University
Budapest, Hungary

Abstract Nowadays practical solutions of engineering problems involve model integrated computing. Due to their flexibility, robustness, and easy interpretability, the application of soft computing (SC) models may have an exceptional role at many fields, especially in cases where the problem to be solved is highly nonlinear or when only partial, uncertain and/or inaccurate data is available.

Anytime techniques are the youngest member of the soft computing family. Systems based on this approach are flexible with respect to the available input data, time, and computational power: Anytime models and algorithms are able to generalize previous input information and to provide short response time if the required reaction time is significantly shortened due to failures or alarms appearing in the modeled system; or if one has to make decisions before sufficient information arrives or the processing can be completed, i.e., anytime methods are able to work in changing circumstances and can ensure continuous operation in resource, data, and time insufficient conditions with guaranteed response time and known error.

The presentation deals with the topic of anytime computing and its possible combinations with other soft computing techniques. During the talk we address the concepts and techniques of anytime systems together with solutions of anytime fuzzy and anytime neural approaches. In order to illustrate the presentation, several applications will also be discussed, focusing on areas such as signal-, image processing, fault diagnosis, and control.



Annamaria R. Varkonyi-Koczy was born in Budapest, Hungary, in 1957. She received the M.Sc. E.E., M.Sc. M.E.-T., and Ph.D. degrees from the Technical University of Budapest in 1981, 1983, and 1996, respectively. In 2010 she presented her D.Sc. Thesis at the Hungarian Academy of Sciences. She was a Researcher of the Research Institute for Telecommunication, Hungary, for six years, followed by four years with the Hungarian Academy of Sciences.

From 1991 till 2009 she has been with the Department of Measurement and Information Systems, Budapest University of Technology and Economics, in assistant professor and associate professor positions. Since 2009 she is full professor at Obuda University, till 2016 at the Institute of Mechatronics and Vehicle Engineering (from 2010 also Head of Institute), while from 2016 at the Institute of Automation. From 2013 she is also full professor at the Department of Mathematics and Informatics, J. Selye University, Slovakia. Her research interests include digital image and signal processing, uncertainty handling, machine intelligence, and intelligent computing. She is author/contributor of 26 books, 54 journal papers, and more than 230 conference papers. She is founding chair of the IEEE-WISP and IEEE-SOFA symposium series, general chair of Interacademia'2004, 2008, and 2012, IPC chair of SAMI'2004, SAMI'2005, SAMI'2006, IEEE-ICCA'2005, ICONS'2008, I2MTC'2010 and 2012. Dr. Varkonyi-Koczy is Fellow of IEEE (from 2011 till 2014 also Distinguished Lecturer), elected Member of the Hungarian Academy of Engineers, former Vice President of the Hungarian Fuzzy Association, and member of the John von Neumann Computer and the Measurement and Automation Societies.

Intelligent Rooftop Greenhouses as Systems Engineering

Subject: Objectives, Models, Control

Prof. Valentina Emilia Balas

“Aurel Vlaicu” University of Arad

Arad, Romania

Abstract One of the weapons we have in the fight against global warming is carbon offset. Although there are various ways to provide this offset, carbon sequestration by plants is the most natural and effective way. To heal our planet, it is desirable to increase the presence of plants in our cities through urban agriculture. Our solution in this regard is the Intelligent Rooftop Greenhouse, which promotes a close symbiosis between the tenants of the building and the plants placed on the roof, as well as the integration into IoT networks to monitor and increase air quality. The presentation will mention the latest research carried out within this topic by our team from “Aurel Vlaicu” University of Arad.



Valentina Emilia Balas is currently Full Professor in the Department of Automatics and Applied Software at the Faculty of Engineering, “Aurel Vlaicu” University of Arad, Romania.

She holds a Ph.D. Cum Laude, in Applied Electronics and Telecommunications from Polytechnic University of Timisoara. Dr. Balas is author of more than 350 research papers in refereed journals and International Conferences. Her research interests are in Intelligent Systems, Fuzzy Control, Soft Computing, Smart Sensors, Information Fusion, Modeling and Simulation.

She is the Editor-in Chief to *International Journal of Advanced Intelligence Paradigms* (IJAIP) and to *International Journal of Computational Systems Engineering* (IJCSE), is a member in Editorial Board member of several national and international journals and is evaluator expert for national, international projects and Ph.D. Thesis.

Dr. Balas is the director of Intelligent Systems Research Centre in “Aurel Vlaicu” University of Arad and Director of the Department of International Relations, Programs and Projects in the same university.

She served as General Chair of the International Workshop Soft Computing and Applications (SOFA) in nine editions organized in the interval 2005–2020 and held in Romania and Hungary.

Dr. Balas participated in many international conferences as Organizer, Honorary Chair, Session Chair, member in Steering, Advisory or International Program Committees and Keynote Speaker.

Recently she was working in a national project with EU funding support: BioCell-NanoART = Novel Bio-inspired Cellular Nano-Architectures—For Digital Integrated Circuits, 3M Euro from National Authority for Scientific Research and Innovation.

She is a member of European Society for Fuzzy Logic and Technology (EUSFLAT), member of Society for Industrial and Applied Mathematics (SIAM), a Senior Member IEEE, member in Technical Committee—Fuzzy Systems (IEEE Computational Intelligence Society), chair of the Task Force 14 in Technical Committee—Emergent Technologies (IEEE CIS), and member in Technical Committee—Soft Computing (IEEE SMCS).

Dr. Balas was past Vice-president (responsible with Awards) of IFSA—International Fuzzy Systems Association Council (2013–2015), is a Joint Secretary of the Governing Council of Forum for Interdisciplinary Mathematics (FIM)—A Multidisciplinary Academic Body, India, and recipient of the “Tudor Tanasescu” Prize from the Romanian Academy for contributions in the field of soft computing methods (2019).

Moore's Laws and the Higher Education

Prof. Marius M. Balas

“Aurel Vlaicu” University of Arad
Arad, Romania

Abstract Inspired by electronics, Moore's laws approximate the development of human activities by exponential functions. Indeed, the volume of our knowledge has grown exponentially, which increasingly hampers the process of training young people, especially at the higher undergraduate education level, resulting in higher school dropout figures. This phenomenon is partly due to the changes in mentality of young people from generation Z. Through the Bologna Declaration, higher education is divided into mass education and elite education. School dropout occurs mainly at the mass level—license degree. In this situation, the effectiveness of learning higher education must be increased. An important first measure is the transition, from the traditional bottom-up approach to Systems Engineering's top-down approach. The presentation also proposes some concrete measures in this direction: increasing orientation towards visual methods of representing and processing knowledge, teaching mathematics in close connection with the Matlab software package, etc.



Marius M. Balas, IEEE Senior Member, is a Habilitated Professor at The Engineering Faculty of “Aurel Vlaicu” University of Arad, Romania. His research topics are in Systems Engineering, Electronic Circuits, Intelligent and Fuzzy Systems, Adaptive Control, Modeling and Simulation. He is the author of 16 books and book chapters, 120 indexed papers and 7 invention patents.

His main scientific contributions are the fuzzy-interpolative systems, the passive greenhouses, the intelligent rooftop greenhouses, the constant time to collision traffic optimization, the imposed distance braking, PWM inverter for railway coaches in tropical environments, the rejection of the switching controllers’ instability by phase trajectory analysis and the Fermat neuron. He is Coordinator for several research projects and student competitions, co-organizing the Control Engineering and Industrial Software Branch of the Engineering Faculty of Aurel Vlaicu University.

Generative Adversarial Networks (GANs): From Perspective of Past, Present, and Future

Köksal Erentürk

Faculty of Engineering
Department of Electrical and Computer Engineering
Atatürk University
Erzurum, Türkiye

Abstract Generative Adversarial Networks (GANs), firstly introduced in 2014 by Ian Goodfellow et al., have revolutionized the field of artificial intelligence. This revolutionary technique consists of two neural networks against each other in a competitive game, fostering the creation of ever-more realistic and sophisticated data. In this presentation, we investigate the fascinating journey of GANs, exploring their past breakthroughs, current applications, and the exciting possibilities they hold for the future.

GANs operate in a fascinating way simple principle. One of the designed networks, called as the generator, strives to produce realistic data, be it images, music, or even 3D models. Meanwhile, the other designed network, named as the discriminator, acts as a discriminating critic, aiming to distinguish the generated data from real-world samples. Through this adversarial process, the generator continuously learns from the discriminator's critiques, refining its ability to create high-fidelity generative models or ever-more convincing forgeries.

Fast forward to today, and GANs have become a cornerstone of AI research. With the improvement of powerful computing resources and advancements in deep learning architectures, GANs have achieved remarkable mastery. They can generate better realistic images of people, landscapes, and even inexistent creatures. The applications are covering broad area of sciences, including creative endeavors like generating artistic styles or composing music, to practical uses like creating realistic medical imaging data or designing novel materials.

However, mode collapse, where the generator gets stuck producing a limited range of outputs, and training instability continue to be areas of active research are the some of the GANs present challenges. Additionally, ethical concerns regarding the potential for generating deepfakes require careful consideration.

Looking ahead, with advancements in techniques like progressive growing and spectral normalization offering promising solutions, researchers are exploring ways to address current limitations. Additionally, the collaborative AI frameworks allow researchers to build upon each other's work, accelerating progress.

Future applications of GANs are also captivating. Imagine personalized healthcare experiences with customized medical simulations or the creation of entirely new forms of art and entertainment. In drug discovery and scientific research, accelerating the pace of innovation is some of the future applications of GANs.



Köksal Erentürk has a B.Sc. degree from Yıldız Technical University, Department of Electronics and Telecommunication Engineering, İstanbul, Türkiye 1994; M.Sc. degree from the Istanbul University, Institute of Graduate Studies in Applied Sciences, İstanbul, Türkiye 1997; and Ph.D. degree in Electrical Engineering from Karadeniz Technical University, Trabzon, Türkiye 2002. Currently he is professor of electrical and computer engineering, Faculty of Engineering, Department of both Electrical and Computer Engineering, Atatürk University, Erzurum, Türkiye. He has also served as a faculty member at different universities in different countries such as Kuwait, Sweden, and Kingdom of Saudi Arabia.

He is a member of IEEE since 2008 and holding more than 10 awards from broad area of scientific societies. He has authored and coauthored more than 150 journal and conference papers. His work has focused on machine learning, artificial intelligence, deep learning, generative adversarial networks, fuzzy logic and control, artificial neural networks, genetic algorithms, development and application of control theory to a variety of mechatronic systems, with a focus on observation and estimation-based control as well as highly efficient power converters, resonant and multilevel converters, including the modeling, control, and development of prototypes for various industrial applications in electric traction, renewable energy, power supplies for drives, energy storage devices, and their optimal control for electrified transportation and renewable energy applications.

Level Set Fuzzy Modeling and Machine Learning

Fernando Gomide

School of Electrical and Computer Engineering

University of Campinas

Campinas, São Paulo, Brazil

Abstract Fuzzy modeling is an important subject of fuzzy systems theory and applications since the pioneering work of Zadeh. The nature of fuzzy models is highly diversified and includes tabular fuzzy models, fuzzy relational models and associative memories, fuzzy trees, fuzzy cognitive maps, fuzzy neural networks, and rule-based fuzzy models.

Rule-based fuzzy models prevail in fuzzy modeling. Two categories of rule-based fuzzy models, linguistic fuzzy models and functional fuzzy models, form the majority of the fuzzy modeling applications.

Actually, there is a third way to construct fuzzy models, which has left out the mainstream fuzzy modeling efforts, the level set-based fuzzy modeling method. Originally introduced by Yager, the level set-based method computes outputs using a function that maps the activation level of the fuzzy rules into a value of the output space. It is simple, easy to use, and very effective modeling approach.

The presentation briefly reviews the level set fuzzy modeling idea, addresses recent data-driven frameworks for level set-based fuzzy modeling, and discusses its modeling performance when compared with state-of-the-art machine learning and computational intelligence modeling frameworks.



Fernando Gomide has a BSEE degree from Polytechnic Institute of the Pontifical Catholic University of Minas Gerais, Belo Horizonte, Minas Gerais, Brazil, 1975; M.Sc. degree from the University of Campinas, School of Electrical and Computer Engineering, Campinas, São Paulo, Brazil, 1979; and Ph.D. degree in Systems Engineering from Case Western Reserve University, Cleveland, Ohio, USA, 1983. Currently he is professor of electrical and computer engineering, School of Electrical and Computer Engineering, University of Campinas, São Paulo, Brazil. He is IEEE Fellow Class 2016, K. S. Fu Awardee North American Fuzzy Information Society 2011, IFSA Fellow Class 2009, and Fellow 1A Brazilian National Council for Scientific and Technological Development 2000.

His areas of interest include modeling, optimization, control, decision-making, artificial intelligence, machine learning, fuzzy systems, and applications. He serves the editorial board of *Fuzzy Optimization and Decision Making*, *Journal of Advanced Computational Intelligence*, *Evolving Systems*, *Granular Computing*, and the *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*.

References

1. R. Pérez-Fernández, M. Rademaker and B. De Baets, Monometrics and their role in the rationalisation of ranking rules, *Information Fusion* 34 (2017), 16–27.
2. R. Pérez-Fernández, B. De Baets and M. Gagolewski, A taxonomy of monotonicity properties for the aggregation of multidimensional data, *Information Fusion* 52 (2019), 322–334.
3. R. Pérez-Fernández and B. De Baets, On the role of monometrics in penalty-based data aggregation, *IEEE Transactions on Fuzzy Systems* 27 (2019), 1456–1468.
4. M. Gagolewski, R. Pérez-Fernández and B. De Baets, An inherent difficulty in the aggregation of multidimensional data, *IEEE Transactions on Fuzzy Systems* 28 (2020), 602–606.
5. R. Pérez-Fernández, M. Gagolewski and B. De Baets, On the aggregation of compositional data, *Information Fusion* 73 (2021), 103–110.
6. R. Pérez-Fernández and B. De Baets, Aggregation theory revisited, *IEEE Transactions on Fuzzy Systems* 29 (2021), 797–804

Contents

Fuzzy Sets and Quantum Computing

What Fuzzy and Quantum Computing Can Learn From the Success of Deep Learning 3
Shahnaz N. Shahbazova and Vladik Kreinovich

How to Fairly Allocate Safety Benefits of Self-Driving Cars 11
Fernando Muñoz, Christian Servin, and Vladik Kreinovich

Paradox of Causality and Paradoxes of Set Theory 19
Alondra Baquier, Bradley Beltran, Gabriel Miki-Silva, Olga Kosheleva, and Vladik Kreinovich

Fuzzy Recognition System

Continuous Growth of Information and the “Forgetting” Function in the Recognition of Sound Information 27
Shahnaz N. Shahbazova

Technology Building a Virtual Room for Recognition of Visual and Sound Information 33
Shahnaz N. Shahbazova

Basic Methods for Recognizing Visual and Audio Information and the Results of Experiments 43
Huseynova Zinyet Revayet

Intelligent Methods and Fuzzy Approach

An Approach to Feature Engineering in a Data-Driven Production Management 55
Aleksey A. Filippov, Gleb Yu. Guskov, Anton A. Romanov, and Nadezhda G. Yarushkina

How to Make a Decision Under Interval Uncertainty If We Do Not Know the Utility Function	65
Jeffrey Escamilla and Vladik Kreinovich	
Number Representation With Varying Number of Bits	71
Anuradha Choudhury, Md Ahsanul Haque, Saeefa Rubaiyet Nowmi, Ahmed Ann Noor Ryen, Sabrina Saika, and Vladik Kreinovich	
Fuzzy Neural Networks	
How to Make a Neural Network Learn from a Small Number of Examples—and Learn Fast: An Idea	79
Chitta Baral and Vladik Kreinovich	
Analysis of Histological Preparations Using Convolutional Neural Networks	85
Guskov G. Yu, N. G. Yarushkina, and A. V. Chekina	
Why Two Fish Follow Each Other but Three Fish Form a School: A Symmetry-Based Explanation	95
Shahnaz Shahbazova, Olga Kosheleva, and Vladik Kreinovich	
Fuzzy Simulation and Data Mining	
Creating a Database for Assessment of the Reliability of the Operational State of Technical Facilities	105
Telman Aliev and Naila Musaeva	
GANs Applications in Chemical Engineering	113
Köksal Erentürk and Saliha Erentürk	
The Method of Group Temperature Measurements of Remote Atmospheric Objects	121
Huseynova Rena and Aliyeva Gunel	
Fuzzy Control Systems	
Composite Indicators Design Based on Concordant Information Using Semi-supervised Learning on Aggregated Data	131
Leonid Lyubchyk, Galyna Grinberg, and Klym Yamkovyi	
Software of Artificial Intelligence Control System for UAV	143
Azad Bayramov and Samir Suleymanov	
Application Fuzzy PID Controller in UAV	151
Azad Bayramov, Samir Suleymanov, and Fatali Abdullaev	

Fuzzy Applications and Analysis System

Data Fusion Is More Complex Than Data Processing: A Proof	161
Robert Alvarez, Salvador Ruiz, Martine Ceberio, and Vladik Kreinovich	

A Comparative Analysis of AI-Based Cryptocurrencies: Fundamental and Technical Perspectives	167
Hilala Jafarova	

Application and Enhancement of New Technologies for Cyber Risk Prediction	177
Shahnaz N. Shahbazova, Semral Hajiyeve, and Hilala Jafarova	

Soft Computing and Fuzzy Logic

From Question to Answer: Innovative AI-Based System for Exam Preparation	189
Shahnaz N. Shahbazova and Agha P. Aghayev	

Recursive Approach to Algorithmic Analyses of Language of Mythes and Legends	199
Mammadova Ana Aga and Leyla Aliyeva	

A CNN-Based Handwritten Digit Classifier: Architecture, Training, and Performance Evaluation	207
Shahnaz N. Shahbazova	

Author Index	217
---------------------------	-----

About the Editors



Dr. Shahnaz N. Shahbazova received the academic degree of Ph.D. in 1995, the academic title of associate professor in 1996, the academic degree of Doctor of Sciences in Engineering 2015. Since 2002, she has been an academician and since 2014 Vice-President of the Lotfi A. Zadeh International Academy of Sciences. Since 2011, she has held the position of General Chairman and organizer of the World Conference on Soft Computing (WConSC), dedicated to the preservation and development of the scientific heritage of Prof. Lotfi A. Zadeh. She is an honorary professor at the University of Obuda in Hungary and the “Aurel Vlaicu” University of Arad in Romania, an Honorary Doctor of Philosophy of Technical Sciences of the UNESCO International Personnel Academy.

She is a member of the Board of Directors of the North American Society for Fuzzy Information Processing (NAFIPS); moderator of the Berkeley Soft Computing Initiative (BISC) group; member of the council 3338.01 “System Analysis, Management and Information Processing” for the defense of Ph.D. and doctoral dissertations, as well as chairman of the seminar of the same council at the Azerbaijan Technical University.

Dr. Shahbazova participated in many international conferences as Organizer, Honorary Chair, Session Chair, member in Steering, Advisory or International Program Committees and Keynote Speaker. She is a president of Fuzzy System Association in Baku.

She's awards: India—3 months (1998), Germany (DAAD)—3 months (1999, 2003, 2010), USA, California, Berkeley University (Fulbright)—(2007–2008, 2012, 2015, 2016, 2017). She is the author of more than 242 scientific articles, 6 methodological manuals, 8 textbooks, 2 monographs and 13 Springer publications. Her research interests include Artificial Intelligence, Soft Computing, Intelligent Systems, Machine Learning Methods for Decision-Making and Fuzzy Neural Network.



Vladik Kreinovich received his MS in Mathematics and Computer Science from St. Petersburg University, Russia, in 1974, and Ph.D. from the Institute of Mathematics, Soviet Academy of Sciences, Novosibirsk, in 1979. From 1975 to 1980, he worked with the Soviet Academy of Sciences; during this time, he worked with the Special Astrophysical Observatory (focusing on the representation and processing of uncertainty in radioastronomy). For most of the 1980s, he worked on error estimation and intelligent information processing for the National Institute for Electrical Measuring Instruments, Russia. In 1989, he was a visiting scholar at Stanford University.

Since 1990, he has worked in the Department of Computer Science at the University of Texas at El Paso. In addition, he has served as an invited professor in Paris (University of Paris VI), France; Hannover, Germany; Hong Kong; St. Petersburg and Kazan, Russia; and Brazil.

His main interests are the representation and processing of uncertainty, especially interval computations and intelligent control. He has published 12 books, 39 edited books, and more than 1800 papers. Vladik is a member of the editorial board of the international journal *Reliable Computing* (formerly *Interval Computations*) and several other journals. In addition, he is the co-maintainer of the international Web site on interval computations <http://www.cs.utep.edu/interval-comp>.

Vladik is Vice President of the International Fuzzy Systems Association (IFSA), Vice President of the European Society for Fuzzy Logic and Technology (EUSFLAT), Fellow of International Fuzzy Systems Association (IFSA), Fellow of Mexican Society for Artificial Intelligence (SMIA), Fellow of the Russian

Association for Fuzzy Systems and Soft Computing, Treasurer of the IEEE Systems, Man, and Cybernetics Society; he served as Vice President for Publications of IEEE Systems, Man, and Cybernetics Society 2015–18 and as President of the North American Fuzzy Information Processing Society 2012–14; he is a foreign member of the Russian Academy of Metrological Sciences; he was the recipient of the 2003 El Paso Energy Foundation Faculty Achievement Award for Research awarded by the University of Texas at El Paso; and he was a co-recipient of the 2005 Star Award from the University of Texas System.



Janusz Kacprzyk is Professor of Computer Science at the Systems Research Institute, Polish Academy of Sciences, WIT—Warsaw School of Information Technology, and Chongqing Three Gorges University, Wanzhou Chongqing, China, and Professor of Automatic Control at PIAP—Industrial Institute of Automation and Measurements in Warsaw, Poland. He is Honorary Foreign Professor at the Department of Mathematics, Yli Normal University, Xinjiang, China. He is Full Member of the Polish Academy of Sciences, Member of Academia Europaea, European Academy of Sciences and Arts, European Academy of Sciences, Foreign Member of the: Bulgarian Academy of Sciences, Spanish Royal Academy of Economic and Financial Sciences (RACEF), Finnish Society of Sciences and Letters, Flemish Royal Academy of Belgium of Sciences and the Arts (KVAB), National Academy of Sciences of Ukraine and Lithuanian Academy of Sciences. He was awarded with 6 honorary doctorates. He is Fellow of IEEE, IET, IFSA, EurAI, IFIP, AAIA, I2CICC, and SMIA.

His main research interests include the use of modern computation computational and artificial intelligence tools, notably fuzzy logic, in systems science, decision-making, optimization, control, data analysis and data mining, with applications in mobile robotics, systems modeling, ICT, etc.

He authored 7 books, (co)edited more than 150 volumes, (co)authored more than 650 papers, including ca. 150 in journals indexed by the WoS. He is listed in 2020 and 2021 “World’s 2% Top Scientists” by Stanford University, Elsevier (Scopus) and SciTech Strategies and published in *PLOS Biology Journal*.

He is the editor in chief of 8 book series at Springer, and of 2 journals, and is on the editorial boards of ca. 40 journals. He is President of the Polish Operational and Systems Research Society and Past President of International Fuzzy Systems Association.



Bernard De Baets is a senior full professor in applied mathematics at Ghent University, Belgium, where he is leading the research unit KERMIT. As a trained mathematician, computer scientist and knowledge engineer, he has developed a passion for multi- and interdisciplinary research. Over the past 25 years, more than 100 Ph.D. students have graduated under his (co-) supervision. Bernard is a prolific writer, with a bibliography comprising over 650 peer-reviewed journal papers, accumulating more than 32000 Google Scholar citations (h-index 88). Several of his works have been bestowed upon with a best paper award.

Moreover, he is a much-invited speaker, having delivered over 250 lectures worldwide. Bernard also actively serves the research community, in particular as co-editor-in-chief of *Fuzzy Sets and Systems*.

Bernard has been an affiliated professor at the Anton de Kom Universiteit (Suriname); he is an Honorary Professor of Budapest Tech (Hungary), a Doctor Honoris Causa of the University of Turku (Finland), a Professor Invitado of the Universidad Central “Marta Abreu” de las Villas (Cuba), a Professor Extraordinarius of the University of South Africa and an Honorary Chair Professor (2023) at Bennett University, India. In 2011, he was elected Fellow of International Fuzzy Systems Association (IFSA) and, in 2012, he was a nominee for the Ghent University Prometheus Award for Research. In 2019, he received the EUSFLAT Scientific Excellence Award and in 2021 he was declared Honorary Member of the EUSFLAT Society.

Fuzzy Sets and Quantum Computing

What Fuzzy and Quantum Computing Can Learn From the Success of Deep Learning



Shahnaz N. Shahbazova and Vladik Kreinovich

Abstract How can we apply the ideas that made deep neural networks successful to other aspects of computing? For this purpose, we reformulate these ideas in a more general form—and we show that this generalization also covers fuzzy and quantum computing. This enables us to suggest that similar ideas can be helpful for fuzzy and quantum computing as well. In this suggestion, we are encouraged by the fact that as we show, to some extent, these ideas are already helpful.

1 Formulation of the Problem

A natural idea is to learn from successes. We all, scientists and practitioners, try our best to solve our problems better—more effectively, more efficiently, etc. We want to better (and faster) predict all aspects of the future state of the world, we want to compute designs and controls that will make the world’s future state even better.

From this viewpoint, every time an approach leads to a success, a natural idea is: how can we use the corresponding successful idea(s) to make improvements in other areas as well?

From this viewpoint, what can we learn from the successes of deep learning. At present, in AI—and, arguably, in computer science in general—one of the most successful directions is deep learning; see, e.g., [4]. It is therefore reasonable to ask: how can we use the main ideas behind this success to make improvements in other directions of computing as well?

This is the problem that we concentrate on in this paper.

S. N. Shahbazova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics UNEC, Baku, Azerbaijan
e-mail: shahbazova.shahnaz@unec.edu.az

V. Kreinovich

Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA
e-mail: vladik@utep.edu

Comment. Of course, some researchers may say (and have said): just abandon your less-successful ideas and jump on deep learning bandwagon—but, in contrast to these researchers, we believe that all approaches have potential and can benefit each other. And we will provide examples showing that our belief is justified.

Why quantum and fuzzy computing? In order to analyze how we can generalize the success of neural networks—in their deep learning form, a natural idea is to reformulate the main ideas behind neural networks and deep learning ideas in a more general form. This is what we will do, and we will show that this generalization naturally leads to fuzzy and quantum computing.

2 Let Us Reformulate the Main Ideas Behind Neural Networks and Deep Learning in a More General Form

What is the main challenge of computing? In order to come up with the desired generalization, let us recall what is the main problem of computing now.

For many practical problems, we have, at least on the theoretical level, algorithms for the desired prediction and/or for the desired control. For many problems, these algorithms are practically useful. For example, the existing algorithms can predict tomorrow's weather reasonably well—not perfectly well, but definitely much better than it was possible even a few years ago. Many of these algorithms are so good that, e.g., modern airplanes are designed and tested on computer simulations—and the following flight tests only confirm the simulation results.

However, in some problems, we have algorithms, but these algorithms require so much computation time that they become practically useless. For example, it is possible to predict in which direction a potentially deadly tornado will turn in the next 15 min—this can be done by using algorithms similar to weather prediction—but this prediction is practically useless since it takes several hours on a high-performance computer. By the time we have computed this prediction, the tornado has already moved.

Such examples are plentiful. The need to make computations faster is one of the main challenges of computing.

This challenge is objective. In some cases, it is possible to come up with faster algorithms for solving the same problem. However, in general, there is a limit to how much we can gain this way. Many practical problems are known to be NP-hard. This means that unless $P = NP$ (which most computer scientists believe to be impossible), no feasible algorithm can solve all the instances of this general problem; see, e.g., [6, 11]. In other words, no matter how clever our algorithms, there will always be instances on which these algorithms will take too long to be practically useful.

Since there is a limit to how much we can speed up computations by coming up with better algorithms, it is crucially important to make sure that the implementation of current algorithms is as fast as possible.

How to make the algorithm implementation faster? An algorithm, by definition, consists of elementary computational steps. What are these elementary steps—i.e., hardware supported elementary operations—depends on the computational device. So, to speed up computations, we need:

- To select elementary computational steps that can be computed in the fastest possible way, and
- To actually implement these elementary steps in the fastest possible manner.

Let us analyze these two aspects one by one.

Which elementary steps are the fastest? In sensors and in computers, information is transferred by electric signals.

Some signals are analog. In these signals, the information is conveyed by different values of current and/or voltage. This is how most sensors operate: e.g., a photosensor transforms the intensity of light into the intensity of the corresponding current.

For electric signals, all basic transformations are linear: Ohm's law—according to which voltage V linearly depends on current I , as $V = I \cdot R$ —is clearly linear, more general Kirchhoff's laws are also linear. For analog signals, any linear function is easy to implement:

- Multiplication by a constant corresponds to using different resistances R , and
- If we bring two currents I_1 and I_2 together, the resulting current I will be equal to their sum $I = I_1 + I_2$.

In principle, we can have non-linear electric devices, but the fastest are linear transformations.

In most modern computational devices, signals are digital, they are represented as a sequence of bits. For numbers represented in binary form, addition or multiplication by a number are still reasonably fast, but they are no longer the fastest possible operations: just like when we add multi-digit numbers, we perform several digit operations, so does the computer. For digital signals, the fewer bits we have to compute, the faster the computations. From this viewpoint, the fastest are operations in which we compute only one bit—i.e., for which, in effect, we decide whether some property is true or false. For numbers, the only such properties are equality and inequality, so the fastest operations are min and max. Indeed, in each of these operations:

- First, we decide—via computations—which number is smaller (or larger), and
- Then return this smaller (or larger) number—without performing any additional computations.

How to actually implement these elementary operations in the fastest possible way? One of the main factors that limits computer speed is the speed of light—since, according to modern physics, no communication can be faster than the speed of light; see, e.g., [3, 12].

This may sound like a remote restriction, not something to worry about – but we are already close to this limit. For a typical computer which is about 30cm across, it takes $30\text{ cm}/300\,000\text{ km/sec} = 1\text{ nanosecond}$ to go across. During this time, the standard 4GHz computer already performs 4 operations! So, the only way to make computers much faster is to make them much smaller—and this means making all their components much smaller.

Already each memory cell is of the size of thousands or even hundreds of molecules. If we make this cell smaller, its size will approach the size of a single molecule. For such small objects, we can no longer use the laws of Newtonian mechanics, we need to use special physics of microworld known as quantum physics.

Computing that takes quantum effects into account is known as *quantum computing*; see, e.g., [9]. Interestingly, one of the main features of quantum physics is that most its processes are linear, so we again go to the need to use linear transformations [3, 9, 12].

So which elementary operations are the fastest? Our analysis shows that the fastest possible elementary operations are:

- Linear transformations, that transform inputs $x_1, \dots, x_n$ into their linear combination $w_0 + w_1 \cdot x_1 + \dots + w_n \cdot x_n$, and
- Min- and max-transformations that transform inputs $x_1, \dots, x_n$ into either $\min(x_1, \dots, x_n)$ or $\max(x_1, \dots, x_n)$.

Unfortunately, fastest transformations are not sufficient. It would be great if we could only use these fastest elementary operations, but, unfortunately, they are not sufficient:

- If we only use linear transformations, then we can only compute linear functions—and many real-life processes are non-linear;
- If we only use min and max—both of which select one of the input values as the output—we will always return one of the inputs, and we will never be able to compute anything else.

So, In addition to the fastest elementary operations, we need something else:

- In addition to linear transformations, we need to use some non-linear transformations, and
- In addition to min and max, we need to use some operations that return the value which is different from one of the inputs.

This is exactly what leads us to neural, fuzzy, and quantum computing. Let us first consider the case when we add some non-linear transformation to linear operations. In this cases, computations means that we interchangingly apply linear transformations and some non-linear operations $y = f(x)$. This is exactly what neural networks do. On each layer, each processing unit—called a *neuron*—first computes a linear combination of inputs (or, for next layers, outputs from the previous layer), and

then applies some non-linear transformation to the result. For neural networks, the corresponding transformation is known as an *activation function*; see, e.g., [2].

For artificially set up neural networks, we can select any activation function we want. For quantum computing, we have to rely on nature's non-linear process. Such a process is known as *measurement*. If we have a quantum state s which is a linear combination $s = c_1 \cdot s_1 + \dots + c_n \cdot s_n$ of the basic states $s_1, \dots, s_n$, then measurement transforms this state into one of the states s_i with probability $|c_i|^2$; see, e.g., [3, 9, 12].

The operations min and max correspond to the most widely used operations of *fuzzy logic*, corresponding to “and” and “or”. In this case, the corresponding additional operation is known as *defuzzification*; see, e.g., [1, 5, 7, 8, 10, 13].

Thus, our general approach to computations indeed leads to neural, fuzzy, and quantum computing.

3 From This General Viewpoint, What Can We Learn from Deep Learning?

A seemingly reasonable idea. At first glance, the situation is straightforward:

- In addition to the fastest elementary computational steps, we need at least one not-so-fast more complex one;
- Overall, we want to speed up computations;
- Thus, we need to limit the number of not-so-fast computational steps to a minimum—if possible, to just one such step.

As a result, it is reasonable to use exactly one not-so-fast computational step.

This is exactly how traditional neural networks worked. This is exactly how traditional neural networks worked (see, e.g., [2]):

- First, each neuron k transformed the inputs $x_1, \dots, x_n$ into their linear combination $y_k = w_{k0} + w_{k1} \cdot x_1 + \dots + w_{kn} \cdot x_n$;
- then, we apply an appropriate nonlinear transformation $z = f(x)$ to each of these results, getting $z_k = f(y_k)$ for each k , and
- Finally, we apply a linear transformation to the values z_k , resulting in $y = W_0 + W_1 \cdot z_1 + W_2 \cdot z_2 + \dots$.

This is exactly how fuzzy and quantum computing work now. How does a usual fuzzy control works?

- We apply min and max operations to the original values of the measurement functions, and then
- We apply an appropriate defuzzification procedure to the result, generating the recommended control value.

This is also how usual quantum computing algorithms work:

- We perform some linear quantum operations, and then
- We perform a measurement to get the result.

The above seemingly natural idea is only a first approximation. The above idea would work perfectly if computing fastest elementary operation would require no time at all, and the only computing time would be spent of computing other not-so-fast operations. In reality, computing fastest elementary operations also spends time, and these times accumulate. Sometimes, it may be more efficient to perform more not-so-fast operations.

A simple analogy can explain this. Suppose that you own a car. For a car, repairs are usually cheaper than buying a new car. So, by the same logic as we described above, to save money, we should postpone buying a new car until it is absolutely necessary: e.g., when our previous car stops working. This advice is reasonable at the beginning, when the car is not very old, but with time, it becomes more expensive to continuously spend more and more money on more and more frequent repairs than simply to buy a newer car.

Similarly here, while in general, the above logic sounds reasonable in the first approximation, in reality, it may be faster to have two or more not-so-fast elementary operations than to limit ourselves to only one such operation.

When does it make sense to use several not-so-fast operations. When we have a small number of inputs and simple computations, probably the above argument works: not-so-fast elementary operations are still much slower than all corresponding faster elementary operations taken together. In this case, limiting ourselves to only one not-so-fast operation makes perfect sense.

However, as the number of inputs grows and the computations become more complex, the overall time needed for all elementary operations becomes comparable with—and even larger—than the time needed for a single not-so-fast operation. In this case, using more than one not-so-fast operation may be beneficial.

What deep learning showed. This is exactly what happened in neural networks: it turned out that in many complex applications, it is more beneficial to use several non-linear layers. This is what is called *deep learning*—when we have several nonlinear layers.

So what can fuzzy and quantum computing learn from this success. In view of the above analysis, a natural idea is to allow several not-so-fast layers.

In fuzzy control, it means that instead of a single one-stage fuzzy controller, we should use multi-stage (e.g., hierarchical) fuzzy controllers, where results of some stages are used as input to other controllers. This indeed works—hierarchical fuzzy controllers have been shown to be efficient in controlling complex systems.

In quantum computing, it means that instead of trying to fit everything into a single quantum algorithm, it may be better to have a multi-stage algorithms, in which we first perform measurements, and then do something with this measurement results. This also works—e.g., this is how quantum cryptography works: we first perform measurements, and then process the results of these measurements; see, e.g., [9].

Our examples show that this idea is not as new and as revolutionary as it may seem at first glance. Our point is that at present, this idea is not the mainstream neither in fuzzy nor in quantum computing. Our argument is that, as we handle more and more complex problems, with more and more inputs, this idea should become mainstream.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), and HRD-1834620 and HRD-2034030 (CAHSI Includes), and by the AT&T Fellowship in Information Technology.

It was also supported by the program of the development of the Scientific-Educational Mathematical Center of Volga Federal District No. 075-02-2020-1478.

References

1. R. Belohlavek, J.W. Dauben, G.J. Klir, *Fuzzy Logic and Mathematics: A Historical Perspective* (Oxford University Press, New York, 2017)
2. C.M. Bishop, *Pattern Recognition and Machine Learning* (Springer, New York, 2006)
3. R. Feynman, R. Leighton, M. Sands, *The Feynman Lectures on Physics* (Addison Wesley, Boston, Massachusetts, 2005)
4. I. Goodfellow, Y. Bengio, A. Courville, *Deep Learning* (MIT Press, Cambridge, Massachusetts, 2016)
5. G. Klir, B. Yuan, *Fuzzy Sets and Fuzzy Logic* (Prentice Hall, Upper Saddle River, New Jersey, 1995)
6. V. Kreinovich, A. Lakeyev, J. Rohn, P. Kahl, *Computational Complexity and Feasibility of Data Processing and Interval Computations* (Kluwer, Dordrecht, 1998)
7. J.M. Mendel, *Uncertain Rule-Based Fuzzy Systems: Introduction and New Directions* (Springer, Cham, Switzerland, 2017)
8. H.T. Nguyen, C.L. Walker, E.A. Walker, *A First Course in Fuzzy Logic* (Chapman and Hall/CRC, Boca Raton, Florida, 2019)
9. M.A. Nielsen, I.L. Chuang, *Quantum Computation and Quantum Information* (Cambridge University Press, Cambridge, U.K., 2000)
10. V. Novák, I. Perfilieva, J. Močkoř, *Mathematical Principles of Fuzzy Logic* (Kluwer, Boston, Dordrecht, 1999)
11. C.H. Papadimitriou, *Computational Complexity* (Addison Wesley, San Diego, CA, 1994)
12. K.S. Thorne, R.D. Blandford, *Modern Classical Physics: Optics, Fluids, Plasmas, Elasticity, Relativity, and Statistical Physics* (Princeton University Press, Princeton, New Jersey, 2017)
13. L.A. Zadeh, Fuzzy sets. Inf. Control **8**, 338–353 (1965)

How to Fairly Allocate Safety Benefits of Self-Driving Cars



Fernando Muñoz, Christian Servin, and Vladik Kreinovich

Abstract In this paper, we describe how to fairly allocated safety benefits of self-driving cars between drivers and pedestrians—so as to minimize the overall harm.

1 Formulation of the Practical Problem

Our roads are still not perfectly safe. In spite of all the efforts of transportation engineers, there are still about 6 million car accidents every year in the US only.

How can we make roads safer: the promise of self-driving cars. Most road accidents are caused by human errors—which, in their turn, are caused by fatigue, distractions, etc. Because of this, a natural way to decrease this number of accidents and to make roads safer is to replace some—or even all—human decision making with decision made by automated systems. The current self-driving cars have not yet reached their safety potential, but they are getting better, and we all hope that in the nearest future they will be safer to drive than the usual human-driven vehicles.

Resulting problem. Self-driving cars will (hopefully) bring more safety. When an automatic system is in charge, a system that does not know fatigue, that is never under the influence, that does not violate traffic rules, there will be hopefully much fewer accidents, and most of the remaining accidents will be less severe than they are now—since the reaction time of an automatic system is much smaller than the reaction time of a human driver.

F. Muñoz · V. Kreinovich (✉)

Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA

e-mail: vladik@utep.edu

F. Muñoz

e-mail: fmunoz9@miners.utep.edu

C. Servin

Information Technology Systems Department, El Paso Community College (EPCC), El Paso, TX, USA

e-mail: cservin1@epcc.edu

It would be great to eliminate all accidents, but we do not think that this will happen in the nearest future: while some cars will be self-driving, many other cars will remain driven by humans who are more error-prone, not to mention that the roads are also shared by pedestrians who do not always obey the traffic rules.

In pedestrian-car situations when an accident is inevitable and someone will get hurt, the automatic system often has a choice. For example, it can swerve into the wall thus injuring the car and the driver but leaving the pedestrian intact, or it can try to make the driver safer but risking hitting and injuring the pedestrian. In human-driven cars, this decision is left to the driver (and the car manufacturers try to minimize the impact both on the driver and on the pedestrian by designing safer cars and cars with bumpers that bring the least harm to pedestrians). However, in a self-driving car, this decision needs to be made by the automatic system; see, e.g., [1] and references therein.

How to allocate safety benefits: a practical problem. What this system will do depends on how we program it. So what decision should we program?

What we proposed earlier. In our previous paper [1], we proposed to follow social norms when designing the system.

What we would prefer to do. Instead of relying on social norms—which do not always perfectly describe what is the best for the society—it is desirable to select such programs for which the overall harm to the society—i.e., the overall harm to pedestrians and to drivers—is minimized.

What we do in this paper. In this paper, we propose a solution to this problem.

The structure of the paper. We start, in Sect. 2, with a simple seemingly adequate model of the situation. In Sect. 3, we explain why this—as it turned out, oversimplified—model is not fully adequate for this problem, and what else needs to be taken into account to come up with an adequate model. In Sect. 4, we describe our final model. In Sect. 5, we show that this model leads to a simple straightforward algorithm for the optimal allocation of safety benefits.

2 First-Approximation Model

What causes pedestrian-car accidents? Our goal is to minimize the overall expected harm caused by pedestrian-car accidents. To gauge the related harm, let us first recall what causes these accidents.

- Some accidents are caused by drunk drivers—this problem will be completely eliminated if we switch to self-driving cars.
- Some accidents are caused by an unexpected event: a person absent-mindedly walks into a street, is accidentally pushed into the traffic area, etc. Such accidents are relatively rare. Let us denote by n_a the proportion of such accidents per single street crossing.

- The vast majority of car incidents involving pedestrians occur when a pedestrian tries to cross a street at the wrong place—or at red light. The ability of cars to self-drive will, unfortunately, not cause pedestrians to be more disciplined—so this needs to be taken into account.

How to gauge the average harm caused by pedestrian-car accidents: current situation. Let us estimate the average harm per one street crossing. In view of the above, to estimate this average harm, we need to know:

- what is the proportion of street crossings that are not following the rules; let us denote this proportion by n ;
- what is the proportion of not-following-the-rules crossings that result in an accident; let us denote this proportion by a ; and
- what is the average harm done to the pedestrian and to the driver; let us denote these harms by, correspondingly, h_p and h_d .

In these terms, the average harm per street crossing is equal to the product $n \cdot a \cdot (h_p + h_d)$.

What will happen when we use self-driving cars. When we introduce effective self-driving cars, the following will happen:

- some proportion p_+ of accidents will be avoided completely;
- some proportion p_- of accidents we will still not be able to avoid at all; and
- the remaining proportion $p_0 = 1 - p_+ - p_-$ is when we can decrease the harm, but not fully eliminate it.

It is in this last category, in which we cannot fully eliminate the harm, that we need to decide how to fairly distribute the benefits of increase safety.

Let $b_p \in [0, 1]$ indicate the proportion of benefits that goes to the pedestrian, then $b_d = 1 - b_p$ will be the proportion of benefits that goes to the driver. In these terms:

- the expected harm to the pedestrian will be equal to

$$(n_a + n) \cdot a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p)), \text{ and}$$

- the expected harm to the driver will be equal to

$$(n_a + n) \cdot a \cdot h_d \cdot (p_- + p_0 \cdot (1 - b_d)).$$

Thus, the overall expected harm is equal to the sum of these harms:

$$\begin{aligned} & (n_a + n) \cdot a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p)) \\ & + (n_a + n) \cdot a \cdot h_d \cdot (p_- + p_0 \cdot (1 - b_d)). \end{aligned} \tag{1}$$

3 Limitations of This Model

What does this model recommend? At first glance, our resulting formula (1) provides a reasonable description of the average harm. So, all we need to do now is to find the value b_p that minimizes the average harm.

To find this optimal value, let us plug in the expression $b_d = 1 - b_p$ into this formula. In this case, $1 - b_d = b_p$ and thus, we get the following expression:

$$(n_a + n) \cdot a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p)) + (n_a + n) \cdot a \cdot h_d \cdot (p_- + p_0 \cdot b_p). \quad (2)$$

We get a linear expression in terms of the unknown b_p . Separating terms depending and not depending on b_p in this expression, we can transform the expression (2) into the following form:

$$\text{const} + b_p \cdot (n_a + n) \cdot a \cdot p_0 \cdot (h_d - h_p), \quad (3)$$

where by const, we mean terms that do not depend on b_p at all.

In a pedestrian-car accident, the harm to the pedestrian is usually much larger than the harm to the driver: $h_p \gg h_d$. Thus, to minimize the expression (3) that describes the overall harm, we need to maximize the benefits b_p allocated to the pedestrian, i.e., to select $b_p = 1$.

What is wrong with this recommendation? What will happen if we follow this recommendation? Let us consider the extreme case when $p_- = 0$, i.e., when the effect of *all* accidents will be decreased. In this case, if we select $b_p = 1$, this means that even when a pedestrian crosses the street in the wrong place, he/she will not be harmed: the only person who may be harmed will be the driver.

In our cities, we already have a lot of people jaywalking—even though sometimes this causes them harm. If we decrease the risk to pedestrians, what will prevent even more people to start jaywalking and thus, increasing the number of accidents even more?

This is clearly not what we expected.

So what was wrong with the model? Since the recommendations provided by our model are not fully adequate, this means that something in this model was not adequate.

The above reasoning explains what was wrong—we implicitly assumed that the proportion n of people who do not follow the rules when crossing the street is a constant. In reality, as the expected harm decreases, this proportion will increase. To make our model more adequate, we need to take this increase into account.

4 Final Model

How can we determine the proportion n of pedestrians who do not follow the rules? To answer this question, let us recall why pedestrians sometimes do not follow the rules: because if they do not wait for the green light and/or do not walk to the official crossing, they gain some time. Let us denote the average gain per each crossing by g .

Each pedestrian has a choice:

- either to follow the rules and thus, lose the potential gain amount g ,
- or not to follow the rules and thus, face the loss $n \cdot a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p))$.

If the second loss is smaller, more pedestrians will select not following the rules and thus, the proportion n will increase—until these numbers even out. Similarly, if the first loss is smaller, more pedestrians will decide to follow the rules, and thus, the proportion n will decrease—until these numbers even out. Thus, eventually, these losses will become equal:

$$g = n \cdot a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p)).$$

Based on this equality, we can find the desired proportion n :

$$n = \frac{g}{a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p))}. \quad (4)$$

Our final model. Substituting the expression (4) into the formula (2), we get our final model, that describes the expected average harm corresponding to the given value b_p :

$$\left(n_a + \frac{g}{a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p))} \right) \cdot H, \text{ where}$$

$$H \stackrel{\text{def}}{=} a \cdot h_p \cdot (p_- + p_0 \cdot (1 - b_p)) + a \cdot h_d \cdot (p_- + p_0 \cdot b_p). \quad (5)$$

To find the best allocation of safety benefits between pedestrians and drivers, we need to select the value b_p from the interval $[0, 1]$ that minimizes the expression (5).

Why we believe that our new model is more adequate. We dismissed our original simplified model because it led to an inadequate solution $b_p = 1$. Let us show that our new, more adequate model does not have this problem, at least when the value p_- —that describes that remaining imperfection of the self-driving cars—is sufficiently small.

Indeed, in the case of $p_- = 0$, when b_p tends to 1, the first factor in the formula (5) tends to infinity, while the second factor remains non-zero. Thus, the product (5) also tends to infinity when b_p tends to 1. So, in contrast to the above simplified model, the minimum of the expression (5) cannot be attained when $b_p = 1$.

5 How to Find the Optimal Allocation b_p

Analysis of the problem. How can we actually find the optimal allocation b_p for which the expression (5)—that describes the expected average harm—attains its smallest possible value? Good news is that we can actually find an explicit expression for this optimal value. Let us explain how to do it.

The dependence of the expression (5) on the unknown b_p has the following form:

$$\left(a_0 + \frac{a_1}{a_2 + a_3 \cdot b_p} \right) \cdot (a_4 + a_5 \cdot b_p), \quad (6)$$

for the following coefficients a_i :

$$a_0 = n_a, \quad a_1 = g, \quad a_2 = a \cdot h_p \cdot (p_- + p_0), \quad (7)$$

$$a_3 = -a \cdot h_p \cdot p_0, \quad a_4 = a \cdot h_p \cdot (p_- + p_0) + a \cdot h_d \cdot p_i, \quad a_5 = -a \cdot h_p \cdot p_0 + h_d \cdot p_0. \quad (8)$$

According to calculus, the minimum of a differentiable function inside an interval $b_p \in (0, 1)$ is attained at a point where the derivative of this function is equal to 0. Differentiating the expression (6) with respect to b_p and equating the derivative to 0, we get the following equation:

$$-\frac{a_1 \cdot a_3}{(a_2 + a_3 \cdot b_p)^2} \cdot (a_4 + a_5 \cdot b_p) + \left(a_0 + \frac{a_1}{a_2 + a_3 \cdot b_p} \right) \cdot a_5 = 0. \quad (9)$$

Multiplying both sides of this equation by $(a_2 + a_3 \cdot b_p)^2$, we get the following quadratic equation:

$$-a_1 \cdot a_3 \cdot (a_4 + a_5 \cdot b_p) + a_0 \cdot a_5 \cdot (a_2 + a_3 \cdot b_p)^2 + a_1 \cdot a_5 \cdot (a_2 + a_3 \cdot b_p) = 0,$$

i.e., the equation

$$A \cdot b_p^2 + B \cdot b_p + C = 0, \quad (10)$$

where we denoted

$$A = a_0 \cdot a_5 \cdot a_3^2, \quad (11)$$

$$B = -a_1 \cdot a_3 \cdot a_5 + 2a_0 \cdot a_5 \cdot a_2 \cdot a_3 + a_1 \cdot a_5 \cdot a_3 = 2a_0 \cdot a_5 \cdot a_2 \cdot a_3, \quad (12)$$

$$C = -a_1 \cdot a_3 \cdot a_4 + a_0 \cdot a_5 \cdot a_2^2 + a_1 \cdot a_5 \cdot a_2. \quad (13)$$

Now, we can use the general formula for solving a quadratic equation to find the optimal value of b_p :

$$b_p = \frac{-B \pm \sqrt{B^2 - 4A \cdot C}}{2A}. \quad (14)$$

Out of two possible solutions, we should select the one that is inside the interval $(0, 1)$ —or, if both solutions are inside this interval, the one for which the value of the expression (5) is the smallest.

So, we arrive at the following algorithm.

Algorithm for computing the parameter b_p that describes optimal allocation of safety benefits of self-driving cars. To compute b_p , we need to know the following values:

- the proportion a of not-following the rules crossings that currently lead to an accident;
- the proportion n_a of crossings that result in accidental violation of the rules;
- the current average per-accident harm h_p to a pedestrian;
- the current average per-accident harm h_d to a driver;
- the proportion p_+ of the accidents that will be completely avoided by the self-driving cars;
- the proportion p_- of the accidents about which the self-driving car system cannot do anything.

Based on these values, we can compute the proportion $p_0 = 1 - p_+ - p_-$ of the accidents for which the self-driving car system can decrease—but not fully eliminate—the harm. What we want to compute is the proportion b_p of the resulting safety benefits that will be allocated to pedestrians. To compute the optimal proportion b_p , we do the following:

- first, we use the formulas (7)–(8) to compute the auxiliary coefficients a_i ;
- then, we use the formulas (11)–(13) to compute the auxiliary coefficients A , B , and C ;
- finally, we use the formula (14) to compute the optimal value of b_p .

Out of two possible solutions, we should select the one that is inside the interval $(0, 1)$ —or, if both solutions are inside this interval, the one for which the value of the expression (5) is the smallest.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

Reference

1. C. Servin, V. Kreinovich, S. Shahbazova, Ethical dilemma of self-driving cars: conservative solution, in *Recent Developments and the New Directions of Research, Foundations, and Applications*, vol. 2, eds. by S.N. Shahbazova, A.M. Abbasov, V. Kreinovich, J. Kacprzyk, I. Batyrshin (Springer, Cham, Switzerland, 2023), pp. 93–98

Paradox of Causality and Paradoxes of Set Theory



Alondra Baquier, Bradley Beltran, Gabriel Miki-Silva, Olga Kosheleva,
and Vladik Kreinovich

Abstract Logical paradoxes show that human reasoning is not always fully captured by the traditional 2-valued logic, that this logic's extensions—such as multi-valued logics—are needed. Because of this, the study of paradoxes is important for research on multi-valued logics. In this paper, we focus on paradoxes of set theory. Specifically, we show their analogy with the known paradox of causality, and we use this analogy to come up with similar set-theoretic paradoxes.

1 Introduction

Paradoxes: a brief reminder. Informal intuitive reasoning often leads us to paradoxes. For example, there is a famous *heap paradox*. It is based on the following simple reasoning:

- one grain of sand does not form a heap, and
- if we have several grains of sand that do not form a heap, and we add one more grain, we will also not get a heap.

By mathematical induction, we can easily conclude that no matter how many grains of sand we combine, we will never get a heap—but in reality, it is easy to form a heap.

A. Baquier · B. Beltran · G. Miki-Silva · O. Kosheleva · V. Kreinovich (✉)
University of Texas at El Paso, El Paso, TX, USA
e-mail: vladik@utep.edu

A. Baquier
e-mail: abaquier1@miners.utep.edu

B. Beltran
e-mail: bjbeltran@miners.utep.edu

G. Miki-Silva
e-mail: gmikisilva@miners.utep.edu

O. Kosheleva
e-mail: olgak@utep.edu

Paradoxes and multi-valued logics. Logical paradoxes show that human reasoning is not always fully captured by the traditional 2-valued logic, that this logic's extensions—such as multi-valued logics (an example of which is fuzzy logic; see, e.g., [1–6])—are needed.

For example, in the heap case, a usual solution to this paradox is that “being a heap” is not a crisp 2-valued notion. Yes, one grain of sand is definitely not a heap, and big heaps are definitely heaps, but on the intermediate stage, we have a configuration about which we are not certain. By adding a single grain of sand, we never jump from “definitely not a heap” to “definitely a heap”. Instead, we slightly increase our degree of confidence that the resulting collection of grains of sand is a heap.

It is therefore important to study paradoxes. Because of this relation between paradoxes and multi-valued logics, the study of paradoxes is important for research on multi-valued logics.

What we do in this paper. In this paper, we focus on paradoxes of set theory. Specifically:

- we show their analogy with the known paradox of causality, and
- we use this analogy to come up with modified paradoxes of set theory.

The structure of this paper. In Sect. 2, we recall the paradox of causality. In Sect. 3, we recall the most well-known paradox of set theory—Russell's paradox. In Sect. 4, we describe the relation between these paradoxes. Finally, in Sect. 5, we use this relation to come up with modified versions of Russell's paradox.

2 Paradox of Causality: A Brief Reminder

Paradox of causality: in brief. Anyone who watched a movie about time travel or read a book about it knows that time travel leads to paradoxes. By using a time machine, a person can go into the past and kill his/her grandfather before the time traveler's father was conceived. Then:

- the traveler's father could not have been born and thus, the traveler him/herself could not have been born,
- however, the time traveler *was* actually born, he/she is here.

This is clearly a paradox.

Let us describe this paradox in more precise terms. This paradox describes the situation when:

- the event e_1 – the time traveler making a decision to use the time machine – affects the event e_2 – life of his/her grandfather, and
- at the same time the event e_2 clearly affects the existence of the time traveler and thus, affects the event e_1 .

Let us denote the relation “event e_1 affects the event e_2 ” by $e_1 < e_2$. In this notation, we get $e_1 < e_2$ and $e_2 < e_1$. Here we have a loop $e_1 < e_2 < e_1$ of size 2:

- the event e_1 affects the event e_2 , and
- the event e_2 affects the event e_1 .

We can consider longer loops. A similar paradox appears also when we have a longer loop:

- the event e_1 affects the event e_2 ,
- the event e_2 affects some other event e_3 – for example, the birth of the time traveler’s father, and
- the event e_3 affects the event e_1 .

So, here, $e_1 < e_2 < e_3 < e_1$.

We can have even longer loops.

3 Paradox of Intuitive (“naive”) Set Theory

What this paradox is about. In set theory, there is a similar paradox—discovered by the famous philosopher Bertrand Russell – that is related to the loop of size 1, namely, to the possibility that a set x may be an element of itself: $x \in x$. Examples of such sets are easy to find: for example:

- the set U of all sets is a set and
- thus, this set is its own element: $U \in U$.

What exactly is the paradox. The paradox appears if we consider the set

$$S = \{x : x \notin x\}$$

of all the sets that *do not* belong to itself.

Specifically, the paradox appears when we check whether this set S belongs to itself.

Why is this a paradox. Indeed, we have only two possible options:

- either the set S belongs to itself: $S \in S$,
- or the S does not belong to itself $S \notin S$.

Let us show that in both cases, we get a contradiction.

What if $S \in S$? Let us assume that the set S belongs to itself. By definition, the set S includes all the sets x that do not belong to themselves.

- Since the set S itself is an element of S , it means that the set S *does not* belong to itself.
- However, we assumed that it *does* belong to itself.

So, in this case, we have a contradiction.

What if $S \notin S$? Let us now assume that the set S does not belong to itself. This means that the set S does not have the property that describes all elements x of the set S – that each of these elements x does not belong to itself. Since the set S *does not* have this property, this means that the set S *does* belong to itself. However, we assume that it *does not* belong to itself.

So, in this case, we also have a contradiction.

So, we have a paradox. In both possible cases – when $S \in DS$ and when $S \notin S$ – we get a contradiction – which means that none of these two cases is possible. So, we indeed have a paradox.

4 How Are These Paradoxes Related?

How are these paradoxes similar. In both cases – for time travel and for sets – we have a paradox. Both paradoxes relate to closed loops.

How are these paradoxes different. The main difference between these two paradoxes is as follows:

- in the case of time travel, the paradox appears for loops of any size, while
- for sets, the paradox is only known for a loop consisting of a single element:

$$x \in x.$$

A natural question. So, a natural question is: what if in the set case, we have a longer loop, will we still get a paradox?

Our answer. Our answer to this question is: yes, it is possible to have versions of Russell’s paradox that include loops of any size. We will show the details in the next section.

5 Resulting Modifications of Russell’s Paradox

What is the desired paradox. In this section, we show that yes, in set theory we can get a similar paradox if we consider loops of any length.

In other words, we get a paradox if we consider the possibility that:

- a set x is an element of a set x_1 ,
- the set x_1 is an element of set x_2 , and so on,
- and at the end, the set x_n is an element of the set x with which we started.

To show that there is a paradox let us consider a new set S_n : the set of all elements x for which such a loop does not exist:

$$S_n = \{x : \neg \exists x_1, \dots, x_n (x \in x_1 \in x_2 \in \dots \in x_n \in x)\}.$$

Just like before, let us check whether this set S_n belongs to itself.

Why is this a paradox. Of course, the set S_n either belongs to itself or it does not. Let us prove that in both cases, we get a contradiction.

What if $S_n \in S_n$? Let us first assume that the set S_n belongs to itself: $S_n \in S_n$. By definition of the set S_n , this means that we *cannot* have a loop in which S_n belongs to some set x_1 , x_1 belongs to some set x_2 , and so on, until we get a set x_n that belongs to S_n :

$$\neg \exists x_1, \dots, x_n (S_n \in x_1 \in x_2 \in \dots \in x_n \in S_n).$$

But we *do* have such a loop – we can take x_1, x_2 , and so on all equal to S_n , then we have

$$S_n \in x_1 = S_n \in x_2 = S_n \in \dots \in x_n = S_n \in S_n.$$

So, in this case, we get a contradiction.

What if $S_n \notin S_n$? On the other hand, let us consider the case when S_n is *not* an element of itself, i.e., when $S_n \notin S_n$. This means that the set S_n does not have the property that defines this set – that no loop exists containing this set. The fact that this property is not satisfied means that for S_n such a loop does exist. In other words, there exist sets $x_1, \dots, x_n$ such that:

- the set S_n belongs to x_1 ,
- the set x_1 belongs to x_2 , and so on, and
- the set x_n belongs to S_n :

$$S_n \in x_1 \in x_2 \in \dots \in x_n \in S_n.$$

Here:

- Since the set x_n belongs to S_n , this means that for this x_n , the property describing the set S_n is true, i.e., for the set x_n , *such a loop is not possible*.
- However, *we can build such a loop*: indeed, x_n belongs to S_n , S_n belongs to x_1 , and so on, and at the end, x_{n-1} belongs to x_n :

$$x_n \in S_n \in x_1 \in x_2 \in \dots \in x_{n-1} \in x_n.$$

So, in this case, we also get a contradiction.

So, we have a paradox. In both possible cases – when $S_n \in S_n$ and when $S_n \notin S_n$ – we get a contradiction. This means that none of these two cases is possible – so, we indeed have a paradox.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

References

1. R. Belohlavek, J.W. Dauben, G.J. Klir, *Fuzzy Logic and Mathematics: A Historical Perspective* (Oxford University Press, New York, 2017)
2. G. Klir, B. Yuan, *Fuzzy Sets and Fuzzy Logic* (Prentice Hall, Upper Saddle River, New Jersey, 1995)
3. J.M. Mendel, *Explainable Uncertain Rule-Based Fuzzy Systems* (Springer, Cham, Switzerland, 2024)
4. H.T. Nguyen, C.L. Walker, E.A. Walker, *A First Course in Fuzzy Logic* (Chapman and Hall/CRC, Boca Raton, Florida, 2019)
5. V. Novák, I. Perfilieva, J. Močkoř, *Mathematical Principles of Fuzzy Logic* (Kluwer, Boston, Dordrecht, 1999)
6. L.A. Zadeh, Fuzzy sets. *Inf. Control.* **8**, 338–353 (1965)

Fuzzy Recognition System

Continuous Growth of Information and the “Forgetting” Function in the Recognition of Sound Information



Shahnaz N. Shahbazova 

Abstract In general, classes of visual objects are being studied, which have not yet received a stable solution, and particular methods need artificial adjustments with subsequent difficulties in integrating into the general and quite harmonious mode of operation of the system. In the paper, the “forgetting” function being developed has three or four implementation methods, and the problems being solved allow us to look from the inside (drawing analogies from the software point of view) at the ways in which human intelligence resorts to solve this problem in order to build similar abilities.

Keywords Forgetting function • Recognition • Motor component • Recognition system • Flat shape • Disassemble and recognize • Functioning associative memory • Various drawn geometric shapes

1 Introduction

Problem statement visual and audio information are the only currently available intelligent means of delivering information to the system therefore, effective research on this relatively simple intellectual task can lay the foundation for further deepening experimental and analytical research on the topic of artificial intelligence development.

Visual and audio information are the only currently available intelligent means of delivering information in a computer system, therefore, effective research on this relatively simple intellectual task can lay the foundation for further deepening experimental and analytical research on the topic of artificial intelligence development [1].

S. N. Shahbazova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, UNEC, Baku, Azerbaijan

e-mail: shahbazova.shahnaz@unec.edu.az

In the case of visual information recognition, this work does not aim to create a special recognition system for any types of images, or implement a specific application program, but is a model of a computer baby who does not know how to think and does not understand the meaning of objects, but is able to see (disassemble and recognize).

Accordingly, as in the case of identification of visual information, it is necessary to clarify that this work does not aim to create a speech recognition system, but considers the general flow of audio information (containing mainly speech) also based on the model of a computer baby who does not understand the meaning, but hears and is able to recognize sound objects [2].

2 Mechanism for Linking Information Objects

The research does not intend to develop any kind of robotic mechanism, since the research is aimed solely at algorithmizing certain primitive functions of natural intelligence related to the recognition of visual and sound objects of environmental information [3].

The target area of research is the motor component of human intelligence, i.e. the functions of understanding visual and audio information on which a person spends less than 0.5–2 s without using the functions of cognition, reasoning and logic [4].

The next issue requiring a special solution is the creation of a mechanism for linking information objects with each other, which does not impose any restrictions on the structure of the organization of the library of objects of natural origin, and at the same time allows developing additional intellectual add-ons on its basis in the future. Looking ahead, we note that all these requirements were implemented by the method of labeling with dictionary concepts. At the learning stage, visual and sound objects are marked with dictionary concepts that carry out associative communication (see Fig. 1) “images and audio” [5].

Accordingly, the field of research includes the task of creating a functioning associative memory, without which it is impossible to imagine the functions of identifying information objects, as well as the learning process, not to mention the higher intellectual properties of natural intelligence [6].



Fig. 1 The mechanism of vocabulary concepts is the relationship between visual and audio information

In parallel, the task of investigating the classification of information objects is set in the conditions of the natural learning process (without artificial design of the classroom space).

3 Continuous Growth of Information

The currently used mechanism for accumulating information samples by the forced method of learning is a forced measure that allows you to solve the problem of lack of available hardware, and is the main limitation in the development and complication of the system as a whole, for example, to switch to self-learning mode [7].

In connection with this property, in order to reduce the number of instances of objects of stored objects, the possibilities of solving the problem in three directions were provided [8, 9]:

- The method of reducing “similar” information objects is a method of periodically viewing databases of all or most services in order to search for “similar” objects (currently not in use).
- The purpose of such cleaning is to try to leave the most characteristic samples of objects in databases, but so far, no automated method provides a reliable way of such classification [8, 10].
- The method of simplified memorization is to allocate space for the image, firstly, only when there is an available place, and secondly, after checking for the absence of this instance in a certain class of objects, thirdly, it is necessary that it differs sufficiently from those already recorded. In a sense, a similar method may be used in human memorization. For example, older people remember events and details from their youth in detail, and at the same time they hardly remember the details of recent events, and as a result of filtering out most new events in their memory, it feels like time is running much faster than in their youth [11].
- The method of memorization with learning—conducting ordinary memorization (with by introducing small random distortions into the object by different services) only if the training command is specified. This method is the main one when conducting the learning process for new information objects.

4 The Function of “Forgetting”

There is no doubt that the harmonious development of the system is not possible without limiting factors, but in the matter of the proper functioning of the “forgetting” function, even a small mistake at the beginning of the development of the system will have serious negative consequences, which may eventually require repeating the training from the beginning [12, 13].

A feature of natural intelligence is a difficult-to-explain effect: a person does not remember objects that he already knows (met, learned, etc.), but once he sees an object a second time, in most cases he recognizes it, despite the fact that before that he was not able to clearly remember it or reproduce it [14–16].

If we assume that the capacitive resources of intelligence are quite limited, there must be some kind of intellectual way to isolate individual characteristic features of information objects in the form of some kind of impression, according to which it is difficult to reproduce the object itself as a whole, but belonging to the original object is quite reliably determined [17].

The method of focal lengths used in the work for visual objects (frequency seals for sound objects) gives good recognition indicators after obtaining a sufficiently large number of focal lengths of random areas of the visual image (frequency characteristics of the sound image), which represent significant amounts of information [1, 18].

In parallel, the possibility of the existence of some natural hashing algorithm is being investigated. In mathematical cryptography, there are algorithms for one-way transformations of information—hashing (hash), generating keys using such an algorithm that guarantees an extremely low probability of detecting other information of a similar type with the same hash value [19, 20].

In our case, most likely, there is a calculation of some kind of compressed hash image, which is not as strict as the mathematical one, but quite economical and which, perceiving full-fledged information from the organs of vision (and hearing), highlights random details and performs a transformation, allowing you to remember only a negligibly small amount of information characterizing significant visual or sound images [21].

It is quite possible that this is how natural memory works, of course, there are exceptions—rare people have a special memory (photographic, musical, etc.), but in general, for this part of the work, the most important thing is to find the most economical representation of information objects with minimal loss of characteristic information. Work in this direction has not yet been completed, but according to preliminary results, the specificity, as well as in recognition, is determined by the focal map of the visual object of information (the characteristic of the spot of frequency densities of sound images) [2, 18].

It should be noted that the self-evident algorithm of simply reducing the number of random measurements leads to a sharp drop in recognition reliability, which forces us to look for an alternative mechanism for hashing images. Currently, there are still many difficulties with the color gradations of visual information, as well as the emotional “coloring” of sound information. Although many human memories contain a color component (in both types of information), in all attempts to use this information in software algorithms, only they reduce the quality of decisions made, since, in fact, from the point of view of informativeness, “coloring” is a kind of noise that increases the dispersion of reliability and, accordingly, reduces the stability of recognition [22].

Perhaps a person remembers everything in neutral colors (grayscale and neutral tone), and the color scheme is superimposed on top of images at the highest level of intelligence. This assumption is very doubtful and has not been found (still) its experimental confirmation, but algorithmic expediency points to this as a completely optimal and effective approach.

In any case, the quality of recognition, the method of storage and the system of organizing the library of information objects are closely interrelated even when making any changes, which helps a lot in research, since the results of the system's functioning are a means of calibration, debugging and testing [23, 24].

5 Conclusion

In conclusion, we note that, in general, classes of visual objects are being studied that have not yet received a stable solution, and specific methods need artificial adjustments with subsequent difficulties in integrating into the general and quite harmonious mode of operation of the system. The article develops the “forgetting” function, which has three or four ways of implementation, and the problems being solved allow us to look inside at the ways in which human intelligence resorts to solving this problem in order to develop similar abilities.

References

1. T. Vetter, T. Poggio, Linear object classes and image synthesis from a single example image. *IEEE Trans. Pattern Anal. Mach. Intell.* **19**, 733–742 (1997)
2. Wilson I.D. Foundations of hierarchical control. “*International Journal of Control*” 29, N6, 1979, p.899–933.
3. A. Marjani, M. Babanezhad, and S. Shirazian, Application of adaptive network-based fuzzy inference system (ANFIS) in the numerical investigation of Cu/water nanofluid convective flow. *Case Stud. Therm. Eng.* **22** (November), 100793 (2020)
4. R.C. Gonzalez, P. Wintz, *Digital Image Processing* (Addison–Wesley. Reading, Massachusetts, 1987), pp. 505
5. J. Kralik, P. Stiegler, Z. Vostry, J. Zavorka, Modelovani dynamiky rozsahlych siti. Praha, Akademia 364 (1984). by neural networks with single-layer training. *IEEE Trans. Neural Netw.* **3**, 962–968
6. Y. Lee, Handwritten digit recognition using K nearest-neighbor, radial-basis function, and back-propagation neural networks. *Neural Comput.* **3**, 440–449 (1991)
7. D.J. Ketcham, Real time image enhancement technique, in *Proceedings SPIE/OSA Conference on Image Processing*, vol. 74 (Pacific Grove, California, 1976), pp. 120–125
8. S. Gutta, H. Wechsler, Face recognition using hybrid classifiers. *Pattern Recogn.* **30**, 539–553 (1997)
9. A. Krzyzak, W. Dai, C.Y. Suen, Unconstrained handwritten character classification using modified backpropagation model, in *Proc. 1st Int. Workshop on Frontiers in Handwriting Recognition* (Montreal, Canada, 1990), pp. 155–166

10. Y. LeCun, O. Matan, B. Boser, J.S. Denker, D. Henderson, R.E. Howard, W. Hubbard, L.D. Jackel, H.S. Baird, Handwritten zipcode recognition with multilayer networks, in *Proc. of International Conference on Pattern Recognition* (Atlantic City, 1990)
11. M. Mariton, M. Drouin, H. Abou-Kandil, G. Duc, Une nouvelle method de decomposition-coordination. 3 e partie: application a la commande coordonnees-hierarchisee des procesus complexes. *APII*, N3, pp. 243–259 (1985)
12. M. Petrou, Learning in pattern recognition. *Lect. Notes Artif. Intell. Mach. Learn. Data Min. Pattern Recognit.* 1–12 (1999)
13. S.A. Sharov, Tools for computer representation of linguistic information (1996), <http://www.rriai.org.ru/~sharoff/sharoff@mx.iki.rssi.ru>
14. W. Findeisen, K. Malinowski, Two-level control and coordination for dynamical systems. *Arctium automatiki i telemechaniki. T. XXIV*, N1, pp. 3–27
15. R. Foltyniewicz, Efficient high order neural network for rotation, translation and distance invariant recognition of gray scale images. *Lect. Notes Comput. Sci. Comput. Anal. Images Patterns*, pp. 424–431 (1995)
16. W. Frei, Image enhancement by histogram hyperbolization. *Comput. Graph. Image Process.* **6**(3), 286–294 (1977)
17. E.L. Hall, Almost uniform distribution for computer image enhancement. *IEEE Trans. Computers.* **23**(2), 207–208 (1974)
18. W.W. Translation, T. Report, *Reprinted in Machine Translation of Languages, 1955* (MIT Press, Cambridge, MA, 1949), pp. 15–23
19. S.G. Tzafestas, Large-scale systems modeling in distributed-parameter control and estimation, in *Modeling and Simul. Eng. 10th IMACS World Congr. Syst. Simul. and Sci. Comput., Montreal, 8–13 Aug. 1982, v.3* (Amsterdam e.a., 1983), pp. 69–77
20. D. Valentin, H. Abdi, A.J. O'Toole, G.W. Cottrell, Connectionist models of face processing: a survey. *Pattern Recognit.* **27**, 1209–1230 (1994)
21. T. Winograd, *Language as Cognitive Process*, vol. 1 (Syntax. Addison, Wesley, 1983)
22. K.S. Yoon, Y.K. Ham, R.-H. Park, Hybrid approaches to frontal view face recognition using the hidden Markov model and neural network. *Pattern Recogn.* **31**, 283–293 (1998)
23. H. Michalska, J.E. Ellis, P.D. Roberts, Joint coordination method for the steady-state control of large-scale systems. *Int. J. Syst. Sci.* **N5**, 605–618 (1985)
24. M. Milanova, P.E.M. Almeida, J. Okamoto, M.G. Simoes, Applications of cellular neural networks for shape from shading problem. *Lect. Notes Artif. Intell. Mach. Learn. Data Min. Pattern Recognit.* 51–63 (1999)

Technology Building a Virtual Room for Recognition of Visual and Sound Information



Shahnaz N. Shahbazova 

Abstract The paper provides a way to solve the problem of separating information objects from the overall picture and relative to each other, using the method of local foci for visual information and the method of frequency densities for audio information. The rationale given boils down to the fact that the applied one is quite capable of making out a visual scene and a sound fragment of high complexity, as well as recognizing their constituent objects. When using hardware with sufficient computing power, the algorithmic parallelization method will make it possible to bring real computing time to virtual.

Keywords Recognition · Complex of type · Artificial personality · Two-dimensional model · Imitation of intellectual behavior · Artificial intelligence model

1 Introduction

In this paper, the problem of recognizing visual and audio information of the natural environment is investigated by constructing an algorithmic model that simulates the human abilities of motor (unconscious or non-intellectual) pattern recognition.

Although the vast majority of modern developments in the field of information recognition are closed commercial projects and their theoretical achievements are practically not disclosed, the entire field of research is in decline, and this, given that recent decades have been characterized by the rapid development of information technology in general, indicates its theoretical and practical impasse.

Almost all more or less successful projects have only confirmed the final scope of applicability and limited reliability of recognition results. The long absence of significant mathematical progress raises the question of the need for a radical revision

S. N. Shahbazova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, UNEC, Baku City, Azerbaijan

e-mail: shahbazova.shahnaz@unec.edu.az; shahbazova@gmail.com

of the directions of research on the problems of recognition of visual and audio information, in order to find alternative ways to solve this problem [17].

As has often happened in the history of science, if theoretical research comes to a dead end when solving a certain problem, then experimental research comes forward, conducting hundreds and thousands of different experiments, groping for a thread that gives correct results, which are further justified and described mathematically.

Practice has shown that it is extremely difficult or practically impossible to mathematically describe the process of information recognition since by its nature. It is more an algorithm of a program than mathematical transformations. This observation was the basis for this research work [7].

2 Behavioral Model of an Artificial Personality (AP)

The research, which this work is a part of, is generally devoted to the development of a machine intelligence system with a gradual increase in the complexity of the behavioral model of an artificial personality (AP) in order to experimentally study the issues of artificial intelligence.

The development of computer technology, as well as the speed of development and integration into all spheres of human society, has made it possible to qualitatively algorithmize almost all areas of human life, except for intellectual functions, the research of which is far from complete.

Most studies of human intellectual functions are conducted either in the narrow commercial interests of specific developments or are at too high of a level of research on the human brain or the creation of universal theories of cognition, memory, etc. For example, a number of projects have achieved good results in the direction of robotics and imitation of intellectual behavior (mainly by a system of conditional templates).

At a certain stage in the development of such projects, developers are forced to use various optimization methods (templates, structurally fixed dictionaries, restriction of features, properties of objects, etc.) to increase the effective performance of the system and improve the quality of decisions (in particular recognition) [12].

Such optimization goes against the algorithms of wildlife, which is quite acceptable and effective for solving applied problems, but is not suitable for the development of intelligent systems. Unfortunately, most of the existing theories of recognition of visual and audio information have been subjected to the same trend. From a biological point of view, brain function, information extraction and processing (including decision-making) are almost inseparable concepts, but all attempts to implement such a system have yet to achieve success. Thus, solving the problem of extracting information from the environment will bring us one step closer to obtaining a full-fledged artificial intelligence model [13].

These studies are aimed at obtaining an algorithmic model capable of extracting “significant” objects from the surrounding world in order to link them with

the corresponding vocabulary concepts, which are atomic (indivisible) bricks of intelligence.

The question of the basic language for describing dictionary concepts is not fundamental, it is enough that the concepts could be described in the target language: object, property, space, time and action, i.e. almost all languages existing in the world. In the future, after obtaining a reliable way to convert “meaningful” (visual and audio) information into words and concepts, scientists and researchers will be able to focus on converting vocabulary structures into some simulation models, systems and functions of memory, cognition, logic and decision-making [4].

3 Technology for Building a Virtual Room

When designing the virtual environment, separate modules of various software products from leading computer industry companies were used:

- Electronic Arts;
- Microsoft;
- Corel;
- AutoDesk;
- and some others.

The program codes of the modules, in which 3D and 2.5D projections of any scenes and objects were used [18].

The standard construction methods that were used have proven themselves well when creating many 2 and 3-dimensional objects. It is assumed that the reader has an understanding of linear algebra, computer graphics, game programming and 3-dimensional scenes.

In [20], along with the description of standard techniques for creating game programs, the internet addresses of additional information resources of a similar subject were published. By having a team of programmers, it is quite possible to create a virtual room system from scratch.

The computing needs of the virtual environment depend on the details with which the scene is rendered, as well as giving it additional realism by adding additional visual noise. In general, the hardware of a modern computer is based on a Pentium 4 processor with sufficient RAM (depends on the size of the room about 1 MB per 1m² of virtual volume) and a good performance video card (preferably 128 + Mbytes of video memory) [3].

This computer simultaneously performs the functions of the researcher's workplace (Fig. 1.), since it should allow the researcher to observe the progress of the experiment from the outside (Fig. 2).

The internal communication protocol is a simple exchange of code sequences of commands and text messages in UTF-8 encoding, with a TCP/IP connection.

There are 4 types of codes:

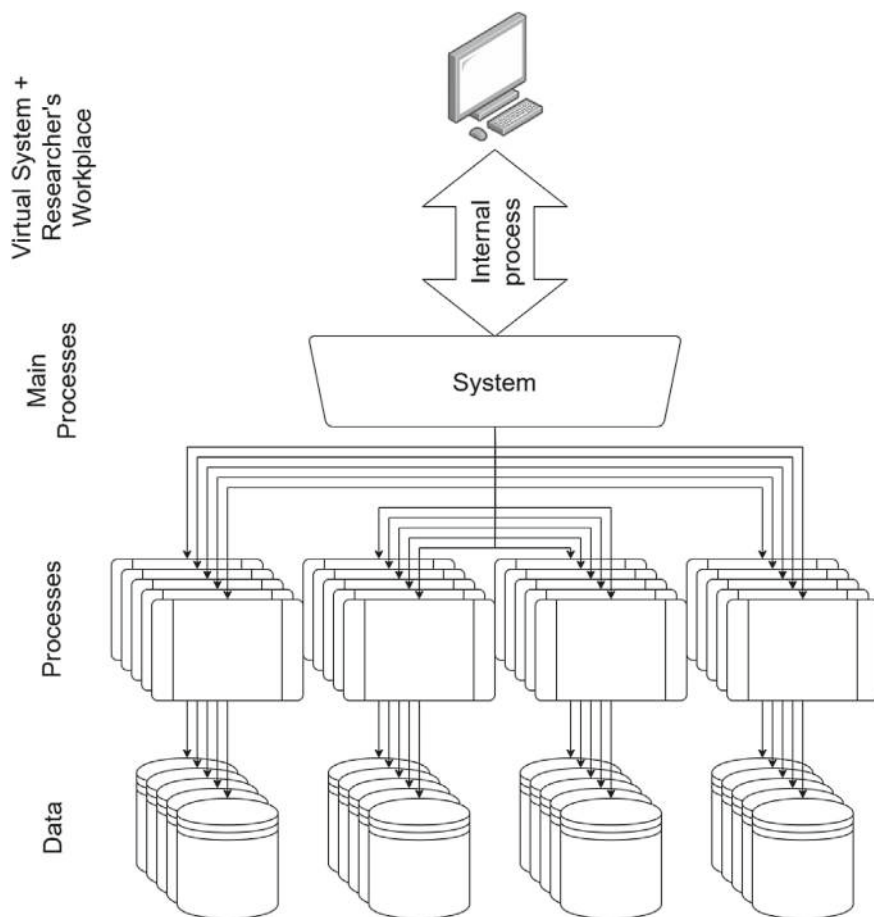


Fig. 1 General structure of the recognition system

- starting—indicating the beginning and the end (including forced) [16];
- tracing—providing information that can be used to monitor the operation of the system, including those that make up the experiment report;
- signaling—performing signaling functions generated at the beginning of object detection, at the beginning of movement to the target object, at the beginning of recognition, etc. [15];
- design—using these commands, the internal content of the virtual room is built before the experiment, including AP, these are the coordinates of the location of objects and their orientation in space.

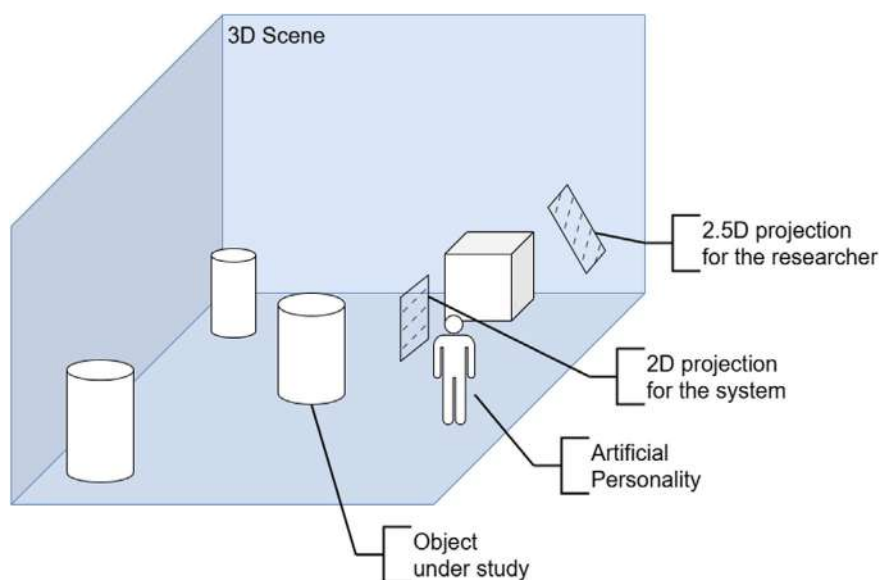


Fig. 2 Functional structure of the recognition system

4 Experimental Results

The results of several experiments with different conditions conducted under different internal states of the system (virtual room objects and sound information) are presented [1].

The block diagram (Fig. 3) shows the stages that describe the logic of the “behavior” of the IL when searching for and recognizing objects in the virtual room. This algorithm is an implementation of simple logic, which is to first find an unrecognized object by rotating around its axis, then move to the target object and begin the recognition procedure. Then, upon completion of recognition, request a sound fragment (if it is attached to an object) and proceed to recognition of the sound object. After this, continue the experiment by returning to the rotation stage or end the experiment [10].

The extreme simplification of the experimental procedure has a number of disadvantages, which are compensated by ease of implementation, clarity, a certain naturalness, power and flexibility - qualities that at this stage of work are of priority importance.

In general, improving the logic of studying the available space and the objects placed there, including the concepts of purposefulness, elements of curiosity, methods of preliminary review, etc., will be required only in the later stages of project development, when it becomes possible to come close to simulating behavioral models of natural intelligence [9].

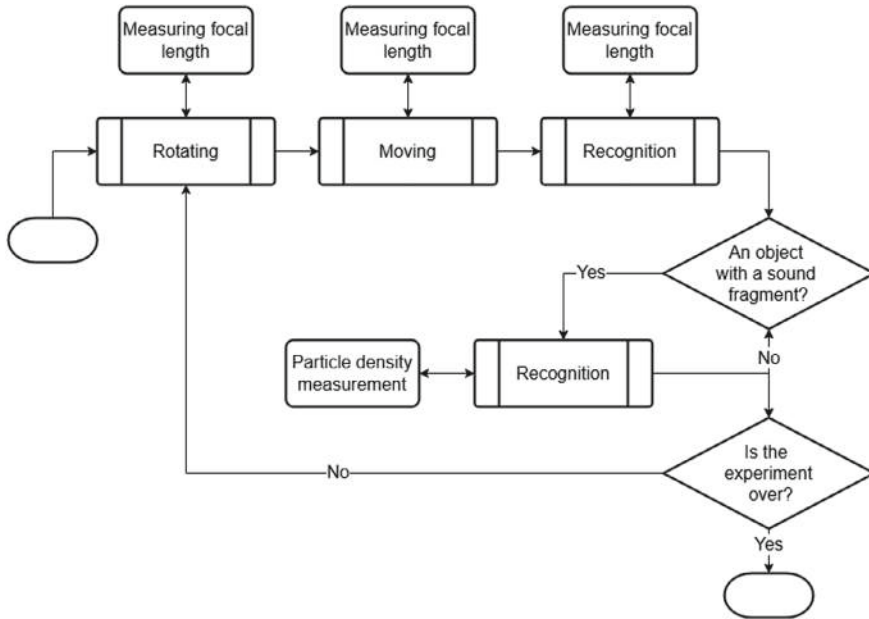


Fig. 3 Simplified functional block diagram conducting experiments

The speed of conducting experiments is greatly influenced by the computing performance of the available hardware, so when conducting experiments, several restrictions were introduced [2]:

- In all experiments the number of objects placed in the virtual room did not exceed 7 pieces;
- If a flat image was used in the experiment, then the number of visual objects depicted on it was summed up with 3-dimensional ones, so that their total number did not exceed 7 pieces;
- The maximum, but not the maximum, time for conducting an experiment is considered to be 10 days, which is quite enough time to conduct a comprehensive experiment (taking into account the available computing resources), while simultaneously developing additional functions and modules of the system that need to be tested for quite possible situations of “intelligent” errors and loops, as well as for common software errors;
- The Z coordinate (height) in all experiments is equal to 0, thus, all objects are
- Located either on the virtual floor or on other objects;
- When creating sound fragments participating in experiments, only one source with speech information (with background or pure) can sound at a time; there can be several noise and background sources.

When conducting experiments, the system keeps records (Table 1) about the events occurring in the virtual room, which further serves to reconstruct a picture of the

history of processes and reactions, which serve as the main tool in the hands of the researcher for testing and monitoring the functioning of various software modules [8] (Table 2).

As you can see, records are kept based on the fact that occurred, if the AP moves from coordinate (3428, 5655) to coordinate (3435, 5675), then the very fact of movement is recorded. Thus, records of events in a virtual room have limited accuracy [5, 6].

For example, when moving from one point to another, an AP can walk around the room 100 times and this will not be reflected in the reporting records. Of course, this is an exaggerated example, but quite possible (regarded as an algorithm error), since the time, coordinates and rotation angle are mainly recorded in the report when checking the current focal length. Due to the fact that the focal length is calculated quite often, the report on the experiment is therefore quite detailed, and algorithmic errors are quite easily recognized [11, 14].

To restore the progress of work (the behavior of the AP) of a particular experiment, a trace of the full report is required. To demonstrate the real results of the recognition algorithm, only the report of one (the smallest) of the four experiments is presented in full, and the rest are presented in brief form (only the main points of the experiments) [19].

Table 1 Summary report template

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	Recognized visual or audio object

Table 2 Summary report template

Real time	Virtual time	Visual object	Sound object	Credibility

5 Conclusion

The work builds an original algorithm for the allocation and stratification of visual and audio information into objects of the real environment, on an unconventional basis. A method of storing and comparing (and, as a result, searching for) objects of visual and audio information based on linking with verbal concepts has also been developed. The designed virtual room allows you to fully control the visual and sound environment and conduct fully controlled experiments. Therefore, by creating a library of prepared experiments, you can maximize the testing of various components of the system, as a new economical and effective training option has been developed, developing the system and conducting experiments by placing it in a virtual environment with artificial objects and sounds.

References

1. Altunin A.E., Semukhin M.V. Models and algorithms of decision-making in fuzzy conditions: Monograph. Tyumen: Tyumen State University Press, 2000. 352 p. <http://www.plink.ru/tnm/gl51.htm>
2. Belikova T.P., Yaroslavsky L.P. The use of adaptive amplitude transformations for image preparation // Radio electronics issues, ser. General technical. - 1974. issue 14. - pp. 88–98.
3. Fast algorithms in digital image processing / T. S. Huang, J.-O. Eklund, G. J. Nussbaumer et al. / Ed. Huang T. S. – M.: Radio and Communications, 1984.
4. Volkov V.V., Luizov A.V., Ovchinnikov B.V., Travnikova N.P. Ergonomics of human visual activity. – L.: Mashinostroenie, 1989.
5. Voloshin G.Ya., “Methods of pattern recognition”, 2000. <http://disser.h10.ru/raspobraz.html>, http://users.zebratelecom.ru/~serega123/Metody_r.zip
6. Vorobel R.A., Zhuravel I.M. Image contrast enhancement using a modified piecewise stretching method // Selection and processing of information. – 2000. № 14 (90). – Pp. 116–121.
7. Vorobel R.A., Zhuravel I.M., Opyr N.V., Popov B.O., Derecha V.Ya., Ravlik Ya.M. Method of quantitative assessment of the quality of X-ray images // Proceedings of the Third Ukrainian scientific and technical Conference “Non-destructive testing and technical diagnostics - 2000”. – Dnepropetrovsk. – pp. 233–236
8. Gorelik A.L., Timushev A.G. and others. On-board collision avoidance system for aircraft with power lines // Sensors and systems. – 2001. №3.
9. Gorelik V. A., Kulyasov S. M. Locally unbiased image filtering // Collection of scientific papers “Modeling, decomposition and optimization of complex dynamic processes”. – M: VC RAS, 2002. – pp. 25–30.
10. I.S. Gruzman, V.S. Kirichuk, V.P. Kosykh, G.I. Peretyagin, A.A. Spector, *Digital image processing in information systems: A textbook* (NSTU Publishing House, Novosibirsk, 2000)
11. Yu.K. Kalintsev, *Speech intelligibility in digital vocoders* (Radio and Communications, Moscow, 1991)
12. A course of lectures on the discipline “Artificial intelligence systems”, www.marstu.mari.ru/mmlab/home/AI/10/index.html
13. Pavlovsky Yu.N. Aggregation of complex models and the construction of hierarchical control systems. In the collection: Operations Research. Issue 4, Central Committee of the USSR Academy of Sciences, M., 1974, pp.3–38.
14. Perepelitsyn E.G. Adaptive methods of processing stochastic information in systems for various purposes. M.: TSNIIEISU. 1999.

15. V.V. Pospelov, *Mathematical problems of digital signal processing* (Dis.doctor of Physical and Mathematical Sciences, Moscow, 1994)
16. Pospelov V. V., Chichagov A.V., A method for restoring lost signal fragments // *Autometry*. – 1988. No.1. – pp. 60–64.
17. Fu K., Gonzalez R., Lee K. *Robotics: Translated from English / Edited by V.G. Hradetsky*. – M.: Mir, 1989.
18. T. Vetter, T. Poggio, Linear Object Classes and Image Synthesis from a Single Example Image. *IEEE Trans. Pattern Anal. Mach. Intell.* **19**, 733–742 (1997)
19. T. Winograd, *Language as Cognitive Process*, vol. 1 (Addison-Wesley, Syntax, 1983)
20. K.S. Yoon, Y.K. Ham, R.-H. Park, Hybrid approaches to frontal view face recognition using the Hidden Markov Model and Neural Network. *Pattern Recogn.* **31**, 283–293 (1998)

Basic Methods for Recognizing Visual and Audio Information and the Results of Experiments



Huseynova Zinyet Revayet

Abstract This paper, it is considered to illustrate the experimental results and demonstrate the progress of the parsing and recognition algorithm. Due to the significant volume of reports, only one (the smallest) experiment is presented in full, all others are presented only in key fragments. In general, using the decision-making method used, with equal chances of choosing 1 object of information from 3 variants of vocabulary concepts, improves the quality of the decision by ~ 10–15%, which in most cases is quite enough to achieve a level of reliability (~ 70%).

Keywords Recognition theory · Principle of separation · Decision block · Virtual stage hybrid method · Dynamic generation

1 Introduction

By common definition, recognition is the assignment of a specific object, represented by the values of its properties (features), to one of a fixed list of images (classes) according to a certain decisive rule in accordance with the set goal [1].

It follows from this definition that recognition can be performed by any system that performs the following functions: measuring feature values and performing calculations that implement a crucial rule. It is assumed that the database of images, informative features and decisive rules are either set by the recognizing system from the outside, or are formed by the system itself [2].

It should be noted that many mathematical theories are most effective not in (classical) recognition of graphical information, for which they are completely unsuitable, but in some specific scientific, technical or medical tasks.

Today, recognition as a scientific field, having gone through a long scientific and experimental path, has reached the maximum (or very close to it) values of the quality and reliability of information identification in the conditions of existing theories [3, 4].

H. Z. Revayet (✉)

Nakhichevan State University, Nakhichevan, Republic of Azerbaijan

e-mail: zhuseynova123@gmail.com

Historically, the development of various methods of visual image recognition has been closely related to applied fields (medicine, geology, sociology, chemistry, etc.), and therefore, at the beginning of the development of recognition theory, a fairly rapid development of many methods and algorithms began that could be applied to practical needs and had satisfactory performance. The vast majority of the methods consisted of an algorithm for applying lax but logically thought-out solutions, while the quality of the decisions made in the experiments was the main proof of efficiency [5].

2 Logical Algorithms Adopted During the Experiments

The vast majority of the methods consisted of an algorithm for applying lax but logically thought-out solutions, while the quality of the decisions made in the experiments was the main proof of efficiency.

The second stage in the development of recognition theories is considered to be the transition from an intuitive search for workable algorithms to a purposeful desire, on the one hand, to solve the problem of finding the best algorithm for a specific problem, and on the other, to move from describing individual solutions to describing universal principles of algorithm formation [6].

This approach allowed us to set and solve, within the framework of a certain model, the task of choosing an algorithm with the necessary quality indicators for classification or prediction. In most practical cases, the class of such tasks is small, and it assumes the existence of an accurate model of the objects and phenomena being studied [7].

Many recognition algorithms take into account a probabilistic indicator—the risk of loss, which is used to build optimal decision rules, when choosing the most informative feature system, etc. In general, theories can be divided into the following models of representations (some works use their combinations and combinations or mutual arbitration) [8]:

- Models based on the principle of separation. They differ mainly in defining a class of surfaces, among which a set of surfaces is selected that best separate the elements of different classes [9].
- Statistical models. They are based on the use of the apparatus of statistical decision theory. It is used in cases where the probabilistic characteristics of classes are known or can be determined, for example, the corresponding distribution functions [10].
- Models based on the potential function method. They are based on the idea of a potential defined for any point in space and depending on the location of the potential source. A “potential function” is used as a function of an object’s class membership—a positive and monotonously decreasing distance function everywhere [11].

- Models for calculating proximity estimates. The proximity between parts of previously classified objects and objects that need to be recognized is analyzed. The presence of proximity serves as a partial precedent and is evaluated according to a given rule. Based on a set of proximity estimates, a general estimate of the recognized object for the class is generated, which is the value of the object's class membership function [12, 13].
- Models based on propositional calculus (the apparatus of mathematical logic). The features of objects are described as logical variables, and the description of classes in the language of features is presented in the form of Boolean relations (FAL) [14–18].

3 Summary of the Experiment: 4 Subjects Without Sound

The virtual stage contains 4 objects (Fig. 1). There is no sound information (Tables 1 and 2).



Fig. 1 Experiment: 4 subjects without sound

Table 1 Experiment: 4 subjects without sound

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	A recognized visual or audio object
0:00:00	0:00:00.00	5000, 5000	1.5708, 0	Stage—start
0:05:20	0:00:00.01	5000, 5000	1.5792, 0	Stage—find
0:12:48	0:00:00.12	5000, 5000	1.5825, 0	Stage—find
...				
9:43:19	0:00:03.20	5000, 5000	1.9055, 0	Stage—find
9:46:18	0:00:03.30	4963, 5085	1.9055, 2242	Stage—move
10:05:50	0:00:03.39	4922, 5158	1.9055, 2124	Stage—move
...				
13:14:11	0:00:04.41	4517, 6011	1.9055, 118	Stage—move
13:55:05	0:00:04.61	4517, 6011	1.8957, 109	Stage—research
14:19:21	0:00:04.77	4517, 6011	1.8542, 120	Stage—research
...				
69:19:24	0:00:22.15	4517, 6011	1.9266, 127	Stage—research
69:28:27	0:00:22.26	4517, 6011	1.9068, 112	Stage—research
72:42:24	0:00:22.64	4517, 6011	1.9068, 112	Stage—make decision—possible visual-object “table”
72:43:53	0:00:22.74	4517, 6011	1.9086, 0	Stage—find
72:53:16	0:00:22.85	4517, 6011	1.9154, 0	Stage—find
...				
80:31:53	0:00:25.78	4517, 6011	2.1564, 0	Stage—find
80:51:54	0:00:25.87	4517, 6011	2.1589, 0	Stage—find
81:11:56	0:00:25.95	4485, 6048	2.1589, 1408	Stage—move
81:14:59	0:00:25.99	4408, 6084	2.1589, 1280	Stage—move
...				
82:42:39	0:00:26.41	4015, 6533	2.1589, 256	Stage—move
82:48:06	0:00:26.51	3983, 6597	2.1589, 128	Stage—move
83:03:34	0:00:26.61	3983, 6597	2.1997, 131	Stage—research
83:28:21	0:00:26.79	3983, 6597	2.1755, 126	Stage—research
...				
159:23:11	0:00:53.06	3983, 6597	2.1161, 120	Stage—research
160:58:33	0:00:53.39	3983, 6597	2.1161, 120	Stage—make decision—possible visual-object “chair”
161:13:02	0:00:53.47	3983, 6597	2.1201, 0	Stage—find
161:32:33	0:00:53.51	3983, 6597	2.1277, 0	Stage—find
...				
167:58:46	0:00:55.72	3983, 6597	2.3260, 0	Stage—find
168:06:22	0:00:55.78	3983, 6597	2.3278, 0	Stage—find

(continued)

Table 1 (continued)

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	A recognized visual or audio object
168:08:28	0:00:55.81	3970, 6639	2.3278, 3240	Stage—move
168:22:01	0:00:55.81	3906, 6734	2.3278, 3132	Stage—move
...				
173:24:23	0:00:57.63	2315, 8113	2.3278, 108	Stage—move
174:00:06	0:00:57.79	2315, 8113	2.2935, 108	Stage—research
174:18:58	0:00:57.98	2315, 8113	2.3222, 103	Stage—research
...				
250:05:55	0:01:23.01	2315, 8113	2.2965, 108	Stage—research
250:47:24	0:01:23.17	2315, 8113	2.4052, 106	Stage—research
255:34:27	0:01:23.77	2315, 8113	2.4052, 106	Stage—make decision—possible visual-object “гумбочка”
255:43:23	0:01:23.87	2315, 8113	2.4140, 0	Stage—find
255:54:07	0:01:23.92	2315, 8113	2.4197, 0	Stage—find
...				
263:01:09	0:01:25.83	2315, 8113	2.5667, 0	Stage—find
263:10:59	0:01:25.94	2315, 8113	2.5692, 0	Stage—find
263:11:47	0:01:25.95	2251, 8167	2.5692, 1599	Stage—move
...				
265:14:11	0:01:26.59	1674, 8725	2.5692, 246	Stage—move
265:29:56	0:01:26.59	1589, 8795	2.5692, 123	Stage—move
266:12:06	0:01:26.71	1589, 8795	2.5244, 114	Stage—research
266:45:10	0:01:26.84	1589, 8795	2.5570, 118	Stage—research
...				
348:20:03	0:01:55.78	1589, 8795	2.5437, 115	Stage—research
348:35:55	0:01:55.96	1589, 8795	2.5563, 125	Stage—research
352:18:23	0:01:56.33	1589, 8795	2.5563, 125	Stage—make decision—possible visual-object “chair”
352:18:23	0:01:56.33	1589, 8795	2.5563, 125	Stage—finish
266:12:06	0:01:26.71	1589, 8795	2.5244, 114	Stage—research

4 Experiment 5 Items Without Sound (Undiscovered Item)

The virtual stage contains 5 objects (Fig. 2). There is no sound information.

In this experiment, a possible situation is demonstrated when the IL does not find an object (a chair is highlighted in a light color) that was blocked by another object (a table), and the experiment was forcibly stopped (Tables 3, 4).

Table 2 The final report of the experiment: 4 subjects without sound

Real time	Virtual time	Virtual object	Sound object	Reliability (%)
72:42:24	0:00:22.64	Possible visual-object “table”		73.46
160:58:33	0:00:53.39	Possible visual-object “chair”		72.12
255:34:27	0:01:23.77	Possible visual-object “a bedside table chair”		65.88
352:18:23	0:01:56.33	Possible visual-object “chair”		77.13
72:42:24	0:00:22.64	Possible visual-object “table”		73.46



Fig. 2 Experiment: 5 objects without sound, with a demonstration of an undiscovered object (highlighted in light)

5 Conclusion

The developed function for calculating the similarity of objects, which is the functional center of modular services, is the only algorithm that has shown satisfactory results and high stability of reliability indicators. The paper examines the results of experiments and demonstrates the progress of the parsing and recognition algorithm. Due to the significant volume of reports, only one (the smallest) experiment is presented in full, all others are presented only in key fragments. In general, using the decision-making method used, with equal chances of choosing 1 object of information from 3 variants of vocabulary concepts, improves the quality of the decision by ~ 10–15%, which in most cases is quite sufficient to achieve a level of reliability (~ 70%).

Table 3 Experiment: 5 objects without sound (demonstration of an undiscovered object)

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	A recognized visual or audio object
0:00:00	0:00:00.00	5000, 5000	1.5708, 0	Stage—start
0:01:54	0:00:00.11	5000, 5000	1.5721, 0	Stage—find
...				
8:12:12	0:00:03.14	5000, 5000	1.8320, 0	Stage—find
8:27:28	0:00:03.17	5000, 5000	1.8388, 0	Stage—find
8:40:55	0:00:03.20	4978, 5027	1.8388, 3136	Stage—move
8:42:10	0:00:03.27	4962, 5071	1.8388, 3024	Stage—move
...				
12:26:34	0:00:04.71	4550, 6923	1.8388, 224	Stage—move
12:38:37	0:00:04.77	4544, 6980	1.8388, 112	Stage—move
12:50:58	0:00:04.98	4544, 6980	1.8138, 118	Stage—research
13:17:18	0:00:05.18	4544, 6980	1.8703, 120	Stage—research
...				
84:33:55	0:00:29.23	4544, 6980	1.7734, 109	Stage—research
84:45:17	0:00:29.44	4544, 6980	1.7696, 108	Stage—research
86:12:13	0:00:29.74	4544, 6980	1.7696, 108	Stage—make decision—possible visual-object “гумбочка”
86:19:27	0:00:29.79	4544, 6980	1.7748, 0	Stage—find
86:31:33	0:00:29.84	4544, 6980	1.7801, 0	Stage—find
...				
91:22:11	0:00:31.42	4544, 6980	1.8901, 0	Stage—find
91:36:15	0:00:31.42	4544, 6980	1.8926, 0	Stage—find
91:54:31	0:00:31.51	4527, 7049	1.8926, 1755	Stage—move
92:09:44	0:00:31.57	4508, 7138	1.8926, 1638	Stage—move
...				
0:00:00	0:00:00.00	5000, 5000	1.5708, 0	Stage—start
0:01:54	0:00:00.11	5000, 5000	1.5721, 0	Stage—find
...				
8:12:12	0:00:03.14	5000, 5000	1.8320, 0	Stage—find
8:27:28	0:00:03.17	5000, 5000	1.8388, 0	Stage—find
8:40:55	0:00:03.20	4978, 5027	1.8388, 3136	Stage—move
8:42:10	0:00:03.27	4962, 5071	1.8388, 3024	Stage—move
...				
12:26:34	0:00:04.71	4550, 6923	1.8388, 224	Stage—move
12:38:37	0:00:04.77	4544, 6980	1.8388, 112	Stage—move
12:50:58	0:00:04.98	4544, 6980	1.8138, 118	Stage—research

(continued)

Table 3 (continued)

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	A recognized visual or audio object
13:17:18	0:00:05.18	4544, 6980	1.8703, 120	Stage—research
...				
84:33:55	0:00:29.23	4544, 6980	1.7734, 109	Stage—research
84:45:17	0:00:29.44	4544, 6980	1.7696, 108	Stage—research
86:12:13	0:00:29.74	4544, 6980	1.7696, 108	Stage—make decision—possible visual-object “гумбочка”
86:19:27	0:00:29.79	4544, 6980	1.7748, 0	Stage—find
86:31:33	0:00:29.84	4544, 6980	1.7801, 0	Stage—find
...				
91:22:11	0:00:31.42	4544, 6980	1.8901, 0	Stage—find
91:36:15	0:00:31.42	4544, 6980	1.8926, 0	Stage—find
91:54:31	0:00:31.51	4527, 7049	1.8926, 1755	Stage—move
92:09:44	0:00:31.57	4508, 7138	1.8926, 1638	Stage—move
94:19:04	0:00:32.24	4260, 7908	1.8926, 234	Stage—move
94:21:48	0:00:32.25	4217, 7960	1.8926, 117	Stage—move
94:30:55	0:00:32.36	4217, 7960	1.8064, 120	Stage—research
94:39:39	0:00:32.49	4217, 7960	1.8570, 112	Stage—research
...				
142:46:08	0:00:49.15	4217, 7960	1.9736, 111	Stage—research
143:31:14	0:00:49.30	4217, 7960	1.8117, 117	Stage—research
147:14:46	0:00:49.67	4217, 7960	1.8117, 117	Stage—make decision—possible visual-object “стыл”
147:27:55	0:00:49.72	4217, 7960	1.8191, 0	Stage—find
...				
154:40:10	0:00:51.69	4217, 7960	2.0041, 0	Stage—find
155:00:29	0:00:51.73	4217, 7960	2.0098, 0	Stage—find
155:12:42	0:00:51.82	4193, 8003	2.0098, 1624	Stage—move
155:19:07	0:00:51.93	4175, 8060	2.0098, 1508	Stage—move
...				
157:09:10	0:00:52.32	3929, 8495	2.0098, 232	Stage—move
157:12:45	0:00:52.32	3869, 8522	2.0098, 116	Stage—move
157:46:22	0:00:52.53	3869, 8522	1.9231, 125	Stage—research
158:22:39	0:00:52.74	3869, 8522	2.0424, 114	Stage—research
...				
247:41:28	0:01:23.22	3869, 8522	2.0046, 122	Stage—research
248:14:25	0:01:23.43	3869, 8522	2.0823, 112	Stage—research

(continued)

Table 3 (continued)

Real time	Virtual time	Coordinates AP (X, Y)	Angle and focal length AP (α , R)	A recognized visual or audio object
249:46:07	0:01:23.95	3869, 8522	2.0823, 112	Stage—make decision—possible visual-object “стол”
249:50:26	0:01:23.97	3869, 8522	2.0850, 0	Stage—find
249:53:58	0:01:24.08	3869, 8522	2.0901, 0	Stage—find
...				
252:19:42	0:01:25.00	3869, 8522	2.1823, 0	Stage—find
252:39:31	0:01:25.10	3869, 8522	2.1868, 0	Stage—find
252:49:32	0:01:25.21	3866, 8633	2.1868, 2825	Stage—move
252:54:15	0:01:25.29	3816, 8746	2.1868, 2712	Stage—move
...				
256:51:33	0:01:26.32	2949, 10,164	2.1868, 113	Stage—move
257:02:44	0:01:26.51	2949, 10,164	2.2737, 121	Stage—research
...				
350:23:23	0:01:56.94	2949, 10,164	2.2518, 115	Stage—research
350:37:58	0:01:57.12	2949, 10,164	2.1324, 104	Stage—research
353:40:06	0:01:57.61	2949, 10,164	2.1324, 104	Stage—make decision—possible visual-object “стул”
353:43:48	0:01:57.65	2949, 10,164	2.1340, 0	Stage—find
...				
424:52:10	0:02:19.03	2949, 10,164	2.6407, 113	Stage—find
424:52:10	0:02:19.03	2949, 10,164	2.6407, 113	Stage—abort

Table 4 The final report of the experiment: 5 objects without sound (demonstration of an undiscovered object)

Real time	Virtual time	Virtual object	Sound object	Reliability (%)
86:12:13	0:00:29.74	Possible visual-object “a bedside table chair”		73.87
147:14:46	0:00:49.67	Possible visual-object “chair”		78.65
249:46:07	0:01:23.95	Possible visual-object “table”		78.37
353:40:06	0:01:57.61	Possible visual-object “chair”		75.54
424:52:10	0:02:19.03	Abort		

References

1. Y. LeCun, O. Matan, B. Boser, J.S. Denker, D. Henderson, R.E. Howard, W. Hubbard, L.D. Jackel, H.S. Baird, Handwritten zipcode recognition with multilayer networks, in *Proc. of International Conference on Pattern Recognition* (Atlantic City, 1990)

2. S.A. Sharov, in *Tools for Computer Representation of Linguistic Information* (1996). <http://www.rriai.org.ru/~sharoff/>, sharoff@mx.iki.rssi.ru
3. J. Kralik, P. Stiegler, Z. Vostry, J. Zavorka, Modelovani dynamiky rozsahlych siti. Praha, Akademia 364. by neural networks with single-layer training. *IEEE Trans. Neural Netw.* 3, 962–968 (1984)
4. T. Winograd, in *Language as Cognitive Process*, vol. 1 (Syntax. Addison, Wesley, 1983)
5. S. Ranganath, K. Arun, Face recognition using transform features and neural networks. *Pattern Recogn.* **30**, 1615–1622 (1997)
6. R.C. Gonzalez, P. Wintz, in *Digital Image Processing* (Addison–Wesley. Reading, Massachusetts, 1987), pp. 505
7. T. Winograd, *Language as Cognitive Process*, vol. 1 (Addison-Wesley, Syntax, 1983)
8. A. Krzyzak, W. Dai, C.Y. Suen, Unconstrained handwritten character classification using modified backpropagation model, in *Proc. 1st Int. Workshop on Frontiers in Handwriting Recognition* (Montreal, Canada, 1990), pp. 155–166
9. G. Sabah, Knowledge representation and natural language understanding, vol. 6(3–4) (AI Communications, 1993), pp. 155–186
10. S.-W. Lee, Y.J. Kim, Off-line recognition of totally unconstrained handwritten numerals using multilayer cluster neural network, in *Proc. Of the 12th IAPR International Conference on Pattern Recognition* (Jerusalem, Israel, 1994), pp. 507–509
11. S.G. Tzafestas, Large-scale systems modeling in distributed-parameter control and estimation, in *Modeling and Simul. Eng. 10th IMACS World Congr. Syst. Simul. and Sci. Comput., Montreal, 8–13 Aug. 1982*, v.3 (Amsterdam e.a., 1983), pp. 69–77
12. D.J. Ketcham, Real time image enhancement technique, in *Proceedings SPIE/OSA Conference on Image Processing*, vol. 74 (Pacific Grove, California, 1976), pp. 120–125
13. M. Petrou, in *Learning in Pattern Recognition. Lecture Notes in Artificial Intelligence—Machine Learning and Data Mining in Pattern Recognition* (1999), pp. 1–12
14. D. Valentin, H. Abdi, A.J. O’Toole, G.W. Cottrell, Connectionist models of face processing: a survey. *Pattern Recognit.* **27**, 1209–1230 (1994)
15. T. Vetter, T. Poggio, Linear object classes and image synthesis from a single example image. *IEEE Trans. Pattern Anal. Mach. Intell.* **19**, 733–742 (1997)
16. W.W. Translation, T. Report, *Reprinted in Machine Translation of languages, 1955* (MIT Press, Cambridge, MA, 1949), pp.15–23
17. I.D. Wilson, Foundations of hierarchical control. *Int. J. Control.* **29**(N6),899–933 (1979)
18. S. Gutta, H. Wechsler, Face recognition using hybrid classifiers. *Pattern Recogn.* **30**, 539–553 (1997)

Intelligent Methods and Fuzzy Approach

An Approach to Feature Engineering in a Data-Driven Production Management



Aleksey A. Filippov , Gleb Yu. Guskov , Anton A. Romanov ,
and Nadezhda G. Yarushkina

Abstract This article presents an approach to feature engineering and dataset synthesis, considering a high uncertainty and constraints of a domain. Feature engineering is a creative process that requires the analyst to have deep knowledge of a specific domain and significant experience in data analysis. When solving a new task, the analyst may face a situation with a high uncertainty, which requires the application of modern methods of feature engineering, versioning, documentation, and storage of features. To solve feature engineering in conditions of high uncertainty regarding the constraints of a specific domain, we developed an information model of the context. This model allows an expert to identify important entities of a domain and their attributes from his point of view. There may be a situation where the dataset is small or unbalanced. To address this issue, we proposed a method for synthesizing the dataset based on a multi-agent approach. Within the multi-agent system, each agent can represent an entity or some of its attributes to synthesize a specific fragment of the dataset. The proposed multi-agent system for synthesizing the dataset is a directed acyclic graph and allows modeling of the behavior of entities and processes in a specific domain.

Keywords Feature engineering · Multi-agent system · Dataset synthesis

1 Introduction

A volume of dataset sufficient to identify hidden regularities in the data is required [1–5] to get high-quality object control models. The model cannot achieve the required level of generalization if there is not enough data, which will affect the quality of the model. The quality of the data itself also affects the quality of the model.

A. A. Filippov · G. Yu. Guskov · A. A. Romanov · N. G. Yarushkina (✉)
Ulyanovsk State Technical University, Ulyanovsk, Russia
e-mail: jng@ulstu.ru

A. A. Filippov
e-mail: al.filippov@ulstu.ru

The most important stage in data analysis is the feature engineering stage. Feature engineering means identifying significant entities and their attributes relevant to the task. This is followed by matching each attribute with the part of the data that describes historical values of the analyzed objects. Analysts are required to have knowledge of a specific domain, substantial experience in data analytics, and machine learning.

Currently, there are software tools available for automatic feature engineering for relational data, such as Featuretools [6]. When using such tools, new features are generated based on combining existing ones and applying various data transformation functions and aggregation functions. However, in this way of feature engineering, the specifics of the domain are not considered, and the dimensionality of the feature space is significantly increased.

Various methods of data synthesis and augmentation are applied [1–5, 7–13] in case of lacking dataset volume. Rotation, translation, and noise addition algorithms are used for images, while methods based on hidden Markov models, process simulators for specific domains, and machine learning models are used for relational data [3–5].

Methods such as Zero/Few-Shot Learning and Transfer Learning [14–16] can also be applied, where a previously trained model is used to solve a similar problem with additional training on a few examples.

Feature engineering may also encounter problems related to the high dimensionality of the feature space, imbalanced sample size, and overfitting [3, 17–19]. The size of the feature space affects the speed and computational complexity of the model training process, as well as the speed of operation of the resulting model. Various features can lead to overfitting and reduced model quality. Imbalanced samples lead to artificially inflated evaluation results of the model performance because of the presence of a class with the largest number of examples.

Therefore, there are specific areas where the application of data analysis methods is hindered by the lack of sufficient dataset. However, it is possible to involve a domain expert who can build a model of the analyzed object to synthesize dataset of sufficient volume based on prior domain knowledge.

Thus, an approach to feature engineering in machine learning for data with a high uncertainty and domain constraints is needed, considering the above problems.

2 Feature Engineering Considering High Levels of Uncertainty and Domain Constraints

Feature engineering is a creative process that requires an analyst to have deep knowledge of a specific domain and extensive experience in data analytics and machine learning. An analyst may challenge a high uncertainty when solving a new task, which requires the application of modern methods for constructing, versioning, documenting, and storing features. For example, Uber uses a machine learning platform

called Michelangelo, which includes a feature store. Feature stores allow for tracking changes in both the features themselves and the resulting machine learning models. Analyst can always track the changes and analyze them when the model quality degrades.

Therefore, feature engineering involves considering the characteristics of both the analyzed object and a specific domain constraint. The information model of context was developed to address feature engineering in conditions of high uncertainty and considering the constraints of a specific domain. This model allows an expert to identify important domain entities and processes and their attributes from their perspective.

Descriptive logic is used to represent the proposed context model [20]:

$$D = \langle D^T Box, D^A Box \rangle,$$

where D^{TBox} is a set of terminology definition axioms of the model, and D^{ABox} is a set of facts (assertions, axioms) of the model.

Let us consider the terminology D^{TBox} of the proposed context model in more detail.

The properties (roles) common to all classes of the model are:

$$\top \sqsubseteq \exists hasName.String \sqcap \forall hasName.String \sqcap = 1 hasName.String,$$

where *hasName* is a functional role to determine the model entity name.

D^{TBox} terminology contains the following classes:

- *Object* is a class for describing analyzed objects,
- *Attribute* is a class for describing the attributes of an analyzed object,
- *Dependency* is a class for describing dependencies between attributes of several analyzed objects,
- *DependencyType* is a class for describing types (positive and negative) of dependencies (degree of influence):

$$Object \sqsubseteq \top$$

$$Attribute \sqsubseteq \top$$

$$Dependency \sqsubseteq \top$$

$$DependencyType \equiv \{positive, negative\}.$$

The classes *Object*, *Attribute*, *Dependency* и *DependencyType* are declared as disjoint classes:

$$Object \sqcap Attribute \sqcap Dependency \sqcap DependencyType \sqsubseteq \perp.$$

The *hasAttribute* role of the *Object* class is specified to determine the list of attributes of the analyzed object:

$$Object \sqsubseteq \forall hasAttribute.Attribute.$$

The *Attribute* class has the following roles:

- *hasMin* and *hasMax* are functional roles to determine the minimum and maximum value of an attribute, respectively,
- *hasSource* is a functional role for defining the column of the local storage table in which the values of this attribute are stored,
- *hasDependency* is a role for defining dependencies between attributes:

$$Attribute \sqsubseteq \forall hasMin.Double \sqcap = 1hasMin.Double \sqcap$$

$$\sqcap \forall hasMax.Double \sqcap = 1hasMax.Double \sqcap$$

$$\sqcap \exists hasSource.String \sqcap \forall hasSource.String \sqcap = 1hasSource.String \sqcap$$

$$\sqcap \forall hasDependency.Dependency.$$

The following roles were created for the *Dependency* class:

- *dependsOn* is a functional role to define dependency between attributes,
- *hasType* is a functional role for determining the type of dependency between attributes,
- *hasWeight* is a functional role for determining dependency weight:

$$Dependency \sqsubseteq \exists dependsOn.Attribute \sqcap \forall dependsOn.Attribute \sqcap$$

$$\sqcap = 1dependsOn.Attribute \sqcap$$

$$\sqcap \exists hasType.DependencyType \sqcap \forall hasType.DependencyType \sqcap$$

$$\sqcap = 1hasType.DependencyType \sqcap$$

$$\sqcap \exists hasWeight.Double \sqcap \forall hasWeight.Double \sqcap = 1hasWeight.Double.$$

A set of classes is defined in the context model to describe analyzed objects, attributes, dependencies, and dependency types. Various roles and properties are specified for each class to define their characteristics.

The facts of the proposed context model are automatically generated using specialized forms to facilitate the expert work. The ontological language OWL 2 [21] is used to represent information based on the context model, while a separate local storage table is used to store the history of feature information changes.

3 Multi-agent Approach for a Dataset Synthesis

There may be situations where the dataset is small or imbalanced after feature engineering. As a result, get a model of the control object with an acceptable level of quality becomes difficult [3]. Various approaches are used in modern practice to solve this problem: augmentation [7, 8], synthetic dataset generation [9–13], zero/few-shot learning [14], transfer learning [15, 16], etc. We propose a method (Sect. 2) for generating a synthetic dataset based on feature engineering considering domain constraints to solve these problems.

The proposed method of synthetic dataset generation is based on a multi-agent approach. The multi-agent approach allows to create universal and flexible systems that can adapt to changes both in a domain and in the external environment relative to the system [22, 23].

Each agent can represent the entity or some of its attributes for synthetic dataset generation within the multi-agent system.

System for synthetic dataset generation based can be represented as:

$$Synth = A^{Synth}, E^{Synth},$$

where A^{Synth} is the set of agents, each of which represents a simple phenomenon from the point of view of a specific domain in which the entity takes part or which affects the attributes of the entity. Each agent contains an implementation of an algorithm that takes the external parameters as input and outputs values of individual features, E^{Synth} is the set of connections between agents. Connections allow linking the input and output of individual agents. Some algorithms may be associated with the connections. An algorithm may modify the values passed between agents. This allows for emulation of the impact of the external environment on the system.

The presented system of agents for synthetic dataset generation is a directed acyclic graph and allows for modeling the behavior of entities and processes in a specific domain.

As individual entities operate in isolation, diverse simulation and modeling systems can act as agents, facilitating the creation of a dataset that closely approximates authentic data.

As well as to facilitate the creation of individual entities, mechanisms of fuzzy logic and fuzzy sets can be used when creating agents and connections in the synthetic dataset generation system to account for fuzziness and uncertainty. For example, fuzzy sets can be used to transform categorical values into numerical values during defuzzification. There is also the possibility of using fuzzy logic inference mechanisms to implement agents considering the context [24].

Let us consider an example of using the proposed approach for generating a time series of an indicator of some object, considering the peculiarities of a specific domain. The data is generated based on information received during feature construction. This method allows to approximate the procedure of creating a synthetic dataset to the approximation of time series based on expert information [25].

The following parameters of the agent were determined:

- Base level of the time series y_{base} . In most cases it is the average value of the points of the time series.
- Main trend of the time series ($tend_{base}$).
- Constraints on maximum or minimum values of the series ($bound = (y_{min}, y_{max})$).
- Rate of trend change (inertia) of the time series ($tend_{\Delta}$).
- Range of acceptable values that are not anomalous ($accept = (y_{min}^{accept}, y_{max}^{accept})$).

These parameters can be divided into types:

1. The parameters determine the behavior of the time series: the base level of the series and the main trend. A rough model of the time series can be presented without analyzing the history of the numerical representation of the time series, relying only on information about a domain.
2. The parameters describe the results of modeling: constraints on minimum and maximum values, trend change rate, range of acceptable values. In addition, determining information can be used for validating forecast values.

Formally represent both types of parameters as:

$$Def = \{y_{base}, tend_{base}\},$$

$$Desc = \{bound, tend_{\Delta}, accept, tend_{base}\}.$$

The model of the time series generated based on features from contextual information can be represented by the following mathematical expression:

$$y_{Def} = \alpha * y_{base} + \beta * tend_{base},$$

where α, β are coefficients of the weight components of the model.

Thus, the results show that the model of the generated synthetic time series is an additive combination of components formed from features.

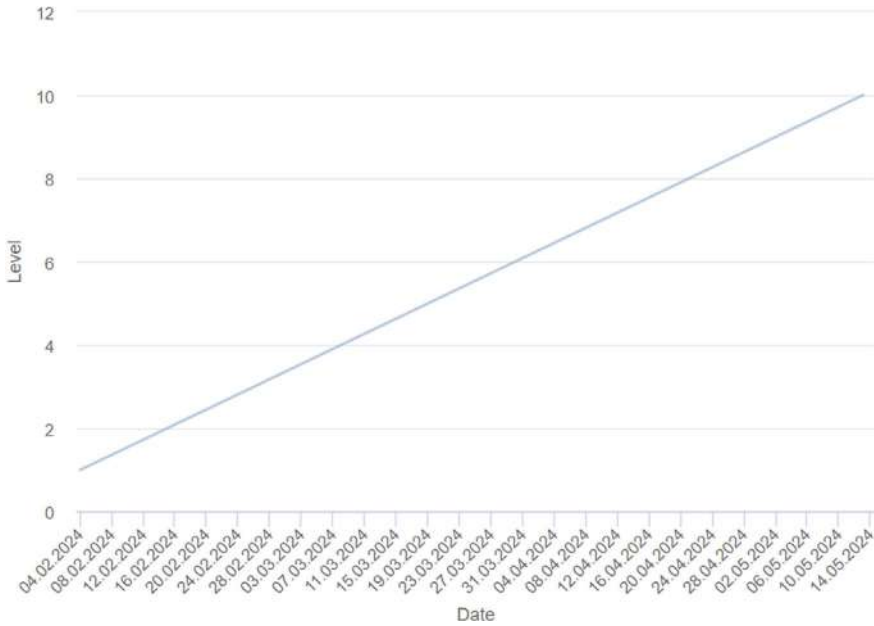


Fig. 1 Generated time series based on features without random component

Let us consider examples of generated time series of development process indicators for the software project development domain. Generation can be based on basic components (Fig. 1) or using a random component r (Fig. 2):

$$y_{Def} = \alpha * y_{base} + \beta * tend_{base} + r.$$

Figure 1 shows the generated time series of the “Number of open tasks” indicator (length of time series = 100, $bound = (0,10)$, $y_{base} = 3$, $tend_{base} = 1$, $tend_{\Delta} = 1$). In time series components, analysts often distinguish a non-systematic component or random noise during modeling. In our approach, when generating time series based on features, we generate this component randomly. We set the value in a wide range, significantly lower of the absolute values of the time series (Fig. 2a) or bigger than them (Fig. 2b). This is important because it creates data to test the adequacy of third-party models.

Thus, the proposed approach allows creating multi-agent systems for modeling the behavior of individual entities and business processes in a specific domain. This system allows to synthesize a dataset in case of insufficient data, or if the sample is unbalanced.

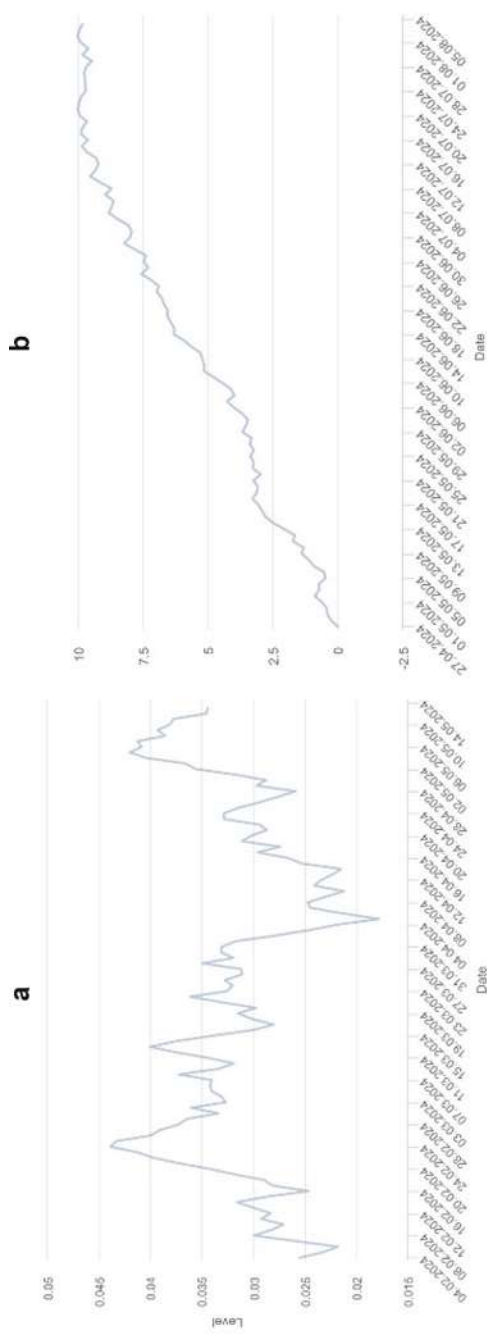


Fig. 2 Generated time series based on features with a different value of random component

4 Conclusion

This article describes a proposed approach to feature engineering in machine learning for data with a high uncertainty and domain constraints. A context model and tools for automating the feature engineering process have been proposed. This method can be applied in situations where the task is tackled under uncertainty, but domain experts familiar with the constraints are involved. The resulting features are stored as an ontology in OWL 2 language, allowing for storage, documentation, and versioning of features. These stored features can be used in solving similar tasks.

In some cases, there may not be enough data to train machine learning models, or the data may be imbalanced. The proposed method of synthesizing a dataset based on a multi-agent approach enables the construction of a model for a specific entity or business process, considering domain constraints and characteristics, and synthesizing the necessary volume of data.

Acknowledgements This study was supported the Ministry of Science and Higher Education of Russia in framework of project No. 075-03-2023-143 «The study of intelligent predictive analytics based on the integration of methods for constructing features of heterogeneous dynamic data for machine learning and methods of predictive multimodal data analysis».

References

1. I. Kraljevski, Y.C. Ju, D. Ivanov, C. Tschöpe, M. Wolff, How to do machine learning with small data? A review from an industrial perspective. arXiv preprint [arXiv:2311.07126](https://arxiv.org/abs/2311.07126) (2023)
2. H.P. Das, R. Tran, J. Singh, X. Yue, G. Tison, A. Sangiovanni-Vincentelli, C.J. Spanos, Conditional synthetic data generation for robust machine learning applications with limited pandemic data. *Proc. AAAI Conf. Artif. Intell.* **36**(11), 11792–11800 (2022)
3. A. Narayanan, T. Hardy, Synthetic data generation for machine learning model training for energy theft scenarios using cosimulation. *IET Gener. Transm. Distrib.* **17**(5), 1035–1046 (2023)
4. J. Dahmen, D. Cook, SynSys: a synthetic data generation system for healthcare applications. *Sensors* **19**(5), 1181 (2019)
5. G. Forestier, F. Petitjean, H.A. Dau, G.I. Webb, E. Keogh, Generating synthetic time series to augment sparse datasets, in *2017 IEEE International Conference On Data Mining (ICDM)*, pp. 865–870 (2017)
6. Featuretools, <https://www.featuretools.com>, last accessed 2023/12/01
7. I. Virkkunen, T. Koskinen, O. Jessen-Juhler, J. Rinta-Aho, Augmented ultrasonic data for machine learning. *J. Nondestr. Eval.* **40**(1), 4 (2021)
8. C. Shorten, T.M. Khoshgoftaar, A survey on image data augmentation for deep learning. *J. Big Data* **6**(1), 1–48 (2019)
9. S.I. Nikolenko, *Synthetic Data For Deep Learning*, vol. 174 (2021)
10. A. Gupta, A. Vedaldi, A. Zisserman, Synthetic data for text localisation in natural images, in *Proceedings of the IEEE Conference On Computer Vision and Pattern Recognition*, pp. 2315–2324 (2016)
11. N. Bannur, V. Shah, A. Raval, J. White, Synthetic data generation for improved covid-19 epidemic forecasting. *medRxiv*, pp. 2020–12 (2020)

12. A. Ghorbani, V. Natarajan, D. Coz, Y. Liu, Dermgan: synthetic generation of clinical skin images with pathology, in *Machine Learning For Health Workshop*, pp. 155–170 (2020)
13. M. Frid-Adar, E. Klang, M. Amitai, J. Goldberger, H. Greenspan, Synthetic data augmentation using GAN for improved liver lesion classification, in *2018 IEEE 15th International Symposium On Biomedical Imaging (ISBI 2018)*, pp. 289–293 (2018)
14. Y. Xian, C.H. Lampert, B. Schiele, Z. Akata, Zero-shot learning—a comprehensive evaluation of the good, the bad and the ugly. *IEEE Trans. Pattern Anal. Mach. Intell.* **41**(9), 2251–2265 (2018)
15. F. Zhuang, Z. Qi, K. Duan, D. Xi, Y. Zhu, H. Zhu, H. Xiong, Q. He, A comprehensive survey on transfer learning. *Proc. IEEE* **109**(1), 43–76 (2020)
16. C. Tan, F. Sun, T. Kong, W. Zhang, C. Yang, C. Liu, A survey on deep transfer learning, in *Artificial Neural Networks and Machine Learning–ICANN 2018: 27th International Conference on Artificial Neural Networks*, vol. III (27), pp. 270–279 (2018)
17. N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, R. Salakhutdinov, Dropout: a simple way to prevent neural networks from overfitting. *J. Mach. Learn. Res.* **15**(1), 1929–1958 (2014)
18. H.P. Das, P. Abbeel, C.J. Spanos, Likelihood contribution based multi-scale architecture for generative flows. arXiv preprint [arXiv:1908.01686](https://arxiv.org/abs/1908.01686) (2019)
19. S.C. Yip, W.N. Tan, C. Tan, M.T. Gan, K. Wong, An anomaly detection framework for identifying energy theft and defective meters in smart grids. *Int. J. Electr. Power Energy Syst.* **101**, 189–203 (2018)
20. S. Rudolph, Foundations of description logics, in *Reasoning Web International Summer School*, pp. 76–136 (2011)
21. OWL 2 Web Ontology Language, <https://www.w3.org/TR/owl2-overview/>, last accessed: 2023/12/01.
22. A. Dorri, S.S. Kanhere, R. Jurdak, Multi-agent systems: a survey. *IEEE Access* **6**, 28573–28593 (2018)
23. J. Qin, Q. Ma, Y. Shi, L. Wang, Recent advances in consensus of multi-agent systems: a brief survey. *IEEE Trans. Industr. Electron.* **64**(6), 4972–4983 (2016)
24. A. Romanov, J. Stroeva, A. Filippov, N. Yarushkina, An approach to building decision support systems based on an ontology service. *Mathematics* **9**(22), 2946 (2021)
25. A. Romanov, A. Filippov, N. Yarushkina, An approach to contextual time series analysis, in *Artificial Intelligence and Soft Computing: 20th International Conference, ICAISC 2021*, vol. I (20), pp. 496–505 (2021)

How to Make a Decision Under Interval Uncertainty If We Do Not Know the Utility Function



Jeffrey Escamilla and Vladik Kreinovich

Abstract Decision theory describes how to make decisions, in particular, how to make decisions under interval uncertainty. However, this theory's recommendations assume that we know the utility function – that describes the decision maker's preferences. Sometimes, we can make a recommendation even when we do not know the utility function. In this paper, we provide a complete description of all such cases.

1 Formulation of the Problem

How a rational person should make a decision: a brief reminder. According to decision theory (see, e.g., [1, 2, 4–8]), decisions of a rational person – i.e., e.g., a person who when preferring A to B and B to C always prefers A to C —are described by a function $u(a)$ called *utility*.

In this description, an alternative in which we gain amount a is better than the alternative in which we gain amount b if and only if $u(a) > u(b)$.

In business situations, when different outcomes can be described by monetary gains, the utility function is (non-strictly) increasing of the corresponding gain:

- If $a \leq b$
- Then $u(a) \leq u(b)$.

Need to take uncertainty into account. In practice, we only know the consequence of each action with uncertainty.

Case of interval uncertainty: what is it and what should we do. In many cases, all we know is the bounds on possible gain, i.e., the interval $[\underline{a}, \bar{a}]$ of possible values of the gain.

J. Escamilla · V. Kreinovich (✉)

Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA

e-mail: vladik@utep.edu

J. Escamilla

e-mail: jescamilla2@miners.utep.edu

In this case, according to decision theory (see, e.g., [3–5]), the decision maker should:

- Select some value $\alpha_H \in [0, 1]$ describing the decision maker's degree of optimism-pessimism, and then
- Select the alternative for which the value $\alpha_H \cdot u(\bar{a}) + (1 - \alpha_H) \cdot u(\underline{a})$ is the largest.

Remaining problem: what if we do not know the utility function? Recommendations of decision theory are based on the assumption that we know the utility function of a decision maker. However, sometimes, we do not know the utility function.

Can we sometimes recommend a decision even if we do not know the utility function? And if yes, when? When can we still conclude that $[\underline{a}, \bar{a}]$ is better than $[\underline{b}, \bar{b}]$?

Sometimes, the solution is easy. The answer is easy:

- When $\alpha_H = 1$ —then we select the alternative with the larger $\bar{a}$, and
- When $\alpha_H = 0$ —then we select the alternative with the larger $\underline{a}$.

But in general, the problem remains open. But what can we suggest in realistic situations, when $0 < \alpha_H < 1$?

What we do in this paper. In this paper, we provide a complete answer to this challenge.

Structure of this paper. In Sect. 2, we formulate and prove our main result. Two auxiliary results form Sect. 3.

2 Main Result

It turns out that the only case when we can make a recommendation without knowing the utility function is when both bounds of the a -interval are larger than the corresponding bounds of the b -interval. Let us describe this result in precise terms.

Main result. Let $\alpha_H \in (0, 1)$ be a real number. Then, for every two intervals $[\underline{a}, \bar{a}]$ and $[\underline{b}, \bar{b}]$, the following two conditions are equivalent:

1. for all non-strictly increasing functions $u(a)$, we have

$$\alpha_H \cdot u(\bar{a}) + (1 - \alpha_H) \cdot u(\underline{a}) \geq \alpha_H \cdot u(\bar{b}) + (1 - \alpha_H) \cdot u(\underline{b});$$

2. $\underline{a} \geq \underline{b}$ and $\bar{a} \geq \bar{b}$.

Proof. To prove the desired equivalence, it is sufficient to prove that:

- If the Condition 2. is satisfied, then the Condition 1. is also satisfied, and
- If the Condition 2. is not satisfied, then the Condition 1. is also not satisfied.

1°. Let us first prove that if the Condition 2. is satisfied, then the Condition 1. is also satisfied.

Indeed, let us assume that the Condition 2. is satisfied. Then, due to monotonicity of the utility function:

- Since $\underline{a} \geq \underline{b}$, we get $u(\underline{a}) \geq u(\underline{b})$; and
- Since $\bar{a} \geq \bar{b}$, we get $u(\bar{a}) \geq u(\bar{b})$.

If we multiply both sides of the first inequality by $\alpha_H > 0$ and both sides of the second inequality by $1 - \alpha_H > 0$, then we conclude that

$$\alpha_H \cdot u(\underline{a}) \geq \alpha_H \cdot u(\underline{b}) \text{ and } (1 - \alpha_H) \cdot \bar{a} \geq (1 - \alpha_H) \cdot \bar{b}.$$

If we add these two inequalities, we get the desired Condition 1.

2°. Let us prove that if the Condition 2. is not satisfied, i.e., if $\underline{a} < \underline{b}$ or $\bar{a} < \bar{b}$, then the Condition 1. is violated for some increasing function $u(a)$.

We will consider the two possibilities $\underline{a} < \underline{b}$ or $\bar{a} < \bar{b}$ separately.

2.1°. Let us first consider the case when $\underline{a} < \underline{b}$.

In this case, let us consider the following utility function $u(a)$:

- We take $u(a) = 0$ for $a \leq \underline{a}$, and
- We take $u(a) = 1$ for all $a > \underline{a}$.

Let us compute the right- and left-hand sides of Condition 1. for this utility function.

- Since $\bar{b} \geq \underline{b}$ and $\underline{b} > \underline{a}$, we have $\bar{b} > \underline{a}$. Thus, $u(\bar{b}) = u(\underline{b}) = 1$, and the right-hand side of the Inequality 1. is equal to $\alpha_H + (1 - \alpha_H) = 1$.
- On the other hand, since $\underline{a} \leq \underline{a}$, we have $u(\underline{a}) = 0$, so the left-hand side is equal to $\alpha_H \cdot u(\bar{a})$. Here, $u(\bar{a}) \leq 1$, so $\alpha_H \cdot u(\bar{a}) \leq \alpha_H < 1$.

Thus, the Inequality 1. is not satisfied.

2.2°. Let us now consider the case when $\bar{a} < \bar{b}$.

In this case, let us consider the following utility function $u(a)$:

- We take $u(a) = 0$ for $a < \bar{b}$, and
- We take $u(a) = 1$ for $a \geq \bar{b}$.

Let us compute the right- and left-hand sides of Condition 1. for this utility function.

- Here, $\bar{b} \geq \bar{b}$, so $u(\bar{b}) = 1$. Thus, the right-hand side is greater than or equal to $\alpha_H \cdot u(\bar{b}) = \alpha_H > 0$.
- In this case, we have $\bar{a} < \bar{b}$, so $u(\bar{a}) = 0$. Since $\underline{a} \leq \bar{a}$ and $\bar{a} < \bar{b}$, we have $\underline{a} < \bar{b}$ and thus, $u(\underline{a}) = 0$. Thus, the left-hand side of the Inequality 1. is equal to 0.

Thus, the Inequality 1. is not satisfied.

The proposition is proven.

3 Auxiliary Results

Discussion. In the previous section, we considered the case when we know α_H and but we do not know $u(a)$. It is reasonable to consider two other cases:

- When we know $u(a)$ but we do not know α_H ; and
- When we do not know neither α_H nor $u(a)$.

In this case, we can prove two similar results.

Comment. In the first case, we have to additionally assume that the utility function $u(a)$ is strictly increasing, i.e., that $a < b$ implies that $u(a) < u(b)$. For such functions, not only $a \geq b$ implies that $u(a) \geq u(b)$, but also, vice versa, $u(a) \geq u(b)$ implies that $a \geq b$ – because otherwise, if we had $a < b$, then we would have $u(a) < u(b)$.

First auxiliary result. *Let $u(a)$ be a strictly increasing function. Then, for every two intervals $[\underline{a}, \bar{a}]$ and $[\underline{b}, \bar{b}]$, the following two conditions are equivalent:*

1. *for all $\alpha_H \in [0, 1]$, we have*

$$\alpha_H \cdot u(\bar{a}) + (1 - \alpha_H) \cdot u(\underline{a}) \geq \alpha_H \cdot u(\bar{b}) + (1 - \alpha_H) \cdot u(\underline{b});$$

2. *$\underline{a} \geq \underline{b}$ and $\bar{a} \geq \bar{b}$.*

Proof.

1°. Similarly to the proof of the main result, we can prove that the Condition 2. implies the Condition 1.

2°. Let us now assume that the Condition 1. is satisfied for all $\alpha_H \in [0, 1]$. In particular, this means that this condition is satisfied for $\alpha_H = 0$ and for $\alpha_H = 1$.

- For $\alpha_H = 0$, the Condition 1. means that $u(\underline{a}) \geq u(\underline{b})$. Since the function $u(a)$ is strictly increasing, this implies that $\underline{a} \geq \underline{b}$.
- For $\alpha_H = 1$, the Condition 1. means that $u(\bar{a}) \geq u(\bar{b})$. Since the function $u(a)$ is strictly increasing, this implies that $\bar{a} \geq \bar{b}$.

Thus, the Condition 2. is indeed satisfied.

The proposition is proven.

Second auxiliary result. *For every two intervals $[\underline{a}, \bar{a}]$ and $[\underline{b}, \bar{b}]$, the following two conditions are equivalent:*

1. *for all $\alpha_H \in (0, 1)$ and for all non-strictly increasing functions $u(a)$, we have*

$$\alpha_H \cdot u(\bar{a}) + (1 - \alpha_H) \cdot u(\underline{a}) \geq \alpha_H \cdot u(\bar{b}) + (1 - \alpha_H) \cdot u(\underline{b});$$

2. *$\underline{a} \geq \underline{b}$ and $\bar{a} \geq \bar{b}$.*

Proof.

1°. Similarly to the proof of the main result, we can prove that the Condition 2. implies the Condition 1.

2°. Let us now assume that the Condition 1. is satisfied for all $\alpha_H \in (0, 1)$. In particular, this means that this condition is satisfied for $\alpha_H = 0.5$. Then, by our Main result, we can conclude that the Condition 2. is satisfied.

The proposition is proven.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

References

1. P.C. Fishburn, *Utility Theory for Decision Making* (John Wiley & Sons Inc., New York, 1969)
2. P.C. Fishburn, *Nonlinear Preference and Utility Theory* (The John Hopkins Press, Baltimore, Maryland, 1988)
3. L. Hurwicz, *Optimality Criteria for Decision Making Under Ignorance*, Cowles Commission Discussion Paper, Statistics, No. 370, 1951
4. V. Kreinovich, Decision making under interval uncertainty (and beyond), in *Human-Centric Decision-Making Models for Social Sciences*. ed. by P. Guo, W. Pedrycz (Springer Verlag, 2014), pp.163–193
5. R.D. Luce, R. Raiffa, *Games and Decisions: Introduction and Critical Survey* (Dover, New York, 1989)
6. H.T. Nguyen, O. Kosheleva, V. Kreinovich, “Decision making beyond Arrow’s ‘impossibility theorem’, with the analysis of effects of collusion and mutual attraction”. *International Journal of Intelligent Systems* **24**(1), 27–47 (2009)
7. H.T. Nguyen, V. Kreinovich, B. Wu, G. Xiang, *Computing Statistics under Interval and Fuzzy Uncertainty* (Springer Verlag, Berlin, Heidelberg, 2012)
8. H. Raiffa, *Decision Analysis* (McGraw-Hill, Columbus, Ohio, 1997)

Number Representation With Varying Number of Bits



Anuradha Choudhury, Md Ahsanul Haque, Saeefa Rubaiyet Nowmi, Ahmed Ann Noor Ryen, Sabrina Saika, and Vladik Kreinovich

Abstract In a computer, usually, all real numbers are stored by using the same number of bits: usually, 8 bytes, i.e., 64 bits. This amount of bits enables us to represent numbers with high accuracy—up to 19 decimal digits. However, in most cases—whether we process measurement results or whether we process expert-generated membership degrees—we do not need that accuracy, so most bits are wasted. To save space, it is therefore reasonable to consider representations with varying number of bits. This would save space used for representing numbers themselves, but we would also need to store information about the length of each number. In view of this, the first natural question is whether a varying-length representation can lead to a drastic decrease in needed computer space. Another natural question is related to the fact that while potentially, allowing number of bits which is not proportional to 8 bits per byte will save even more space, this would require a drastic change in computer architecture, since the current architecture is based on bytes. So will going from bytes to bits be worth it—will it save much space? In this paper, we provide answers to both questions.

A. Choudhury · M. A. Haque · S. R. Nowmi · A. A. N. Ryen · S. Saika · V. Kreinovich (✉)
University of Texas at El Paso, El Paso, TX, USA
e-mail: vladik@utep.edu

A. Choudhury
e-mail: achoudhury@miners.utep.edu

M. A. Haque
e-mail: mhaque3@miners.utep.edu

S. R. Nowmi
e-mail: snowmi@miners.utep.edu

A. A. N. Ryen
e-mail: aryen@miners.utep.edu

S. Saika
e-mail: ssaika@miners.utep.edu

1 Formulation of the Problem

How real numbers are represented in a computer now. Computers usually use the same number of bits to store all real numbers: 64 bits, which is 8 bytes.

This length potentially enables us to represent real numbers with relative precision up to 2^{-64} , which is approximately 10^{-19} .

But do we need that many bits? In some cases, we *do* need this precision—and sometimes, we even need double precision corresponding to 128 bits. However:

- In many practical situations, we process measurement results, that are measured with precision 1% or even less; see, e.g., [9]; in such cases, we do not need that many bits, so additional bits are simply wasted.
- In other practical situations, we process degrees of confidence provided by experts; see, e.g., [1, 5–8, 10]. These degrees are known even with less accuracy, usually the accuracy about 10%, so even more bits are wasted when we use the usual computer representation of real numbers.

It is therefore reasonable to consider number representations with *varying* number of bits. Such representations have indeed been proposed; see, e.g., [2, 3].

Resulting questions. The above idea—of using number representations with varying number of bits—leads to several questions.

First question: how much space will we save? The first question is: how much space do we save?

- On the one hand, we need fewer bits to store each number.
- On the other hand, we will need to store, for each number, information about its length—which also takes a few bits.

Second question: shall we abandon byte structure? The second question is related to the fact that in a computer, bits are usually organized into bytes. From this viewpoint, it is easier to design a computer in which real numbers can use 1, 2, etc. bytes than to allow also any number of bits—including the number of bits that does not divide by 8 and thus, does not constitute several bytes.

The bit-representation complication may be worth it, if it allows us to save a significant portion of memory. So, the second question is: how much memory do we save if we use bits and not bytes?

What we do in this paper. In this paper, we provide answers to both questions.

2 How Do We Answer These Questions: Our Methodology

Challenge. Strictly speaking, to answer these questions, we need to know how frequently we encounter numbers of different length. At present, this information is not available. So what can we do?

Let us use Laplace Indeterminacy Principle. In such situations, when we do not know probabilities of different situations, it makes sense to use what is called *Laplace Indeterminacy Principle* (see, e.g., [4]):

- Since we have no reason to believe that some situations (in our case, some lengths) lengths are more frequent than others,
- It makes sense to assume that all possible lengths are equally frequent.

Let us use this assumption to answer both questions.

3 Case of Byte Representation

What are the options? In the byte representation, a number can occupy 1, 2, ..., 8 bytes.

What are the probabilities of these options? In line with the Laplace Indeterminacy Principle—as described in the previous section—we assume that each of these eight cases has the same probability $1/8$.

In this case, what is the average number of bits taken by a real number. Based on the lengths and probabilities of different options, we can conclude that the average length of the real number is equal to

$$\frac{1 + 2 + \dots + 8}{8} = 4.5 \text{ bytes} = 36 \text{ bits.}$$

How many additional bits do we need to store information about each number's length. In general, with k bits, we can describe 2^k different options. In our case, we have 8 possible options, and $8 = 2^3$. Thus, to store information about the length, we need 3 bits.

How many bits per number do we need overall. For each real number:

- We need, on average, 36 bits to store the number itself, and
- We need 3 bits to store the information about the number's length.

So overall, we need $36 + 3 = 39$ bits.

How much space do we save this way. The resulting amount of 39 bits is much smaller than the current 64 bits—about 40% smaller.

So, is this worth doing? Since this representation can potentially save a lot of space, the answer to the first question is: yes, number representation with varying number of bytes is worth pursuing.

4 Case of Bit Representation

What are the options? In the bit representation, a number can occupy 1, 2, ..., 64 bits.

What are the probabilities of these options? In line with the Laplace Indeterminacy Principle—as described earlier—we assume that each of these 64 cases has the same probability $1/64$.

In this case, what is the average number of bits taken by a real number. Based on the lengths and probabilities of different options, we can conclude that the average length of the real number is equal to

$$\frac{1 + 2 + \dots + 64}{64} = 32.5 \text{ bits.}$$

How many additional bits do we need to store information about each number's length. In general, with k bits, we can describe 2^k different options. In our case, we have 64 possible options, and $64 = 2^6$. Thus, to store information about the length, we need 6 bits.

How many bits per number do we need overall. For each real number:

- we need, on average, 32.5 bits to store the number itself, and
- we need 6 bits to store the information about the number's length.

So overall, we need $32.5 + 6 = 38.5$ bits.

How much space do we save this way. The difference between the resulting average number of 38.5 and 39 bits corresponding to byte representation is very small—about 1%.

So, is this worth doing? The bit representation requires a drastic redesign of computer architecture but saves only a tiny amount of space. So, the answer to the second question is: no, it is probably not worth doing.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

References

1. R. Belohlavek, J.W. Dauben, G.J. Klir, *Fuzzy Logic and Mathematics: A Historical Perspective* (Oxford University Press, New York, 2017)
2. M. Ceberio, C. Lauter, V. Kreinovich, Just-in-accuracy: mobile approach to uncertainty. *J. Mob. Multimed.* **20**(3) (2024)
3. J.F. Gustafson, *The End of Error: Unum Computing* (Chapman & Hall/CRC, Boca Raton, Florida, 2015)
4. E.T. Jaynes, G.L. Bretthorst, *Probability Theory: The Logic of Science* (Cambridge University Press, Cambridge, UK, 2003)
5. G. Klir, B. Yuan, *Fuzzy Sets and Fuzzy Logic* (Prentice Hall, Upper Saddle River, New Jersey, 1995)
6. J.M. Mendel, *Explainable Uncertain Rule-Based Fuzzy Systems* (Springer, Cham, Switzerland, 2024)
7. H.T. Nguyen, C.L. Walker, E.A. Walker, *A First Course in Fuzzy Logic* (Chapman and Hall/CRC, Boca Raton, Florida, 2019)
8. V. Novák, I. Perfilieva, J. Močkoř, *Mathematical Principles of Fuzzy Logic* (Kluwer, Boston, Dordrecht, 1999)
9. S.G. Rabinovich, *Measurement Errors and Uncertainty: Theory and Practice* (Springer Verlag, New York, 2005)
10. L.A. Zadeh, Fuzzy sets. *Inf. Control* **8**, 338–353 (1965)

Fuzzy Neural Networks

How to Make a Neural Network Learn from a Small Number of Examples—and Learn Fast: An Idea



Chitta Baral and Vladik Kreinovich

Abstract Current deep learning techniques have led to spectacular results, but they still have limitations. One of them is that, in contrast to humans who can learn from a few examples and learn fast, modern deep learning techniques require a large amount of data to learn, and they take a long time to train. In this paper, we show that neural networks do have a potential to learn from a small number of examples—and learn fast. We speculate that the corresponding idea may already be implicitly implemented in Large Language Models—which may partially explain their (somewhat mysterious) success.

1 Formulation of the Problem

What is machine learning: a brief reminder. In many practical situations, we want to learn how a quantity y depends on the quantities $x_1, \dots, x_n$. We have several (K) examples in which we know the values both of the values x_i and y , i.e., for each $k = 1, \dots, K$ we know the *patterns* $(x_1^{(k)}, \dots, x_n^{(k)}, y^{(k)})$ for which $y^{(k)} \approx f(x_1^{(k)}, \dots, x_n^{(k)})$ for the unknown function $f(x_1, \dots, x_n)$. Based on such information, we want to reconstruct the desired function $f(x_1, \dots, x_n)$. This will enable us in the future cases when we know the values $x_1, \dots, x_n$, to predict y as $f(x_1, \dots, x_n)$.

Main limitation of current machine learning techniques. Deep learning algorithms have spectacular successes (see, e.g., [2]), they enable computers to play chess and Go much better than humans, they help us solve many practical problems. However, in comparison to humans, the current deep learning techniques have a severe limitation:

C. Baral

Department of Computer Science and Engineering, Arizona State University, Tempe, AZ, USA
e-mail: chitta@asu.edu

V. Kreinovich (✉)

Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA
e-mail: vladik@utep.edu

- To a human being, it is sufficient to have a few examples (patterns), and without practically any delay we can learn to distinguish cats from dogs, etc.;
- In contrast, a deep neural network requires a large number of examples—thousands and even millions—to start producing reasonable results, and even on modern super-fast high-performance computers, it takes a long time to train a neural network.

Comment. To be fair, it should be mentioned that while a neural network usually takes a very long time to *train*, once it is trained, it *produces its results very fast*. For example, even in many applications involving solution to differential equations, where algorithms are known, it is now much faster to use a trained neural network to produce the solution than to use the known algorithms.

What we do in this paper. In view of the above limitation, it is desirable to come up with a machine learning tool that will enable us to learn from a small number of examples—and to learn fast, just like we humans do. In this paper, we describe a proposal for designing such a tool.

Comments.

- We also speculate that this may already be happening for in-context learning systems such as GPT3 and ChatGPT—that produce very reasonable answers to queries already after 5–10 examples.
- Preliminary results from this paper first appeared in [1].

2 What We Want

Let us describe what we want in precise terms. What we would like to have is a *universal* computer system that:

- Given a small number of examples from any area and a new input,
- Would generate a reasonable answer to this new input.

In mathematical terms, this means that this system should take, as an input, the tuple X that contain all the input information, i.e.,

$$X = (x_1^{(1)}, \dots, x_n^{(1)}, y^{(1)}, \dots, x_1^{(K)}, \dots, x_n^{(K)}, y^{(K)}, x_1, \dots, x_n) \quad (1)$$

and this system should return the value y corresponding to the input $x_1, \dots, x_n$. For example:

- We should show this system several images of a cat (i.e., images $x_1^{(k)}, \dots, x_n^{(k)}$ for which $y^{(k)} = 1$), several images of other animals (i.e., images $x_1^{(k)}, \dots, x_n^{(k)}$ for which $y^{(k)} = 0$), and a new image $x_1, \dots, x_n$,

- And the system should decide whether this new image is the image of a cat ($y = 1$) or of some other animal ($y = 0$).

Also:

- We should show, to this same system, many images of vehicles indicating which of them are trucks and which are not trucks, and then show this system a new picture of a vehicle,
- And this system should tell us whether this new image is truck or not a truck.

3 Analysis of the Problem

We know that we humans have this universal ability:

- Given the corresponding input X ,
- To produce the corresponding value y —whether it is about cars, about cars, or about something else.

In other words, we know that the values forming the tuple X largely uniquely determine the desired value y . In mathematical terms, as we have mentioned, this means that there exists a function $y = F(X)$ that takes, as input, a tuple of type (1) and returns the corresponding value y .

This formulation leads to the following natural idea.

4 Resulting Idea

Natural idea. We have an unknown function $F(X)$. We know that neural networks are universal approximators, i.e., that for each reasonable function and each desired accuracy, there exists a neural network that approximates the given function with the desired accuracy.

So why not use a neural network to approximate this unknown function $F(X)$?

Important comment. In the above statement, we overlooked a somewhat minor but important point: that the universal approximation theorem was proven for the case when the number of inputs is fixed. In our case, the number of inputs N forming a tuple X may be different depending on how many examples K we have and how many inputs n are there in each example:

- We have K examples with $n + 1$ values in each of them, and
- We have n values of a new example,

to the total of $N = K \cdot (n + 1) + n$.

So, to be able to apply the universal approximation result, let us limit ourselves to the cases when we have a fixed number of examples K and the fixed number of inputs n in each example.

For example, we can limit ourselves to cases when we have $K = 10$ examples with $n = 4$ inputs each. Then, in each such case, we have tuples X with

$$N = 10 \cdot (4 + 1) + 4 = 54$$

inputs.

Resulting proposal. Suppose that we want to design a computer system that will learn—in all possible application areas—after being presented K examples. Ideally, we should make this system really universal, i.e., it should be applicable to all possible input sizes n . However, in this text, what we are proposing is to do *almost* this—namely, to design a system that only works for situations when we have a fixed number of inputs n .

To design such a system, we do the following.

- First, in each of many different application areas, we collect a large number P of patterns $(x_1^{(i)}, \dots, x_n^{(i)}, y^{(i)})$, $i = 1, \dots, P$ corresponding to this particular area. In most application areas, this is already done, we already have many such example.
- Then, for each application area, many times, we select $K + 1$ of these patterns $i_1, \dots, i_K, i_{K+1}$ and use, as X , the first K of these patterns and the inputs corresponding to the $(K + 1)$ -st pattern:

$$X = (x_1^{(i_1)}, \dots, x_n^{(i_1)}, y^{(i_1)}, \dots, x_1^{(i_K)}, \dots, x_n^{(i_K)}, y^{(i_K)}, x_1^{(i_{K+1})}, \dots, x_n^{(i_{K+1})}).$$

As the value $y = F(X)$ correspond to this tuple X , we use the value y for the $(K + 1)$ -st pattern, i.e., $y = y^{(i_{K+1})}$.

- Finally, we use all the resulting pairs (X, y) —corresponding to different application areas and to different selection of $K + 1$ patterns within each area—to train a neural network that will produce, in all application areas, y based on X .

What we expect. Once trained, this neural network will transform each tuple X of type (1) into an appropriate value y . In other words, it will:

- Take a small number K patterns from some application area and an input $x_1, \dots, x_n$, and
- Generate the output corresponding to what we expect in this particular application area.

In other words, this system will do exactly what we wanted: produce reasonable answer after a small number K of examples.

How fast will this system be? Training this neural network may take forever. However, as we have mentioned earlier, once this neural network is trained, it will

produce its results really fast. In other words, not only will the resulting system learn from a small number of examples—it will also produce its results practically immediately, just like we human do. Again, this is exactly what we wanted.

Speculative comment. We were discussing how to design a neural network that can learn from few examples—and learn fast. But there are already networks that seems to be doing this—namely, modern Large Linguistic Models like GPT3 and ChatGPT. This ability of such models is largely a mystery. So maybe the partial explanation of their success is that these models are, in effect, already implementing the above idea? (Of course, it does not explain their zero-shot ability, when we get an answer without having any examples to train on.)

Indeed, in contrast to the usual neural network that is usually trained on examples from one application area, these models are trained on all kinds of language examples, i.e., examples from all possible application areas—and, as we have argued, such training is one of the main features that can lead to such training-fast-on-few examples.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395, and by the AT&T Fellowship in Information Technology.

It was also supported by the program of the development of the Scientific-Educational Mathematical Center of Volga Federal District No. 075-02-2020-1478, and by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

The authors are thankful to all the participants of the 2023 IEEE International Conference on Systems, Man, and Cybernetics SMC 2023 (Honolulu, Hawaii, October 1–4, 2023) for valuable discussions.

References

1. C. Baral, V. Kreinovich, How to make a neural network learn from a small number of examples—and learn fast: an idea, in *Abstracts of the 2023 IEEE International Conference on Systems, Man, and Cybernetics SMC 2023* (Honolulu, Hawaii, 2023)
2. I. Goodfellow, Y. Bengio, A. Courville, *Deep Learning* (MIT Press, Cambridge, Massachusetts, 2016)

Analysis of Histological Preparations Using Convolutional Neural Networks



Guskov G. Yu , N. G. Yarushkina , and A. V. Chekina

Abstract The article proposes an approach to the analysis of histological preparations using neural networks. Primary tasks are the detection of nests in nevus tissues using a convolutional neural network. Artificial intelligence-based decision support in the field of medicine is a particularly promising area of research in oncology. The main task of the study was to mark up the image and find characteristic objects on it, such as the cell nucleus, nevus node, etc. It was also necessary to take into account the relative position of objects e.g. the depth of the nevus relative to the dermis. The presented results of the models were evaluated for accuracy and completeness relative to expert assessments. The neural network used in the study was previously trained to isolate cell nuclei on a histological preparation. After that, these weights were adjusted taking into account the search for the necessary objects, namely nests of nevus cells. The article discusses examples of images that serve as a training sample, the configuration and parameters of the neural network.

Keywords Convolutional neural networks · Machine learning · Artificial intelligence · Histological preparations

1 Introduction

1.1 The Problem of Diagnosing Skin Melanoma

Skin melanoma is an aggressive tumor, ranked fifth among malignant tumors [1], with a high mortality rate, significant metastatic potential, and low efficacy of therapy in the late stages of the disease.

Worldwide, skin melanoma causes 55,500 deaths (mortality rate 0.7%) each year [2]. In addition, an increase in the incidence and poor prognosis of melanoma accounts for 15% of the 5-year survival rate [3, 4].

G. G. Yu · N. G. Yarushkina (✉) · A. V. Chekina
Ulyanovsk State Technical University, Ulyanovsk, Russia
e-mail: jng@ulstu.ru

In melanomas arising from pre-existing nevi, residual original nevi are usually detected histologically. In patients with multiple nevi, the percentage of transformation into melanoma increases to 50%, especially under conditions of increased exposure to UV radiation [5].

Diagnosis and evaluation of skin lesions is largely based on histopathological criteria in predicting the transformation of a nevus into melanoma [6]. Nevus is a skin neoplasm in a broad sense, such as a mole or birthmark. Over time, a nevus can degenerate into a malignant neoplasm—melanoma. A nevus can be analyzed based on a histological section after removal. If, during the analysis of a histological preparation, specialists find characteristic signs of a nevus degenerating into melanoma, this can serve as a significant argument when making a diagnosis. One of the most characteristic signs of nevus degradation is the presence of nevus nodes. A nevus node is a multitude of cells inside a single membrane. Over time, a nevus node will collapse and cancer cells will enter the bloodstream, which will lead to metastases. Automation of the search and determination of nevus nodes on histological preparations is a relevant and important task given the large volumes of digital images of the histological preparation. Research in this area together with medical specialists will potentially allow us to determine earlier changes in the nevus. In this regard, it is extremely important to identify the pathological process at an early stage with the ability to predict the likelihood of nevus transformation into melanoma, since successful treatment is possible only if melanoma is detected in the early stages of oncogenesis. Moreover, there is a relationship between clinical, morphological and molecular genetic parameters in the diagnosis of melanoma, taking into account melanoma in a family history [7]. AI-based clinical decision support is a particularly promising area of research [8].

Computer analysis of images of histological preparations of skin formations is also used to establish the nosological group [9].

Deep learning methods can be used to recognize morphological patterns in a sample for diagnostic and classification purposes [10]. The use of deep convolutional neural networks (CNN) technology has shown great potential as a tool to aid the digital analysis of skin lesions [11].

The main goal of the work is to form a data set for the analysis of histological preparations using convolutional neural networks.

1.2 Classification of Neoplasms on the Skin

Skin neoplasms of the skin are characterized by benign or malignant evolution. Skin neoplasms are present on the skin of the vast majority of people and can cause physical and moral discomfort. So benign skin formations include such as atheroma, hemangioma, lymphangioma, lipoma or wen, warts, papilloma, fibroma, neurofibroma and nevi.

The skin consists of three main layers: the epidermis, dermis, and subcutaneous fat, each of which is also divided into several layers.

The epidermis is the outer visible layer of the skin. The epidermis protects a person from toxic substances, harmful microorganisms, and fluid loss. It has five layers consisting of keratinocyte cells. These cells originate in the lowest, basal layer of the epidermis and gradually move to the surface of the skin. As they move, they mature and change. This process is called keratinization or keratinization.

The dermis is the middle layer of the skin: relatively wide, elastic, but dense. It consists of two inner layers. Lifestyle and external factors such as solar radiation and temperature changes have a destructive effect on collagen and elastin found in the dermis.

Subcutaneous fat softens blows and retains body heat. It consists of fat cells, collagen fibers and blood vessels.

It has been established that some types of nevi have a special tendency to malignancy. These include:

- Border (junctional) nevus is located in the basal, boundary layer of the epidermis and dermis. The penetration of its structures into the dermis leads to the formation of a complex and intradermal nevus;
- Complex nevus—characterized by the location of its structures both in the epidermis and in the dermis. The evolution of a complex nevus is characterized by processes: transformation into an intradermal one; spontaneous regression; transformation into melanoma;
- Blue nevus—is intradermal. With the development of a blue nevus, two processes are traced: fibrosis and proliferation of melanocytes;
- Cellular blue nevus—biologically more active and differs from a simple blue nevus by a pronounced proliferation of melanocytes, which increases the risk of its development into melanoma;
- Giant pigmented nevus—is most often congenital and increases with the growth of the child, a flat papillary surface, occupies significant areas of the skin of the body. The neoplasm usually has the histological structure of a complex nevus;
- Dysplastic nevi—this type includes epidermal and mixed nevi. Patients with dysplastic nevus syndrome have an increased risk of developing melanomas, including the development of primary multiple tumors.

Precursors of melanoma: congenital non-cellular nevus (giant), dysplastic nevus (giant or small), blue nevus, malignant lentigo [5].

The complex structure of the skin and the broad classification of nevi require constant interaction with specialists to form datasets for training machine learning models.

1.3 Deep Learning Diagnostic Programs for Histological Assessment of Skin Melanoma

Deep learning models can recognize objects in images and other forms of data that might go unnoticed by an expert. The discovery of these new features opens the possibility of using neural network AI models for more efficient clinical decision making by clinicians.

In recent years, there has been active research into deep learning applications for diagnosing and predicting melanoma based on images of histological slides stained with hematoxylin and eosin. Due to advances in information systems, as well as methods for collecting and processing images, data, in recent years there has been a growing interest in machine learning for biomedical research [12].

The study dataset size ranges from 9 to 18,607 images, with an average of 324 images. Sometimes entire images are used, while other studies rely on datasets consisting of smaller image fragments.

Deep learning models using the Convolutional Neural Network (CNN) have been used in such problems more often than other solutions, and showed high results.

Most diagnostic programs designed to differentiate melanoma from nevi are based on the processing of histological images by AI models [9, 12–14]. There are applications for diagnosing melanoma, nevi and normal skin [10] and applications for distinguishing between melanoma and non-melanoma skin cancer [15–17]. Some studies have used deep learning to isolate entire tumor areas [3] or for individual diagnostic markers such as the presence of mitotic cells [18], melanocytes [2], and melanocytic nests [19].

2 Deep Convolutional Neural Network Mask R-CNN

2.1 Mask R-CNN Neural Network

Mask R-CNN is an object detection model based on deep convolutional neural networks (CNN) developed by the Facebook AI research group in 2017. Mask R-CNN is an extension of Faster R-CNN and works by adding a branch for object mask (region of interest) prediction in parallel with the existing branch for bounding box recognition. Benefits of the Mask R-CNN: simplicity, performance, efficiency and flexibility. Mask R-CNN is easy to train on the basis of existing weights to solve problems similar to those already solved.

Mask R-CNN is a model built to compute objects across an entire image without using the traditional sliding window approach used in most object detection methods. It works quite fast compared to traditional sliding window methods. In addition to the benefits of Mask R-CNN offers ample information about each detected object.

To support the Mask R-CNN model with more popular libraries such as TensorFlow, there is a popular open source project called Mask_RCNN which offers an implementation based on Keras and TensorFlow 1.14, which was taken as the basis.

This project was designed to classify objects in an image, which is different from the task of pattern recognition research. Therefore, the system of object classes was excluded from the model, and the algorithm was adapted for recognition.

At the time of the beginning of this study, the solution of such problems using Mark-CNN was the most relevant choice, since it showed excellent results in the task of searching for cell nuclei on histological preparations [20]. The solution that proved to be effective in the competition and was taken as a basis. The neural network was trained both on open data from the competition and data on the selection of another type of objects—nevus nests, which were formed by an expert of the research center.

2.2 Automated Process for Evaluating Histological Preparations in Diagnosing Melanoma

In case of suspicion of the degeneration of a nevus into melanoma during non-invasive observations by a specialist due to the evolution (ABCDE analysis) of a neoplasm, it is removed. The result of the removal of the neoplasm is a histological preparation that allows specialists to evaluate the cellular structure and identify nevus nests.

The specialist examines the histological preparation using a microscope (Axio Imager M2—Carl Zeiss) and also saves the image in the MRXS format. The number of such preparations is large and the preparation itself can be quite large for testing for the presence of nevus nests. That is why it is necessary to use an automated system that will help to focus the attention of a specialist on potentially dangerous areas of the image. Nevertheless, the decision maker always works with the results of an automated system. Systems of this purpose never work in automatic mode.

The proposed solution, based on convolutional neural networks, the definition of nevus nodes presented in this paper is part of a more complex system (Fig. 1).

3 Extending the Dataset for Training a Neural Network

Digitized histological preparations of skin formations were made at the research center “UIGPU named after A.I. I.N. Ulyanova” with registry maintenance, registration and coding of patients. The skin lesions were delivered to the Scientific Research Center as clinical, previously surgically removed, fixed in 10% neutral buffered formalin, material from hospitals, an oncology dispensary, and polyclinics in Ulyanovsk and the Ulyanovsk region.

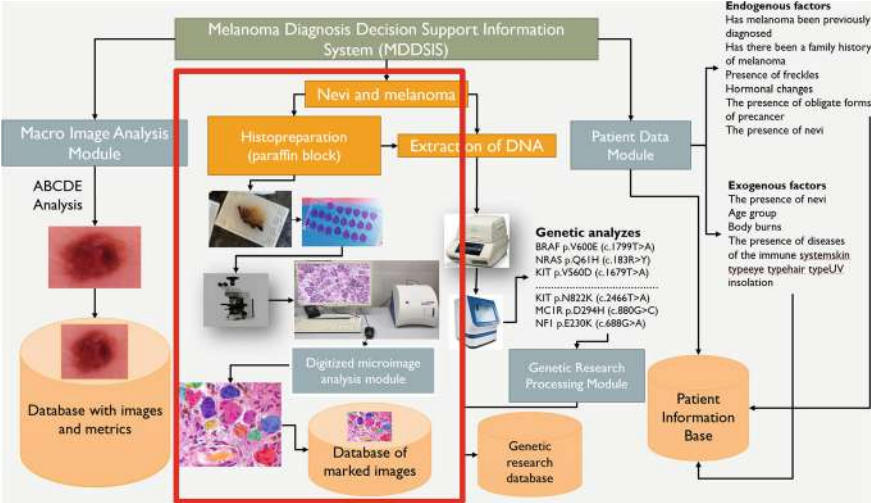


Fig. 1 Melanoma diagnosis decision support information system (MDDISIS) structure

The preparation of histological preparations included the steps of dehydration, paraffin impregnation, embedded in paraffin, serial sections were made, then the sections were stained with Mayer’s hematoxylin and eosin. Microdescription was carried out using a research motorized microscope Axio Imager M2 (Carl Zeiss) with an imaging system: a color digital camera—Axioacam 105 (Carl Zeiss).

Histological preparations were digitized using a digital scanning microscope Panoramic Desk (3DHISTECH, Hungary) using whole slide imaging technology (WSI, full slide images).

The scanning microscope exports digitized images of formations from histological preparations in the form of complex files/slides in the MRXS format. At the same time, slides in the MRXS format have a complex structure and large dimensions. Therefore, the image format obtained during the scan, microscopic images cover the entire skin formation, such an image file has a large weight, reaching up to 1–1.5 gigabytes. Further, the research center experts designated nevus nests, nuclei of nevus cells (Fig. 2).

Further, masks in Fig. 3 were formed for unmarked samples. The dataset that was generated within the study was based on 6 digitized histological preparations (WSI full slide images) with a diagnosis of intradermal nevus of the skin. The image data was cut into 328 fragments for Mark-CNN training. In these images, 1496 nevus nests were found, and a mask image was formed for each of them. Also, using the solution (based on Mark-CNN) from the competition [20], more than 160,000 cell nuclei were identified on the generated dataset, which indicates the absolute need for automation of the analysis of histological preparations.

The results of experiments on training a neural network on a new generated dataset for searching for nevus nests are presented in Table 1.

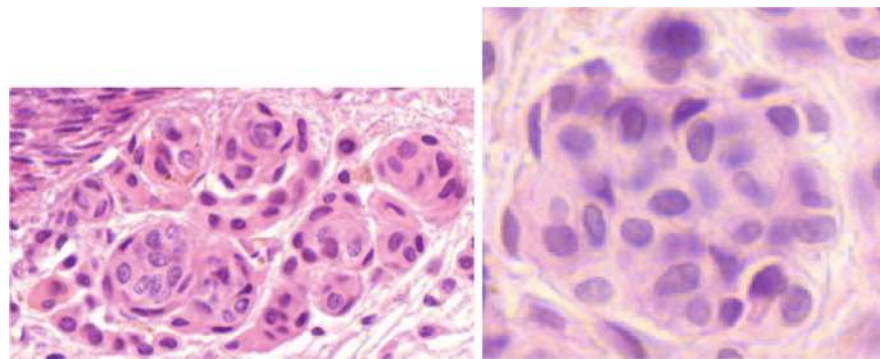


Fig. 2 Nevus nests(left), nuclei of nevus cells (right)



Fig. 3 Mask for nest from Fig. 2

Table 1 The results of the classification of nevus nests

No.	Number of nevus nodes according to expert	The number of nevus nodes according to the neural network classifier	The number of nevus nodes was detected correctly by the neural network	Accuracy	Recall
1	3	5	2	40	66.6
2	3	5	3	60	100
3	2	4	2	50	100
4	2	3	2	66.6	100
5	9	8	7	87.5	77.7
6	5	4	3	75	60
7	6	4	4	100	66.6
8	6	7	5	71.4	83.3
9	5	4	3	75	60
10	3	7	2	28.5	66.6
Average				65.4	78.1

Automated analysis of histological specimens using artificial intelligence methods can be applied to a number of tasks. The tasks are usually very specific to the laboratory, disease, medical standards and requirements.

The study [21] provides evidence that Mask R-CNN is the most effective tool for detecting overlapping objects on histological specimens.

Convolutional neural networks are used for segmentation of histological images [22].

In this study, a well-known neural network model was modified to solve the applied practical problem of automating histological preparations. The novelty of the study is that additional training of the model to solve this problem became possible as a result of collecting and marking a dataset of laboratory materials. The dataset is also a valuable result of the work, based on which it is possible to test the effectiveness of other image segmentation models.

4 Conclusion

The table shows the estimates of the quality of the neural network on the dataset for 10 images. The fact is that only an expert is able to evaluate whether a nevus node has been identified correctly or not. That is why the sample for evaluation is represented by an analysis of 10 images.

It may seem that the indicators of accuracy and recall are low for such tasks. But in fact, the system solves the main task of diagnosis, since if a nevus node confirmed by an expert was found, it confirms the diagnosis. In addition, this task involves a more complex solution for research purposes. In the future, it is necessary to take into account the depth, layer, size, ratio of the space of the nucleus to the space of the cell, ruptures of nevus nodes, and much more.

The main result is the possibility of effective application of AI models, namely convolutional neural networks in such tasks. In addition, the main result of the study can be considered a labeled data set, which is a key value of the work.

Acknowledgements This study was supported the Ministry of Science and Higher Education of Russia in framework of project No. 075-03-2023-143 «The study of intelligent predictive analytics based on the integration of methods for constructing features of hetero-geneous dynamic data for machine learning and methods of predictive multimodal data analysis».

References

1. S. Sankarapandian, S. Kohn, V. Spurrier, S. Grullon, R.E. Soans, K.D. Ayyagari, R.V. Chamarthi, K. Motaparathi, J.B. Lee, W. Shon, et al., A pathology deep learning system capable of triage of melanoma specimens utilizing dermatopathologist consensus as ground truth, in *Proceedings of the 2021 IEEE/CVF International Conference on Computer Vision Workshops*

- (ICCVW), Montreal, BC, Canada, 11–17 October 2021 (IEEE Computer Society, Los Alamitos, CA, USA, 2021), pp. 629–638
2. H. Pehamberger, A. Steiner, K. Wolff, In vivo epiluminescence microscopy of pigmented skin lesions. Pattern analysis of pigmented skin lesions. *J. Am. Acad. Dermatol.* **17**, 571–583 (1987)
 3. C. Lu, M. Mandal, Automated analysis and diagnosis of skin melanoma on whole slide histopathological images. *Pattern Recognit.* **48**, 2738–2750 (2015)
 4. V. Nikolaou, A.J. Stratigos, Emerging trends in the epidemiology of melanoma. *Br. J. Dermatol.* **170**(1), 11–19 (2014). <https://doi.org/10.1111/bjd.12492>. (PMID: 23815297)
 5. D.E. Elder, G.F. Murphy, *Melanocytic Tumors of the Skin. Atlas of Tumor Pathology—3rd Series.—Fasc.2.* (Armed Forces Institute of Pathology, Washington, 1991), p. 216
 6. P. Tschandl, C. Rinner, Z. Apalla, G. Argenziano, N. Codella, A. Halpern et al., Human-computer collaboration for skin cancer recognition. *Nat. Med.* **26**, 1229–1234 (2020)
 7. J.G. Greener, S.M. Kandathil, L. Moffat, D.T. Jones, A guide to machine learning for biologists. *Nat. Rev. Mol. Cell Biol.* **23**, 40–55 (2022)
 8. D. Schadendorf, A.C.J. van Akkooi, C. Berking, et al., Melanoma. *Lancet* **392**(10151), 971–984 (2018)
 9. L. Wang, L. Ding, Z. Liu, L. Sun, L. Chen, R. Jia, X. Dai, J. Cao, J. Ye, Automated identification of malignancy in whole-slide pathological images: identification of eyelid malignant melanoma in gigapixel pathological slides using deep learning. *Br. J. Ophthalmol.* **104**, 318–323 (2020)
 10. D. Kucharski, P. Kleczek, J. Jaworek-Korjakowska, G. Dyduch, M. Gorgon, Semi-supervised nests of melanocytes segmentation method using convolutional autoencoders. *Sensors* **20**, 1546 (2020)
 11. R. Friedman, D. Rigel, A. Kopf, Early detection of malignant melanoma: role of physician examination and self-examination of the skin. *CA Cancer J. Clin.* **35**, 130–151 (1985)
 12. S. Chen, A. Geller, H. Tsao, Update on the epidemiology of melanoma. *Curr. Derm. Rep.* **2**, 24–34 (2013). <https://doi.org/10.1007/s13671-012-0035-5>
 13. H. Wei, W.-H. Xu, J.-X. Wang, J.-M. Hou, H.-L. Zhang, X.-Y. Zhao, G.-L. Shen, et al., Identification, validation, and functional annotations of genome-wide profile variation between melanocytic nevus and malignant melanoma. *BioMed Res. Int.* (2020)
 14. P. Xie, K. Zuo, J. Liu, M. Chen, S. Zhao, W. Kang, F. Li, Interpretable diagnosis for whole-slide melanoma histology images using convolutional neural network. *J. Healthc. Eng.* (2021)
 15. L.J. Loeschner, M. Janda, H.P. Soyer, K. Shea, C. Curiel-Lewandrowski, Advances in skin cancer early detection and diagnosis. *Semin. Oncol. Nurs.* **29**, 170–181 (2013)
 16. P.E. LeBoit, G. Burg, D. Weedon, A. Sarasain, *World health organization classification of tumors. Pathology and Genetics of Skin Tumours* (IARC Press, Lyon, 2006), p. 355
 17. K. Liu, M. Mokhtari, B. Li, S. Nofallah, C. May, O. Chang, S. Knezevich, J. Elmore, L. Shapiro, Learning melanocytic proliferation segmentation in histopathology images from imperfect annotations, in *Proceedings of the 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)* (Nashville, TN, USA, 2021), pp. 3761–3770
 18. J. Nyambal, R. Klein, Automated parking space detection using convolutional neural networks, in *Pattern Recognition Association of South Africa and Robotics and Mechatronics (PRASA-RobMech)*, pp. 1–6 (2017)
 19. E. Antonova, G.G. Yu, N.G. Yarushkina, A. Sapunkov, A. Khambikova, Analysis of micro-images of skin neoplasms using convolutional neural networks in an intelligent medical information system for the early diagnosis of melanoma (2022). https://doi.org/10.1007/978-3-031-19620-1_23.
 20. SIIM-ISIC Melanoma Classification. <https://www.kaggle.com/competitions/siim-isic-melanoma-classification/overview>. Accessed 06 June 2023
 21. L. Rettenberger, F. Münke, R. Bruch, M. Reischl, Mask R-CNN outperforms U-Net in instance segmentation for overlapping cells. *Curr. Dir. Biomed. Eng.* **9**, 335–338 (2023). <https://doi.org/10.1515/cdbme-2023-1084>

22. J. Raufeisen, K. Xie, F. Hörst, T. Braunschweig, J. Li, J. Kleesiek, R. Röhrig, J. Egger, B. Leibe, F. Hölzle, A. Hermans, B. Puladi, Cyto R-CNN and CytoNuke dataset: towards reliable whole-cell segmentation in bright-field histological images. *Comput. Methods Programs Biomed.* **252** (2024). <https://doi.org/10.1016/j.cmpb.2024.108215>

Why Two Fish Follow Each Other but Three Fish Form a School: A Symmetry-Based Explanation



Shahnaz Shahbazova, Olga Kosheleva, and Vladik Kreinovich

Abstract Recent experiments with fish has shown an unexpected strange behavior: when two fish of the same species are placed in an aquarium, they start following each other, while when three fish are placed there, they form (approximately) an equilateral triangle, and move in the direction (approximately) orthogonal to this triangle. In this paper, we use natural symmetries—such as rotations, shifts, and permutation of fish—to show that this observed behavior is actually optimal. This behavior is not just optimal with respect to one specific optimality criterion, it is optimal with respect to any optimality criterion—as long as the corresponding comparison between two behaviors does not change under rotations, shifts, and permutations.

1 Introduction

Formulation of the problem. A recent research [1, 11] showed that, when we place several fish of the same species in an aquarium, then:

- If there are only two fish, they follow each other, but
- If there are three fish, they form a school: they place themselves in positions forming (approximately) an equilateral triangle, and move in the direction (approximately) orthogonal to this triangle.

How can we explain this phenomenon?

What we do in this paper. In this paper, we provide a natural symmetry-based explanation for this phenomenon. Namely, we show that the observed behavior is

S. Shahbazova

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics UNEC, Baku, Azerbaijan

e-mail: shahnaz_shahbazova@unec.edu.az; shahbazova@cyber.az

Department of Teacher Education, University of Texas at El Paso, El Paso, TX, USA

O. Kosheleva

Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA

e-mail: olgak@utep.edu

optimal with respect to all optimality criteria that are invariant with respect to natural symmetries: spatial rotations, shifts, and permutations of the fish.

Comment. In this paper, we will not select any specific optimality criterion, we will not specify any numerical model of the fish behavior. Instead, we will start with a kind-of qualitative natural-language description of the situation—namely, the one described in the previous paragraph. Then, we will show how this natural-language description can be translated into a precise result. In this sense, what we are doing is very similar to what Lotfi Zadeh did when he invented fuzzy techniques [2, 4, 5, 7, 8, 10]. While our techniques are different from what Zadeh used, they still fall under the general rubric of translating natural-language descriptions into precise terms, rubric that can be described as—in very general sense—fuzzy.

The structure of this paper. In Sect. 2, we describe a symmetry-related mathematical result that we will use in our explanation. In Sect. 3, we will apply this result to explain the motion of a pair of fish. In Sect. 4, we apply this result to explain the location and motion of three fish.

2 Symmetry-Related Mathematical Result

What do we mean by an optimality criterion. In general, we have the set A of alternatives, and we want to select the optimal (best) one, i.e., the one which is better than—or of the same quality as—every other alternative. Let us denote the relation “better than or of the same quality as” by $\geq$, so that $a \geq b$ would mean that a is better than b or of the same quality as b . In these terms, the optimal alternative a_{opt} is the one for which $a_{\text{opt}} \geq a$ for all a .

Clearly, each alternative is of the same quality as itself, so we have $a \geq a$. Relations satisfying this property are known as *reflexive*. Also, if a is better or of the same quality as b , and b is better or of the same quality as c , then a is clearly either better or of the same quality as c . Relations satisfying this property are called *transitive*. Thus, by an optimality criterion, we will mean a reflexive and transitive relation. Such relations are known as *pre-orders*. So, we arrive at the following definition.

Definition 1 *Let A be a set. Elements of this set will be called alternatives. By an optimality criterion on the set A , we mean a binary relation $\geq$ that satisfies the following two properties:*

- For all a , we have $a \geq a$, and
- For all a , b , and c , if $a \geq b$ and $b \geq c$, then $a \geq c$.

We say that the alternative a_{opt} is optimal if $a_{\text{opt}} \geq a$ for all $a \in A$.

We will only consider final optimality criteria. Our ultimate goal is to select a single alternative. So, if a current optimality criterion has several equally good optimal alternatives, this means that this criterion is not final: we can use this non-uniqueness to optimize something else.

For example, if we select the fastest algorithm for solving some problem, and several different algorithms have the same average computation time, we can select, among them, the one that has the smallest worst-case computation time. If we still have several equally good algorithms, we can use the remaining non-uniqueness to optimize something else.

At the end, we should end up with a *final* optimality criterion for which there is exactly one optimal alternative.

Definition 2 We say that an optimality criterion is final if there is exactly one alternative that is optimal with respect to this criterion.

Natural symmetries. From the physical viewpoint, there are often some transformations that do not change the relative quality of alternatives. For example, if one dish tastes better than another, this relation does not change if we turn and/or shift the table with the taster. In general, let G denotes the set of all such transformations.

If a transformation does not change the relation, the inverse transformation should not change it either: if we move a person 100m North, then moving the same person 100m back South should not affect his/her taste. Similarly, if each of the two transformations does not change the relation, then their composition

- when we first apply the first transformation and then the second transformation
- also should not change the relation. Sets of transformations that contain inverse and composition are known as *transformation groups*. So, we arrive at the following definition.

Definition 3 Let A be a set. By a transformation group, we mean a set G of functions $g : A \rightarrow A$ for which:

- If the function g is in the set G , then the inverse function g^{-1} exists and is also in the set G ;
- If the functions f and g are in the set G , their composition $g(f(a))$ is also in the set G .

We say that the optimality criterion $\geq$ is G -invariant if for all $g \in G$ and for all $a, b \in G$, we have $a \geq b$ if and only if $g(a) \geq g(b)$.

Terminological comment. In physics—where invariance is one of the main tools—transformations that keep some things invariant are called *symmetries*; see, e.g., [3, 9]. In line with this, we will also call such transformations symmetries.

The result that we will use. Now, we are ready to formulate the mathematical result that we will use to explain the behavior of the fish. To formulate this result, we need to introduce one more definition.

Definition 4 *Let A be a set, and let G be a transformation group on A . We say that an element $a \in A$ is G -invariant if for all $g \in G$, we have $g(a) = a$.*

Proposition 1 *(see, e.g., [6]) Let $\geq$ be a final G -invariant optimality criterion $\geq$, then its optimal alternative a_{opt} is G -invariant.*

Proof To prove this result, we need to prove that, for each $g \in G$, we have $g(a_{\text{opt}}) = a_{\text{opt}}$. Indeed, by definition of an optimal alternative, a_{opt} is better than or of the same quality as any other alternative. In particular, for each a , we have $a_{\text{opt}} \geq g^{-1}(a)$. Since the optimality criterion $\geq$ is G -invariant, we can conclude that $g(a_{\text{opt}}) \geq g(g^{-1}(a))$. By the definition of an inverse function, we always have $g(g^{-1}(a)) = a$, so we conclude that $g(a_{\text{opt}}) \geq a$ for all $a \in A$.

By the definition of an optimal alternative, this means that the alternative $g(a_{\text{opt}})$ is optimal. But our optimality criterion is final, which means that there is only one optimal alternative. Thus, the two optimal alternatives a_{opt} and $g(a_{\text{opt}})$ cannot be different, they must be equal: $g(a_{\text{opt}}) = a_{\text{opt}}$.

The statement is proven.

3 Case of Two Fish

Location and its symmetries. Whatever two locations the two fish select to place themselves in, these two points form a line. One can see that this 2-point spatial configuration is not invariant with respect to any shifts, but it is invariant with respect to several rotations. We can list all the rotations that keep this spatial configuration invariant:

- All rotations around the fish-connecting line; these rotations keep both locations intact, and
- All 180° rotations around a different line—a line which passes through the midpoint between the locations and which is orthogonal to the fish-connecting line; these rotations swap the two locations.

How to describe possible motions. At first glance, it seems that to describe the direction of motion of this spatial configuration, we need to describe the unit vector e in the direction of this motion. This would have been true if we considered a motion with a target destination—e.g., when the fish are pursued by a predator and try to reach a safe zone, which the predator cannot penetrate.

However, in the experiments that we are trying to explain, we are not talking about a clearly time-directed motion, we are talking about moving in circles. In this case, it should not matter whether we consider motions forward in time or the same motions viewed backward in time. Backward in time simply means that we reverse the direction of all velocities, i.e., consider the vector $-e$ instead of the vector e . From this viewpoint, what we want to describe is not so much a unit vector, but rather the direction, the pair $(e, -e)$ consisting of two opposite unit vectors.

Which motion is optimal. It is reasonable to assume that the relative quality of different motions should not change under possible rotations. Thus, by the main result from Sect. 2, the optimal motion—i.e., to be more precise, the optimal dynamic configuration consisting of the fish locations and of the motion-related pair $(e, -e)$ —should be invariant with respect to all the rotations with respect to which the initial spatial configuration is invariant.

One can easily see that if the vector e is not parallel to the fish-connecting line, then the corresponding dynamic configuration is no longer invariant with respect to all the rotations around this line—since each such rotation will change the direction of the vector e . Thus, the only invariant dynamic configuration is the one in which the vector e is parallel to the fish-connecting line, i.e., when fish follow each other.

Since, as we have proved, the optimal motion should lead to an invariant dynamic configuration, this means that the optimal motion is exactly the motion in which the fish follow each other—which is exactly what was observed.

4 Case of Three Fish

How to describe possible locations. To describe the locations of three fish, we need to select three spatial points (x_1, x_2, x_3) . As we have mentioned, rotations and shifts should not change the relative quality of different spatial locations. So, instead of a single triple of points, it makes sense to consider, as alternatives, the sets s of all possible locations that are obtained from a triple (x_1, x_2, x_3) by all possible shifts and rotations.

Which locations are optimal? Which of these sets s is optimal? Since we are talking about similar fish, it should not matter which of them we consider fish number 1 and which fish number 2—the relative quality of different spatial configurations should not change. In other words, the optimality criterion should be invariant with respect to all possible permutations of fish.

According to our main result, this means that the optimal location-describing alternative s should also be invariant under all permutations. This means, for example,

that if we rename fish 1 and 3, then the resulting triple (x_3, x_2, x_1) should belong to the same optimal set s , i.e., it can be obtained from the original triple (x_1, x_2, x_3) by shifts and rotations. Shifts and rotations do not change distance between points, so we conclude that the distance $d(x_3, x_2)$ between the points x_3 and x_2 should be equal to the distance $d(x_1, x_2)$ between the locations of similar fish in the original triple. In other words, two sides of the triangle formed by the three fish should be equal.

By considering a different permutation, we can conclude that the third side should also be equal to the other two sides—so the optimal spatial configuration should indeed be an equilateral triangle, exactly as observed.

What are the symmetries of this spatial configuration. One can see that the only rotations preserving this spatial configuration are rotations by 120 and 240° around an axis α which is orthogonal to the plane formed by the fish and which passes through the center of the fish triangle.

What are the optimal motions? As we have mentioned, in general, a motion can be characterized by a pair $(e, -e)$ of opposite unit vectors. According to our main result, the optimal dynamic configuration—consisting of the fish locations and of the motion-describing pair $(e, -e)$ —must be invariant with respect to all the corresponding symmetries—i.e., in our case, with respect to 120- and 240-degree rotations around α .

One can easily check that if the vector e is not parallel to α , then the corresponding dynamic configuration is not invariant with respect to such rotations: its orthogonal-to- α component changes when we rotate. Thus, the only invariant direction of motion is in the direction of α , i.e., in the direction orthogonal to the fish plane.

And since we proved that the optimal direction should be invariant, we thus conclude that in the three fish case, the motion corresponding to the optimal dynamic configuration is the motion in the direction orthogonal to the fish triangle—which is exactly what was observed.

Summarizing. In both cases, symmetry-based approach to optimization shows that the optimal spatial configuration and the optimal direction of motion are exactly what was observed in the recent experiments. Thus, both observed tendencies can be explained by the fact that fish act optimally.

This conclusion does not depend on what exactly optimality criterion the fish use, as long as it is invariant—as it should be—with respect to all possible rotations, shifts, and permutations of fish.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

References

1. M. Banks, When it comes to fish dynamics, three's a school. <https://physicsworld.com/a/when-it-comes-to-fish-dynamics-threes-a-school>
2. R. Belohlavek, J.W. Dauben, G.J. Klir, *Fuzzy Logic and Mathematics: A Historical Perspective* (Oxford University Press, New York, 2017)
3. R. Feynman, R. Leighton, M. Sands, *The Feynman Lectures on Physics* (Addison Wesley, Boston, Massachusetts, 2005)
4. G. Klir, B. Yuan, *Fuzzy Sets and Fuzzy Logic* (Prentice Hall, Upper Saddle River, New Jersey, 1995)
5. J.M. Mendel, *Explainable Uncertain Rule-Based Fuzzy Systems* (Springer, Cham, Switzerland, 2024)
6. H.T. Nguyen, V. Kreinovich, *Applications of Continuous Mathematics to Computer Science* (Kluwer, Dordrecht, 1997)
7. H.T. Nguyen, C.L. Walker, E.A. Walker, *A First Course in Fuzzy Logic* (Chapman and Hall/CRC, Boca Raton, Florida, 2019)
8. V. Novák, I. Perfilieva, J. Močkoř, *Mathematical Principles of Fuzzy Logic* (Kluwer, Boston, Dordrecht, 1999)
9. K.S. Thorne, R.D. Blandford, *Modern Classical Physics: Optics, Fluids, Plasmas, Elasticity, Relativity, and Statistical Physics* (Princeton University Press, Princeton, New Jersey, 2021)
10. L.A. Zadeh, Fuzzy sets. *Inf. Control.* **8**, 338–353 (1965)
11. A. Zampetaki, Y. Yang, H. Löwen., C. P. Royall, Dynamic order and many-body correlations in zebrafish show that three is a crowd. *Nat. Commun.* **15**, 2591 (2024)

Fuzzy Simulation and Data Mining

Creating a Database for Assessment of the Reliability of the Operational State of Technical Facilities



Telman Aliev  and Naila Musaeva 

Abstract It is known that the degree of reliability of the operational state of technical facility can be different when external factors change. As a result of the influence of external factors the noise uncorrelated with the useful signal appears, and a noisy signal is formed. It is the probabilistic characteristics of the useful signal that allow us to assess the degree of reliability of the operational state. In practice, the samples of the useful signal cannot be extracted from the noisy signal. Therefore, we propose algorithms for calculating the probability of a useful signal falling within the limits that allow the facility to perform the necessary functions. To assess the degree of reliability of the operational state of a facility, technologies for creating a database of probabilistic characteristics of the useful signal are developed. It is noted that depending on the value of the probability of the useful signal falling within the established intervals, it is possible to determine the degree of reliability of the operational state. As a result of use of these algorithms and technologies in information systems it is possible to establish how long the facility can work without failure; when the limit value will come; possibility of the facility to restore operational state after repair; how long the values of parameters can be in the intervals admissible for normal operation.

Keywords Noisy signal · Probability of falling within the admissible interval · Database · Information system

T. Aliev · N. Musaeva
Institute of Control Systems, Baku, Azerbaijan

T. Aliev · N. Musaeva (✉)
Azerbaijan University of Architecture and Construction, Baku, Azerbaijan
e-mail: musanaila@gmail.com

1 Introduction

It is known that technical facilities are characterized by operational and non-operational states [1, 2]. In the operational state, the technical facility meets all the requirements established by the regulations. However, the reliability of the operational state can change as a result of changes in external factors [3, 4]. This means that the values of all parameters remain in the intervals that allow the facility to carry out the assigned functions, but they can be closer to or further from the boundary values, beyond which the facility goes into a non-operational state. As a rule, this happens when the facility is affected by external factors [5, 6]. Being in the operational state, the facility can change the degree of reliability. Traditionally, the reliability of the operational state and the possibility of transition to a non-operational state are monitored by information systems [5–7]. These systems, as a rule, create databases and accumulate knowledge about the states of technical facilities [5, 6]. The database of operational and non-operational states consists of the characteristics of the noisy signal, which comes from the corresponding sensors. In this case, the noisy signal is composed of a useful signal and noise. To create a knowledge base, the algorithms and software of information systems are also supplemented with expert knowledge [7].

However, it is noted in [8] that the operational and non-operational state of the technical condition of the facility should be determined not using the characteristics of the noisy signal, but using the separate characteristics of the noise and the useful signal [9]. Knowledge of the characteristics of the noise allows to reveal the early latent period of the initiation of defects, malfunctions, etc. Knowledge of the characteristics of the useful signal makes it possible to assess the degree of reliability of the operational state of the facility when the influence of external factors changes. For this purpose, it is necessary to estimate the probability of possible values of the useful signal falling within the interval beyond which the degree of reliability of the operational state begins to decrease and there is a danger of transition to a non-operational state.

2 Problem Statement

In information systems for monitoring the state of technical facilities the necessary sensors are installed [9]. Noisy signals $G(t)$ from sensors are received, which are composed of a useful signal $X(t)$ and noise $E(t)$. In [9, 10] it is shown that noise can be caused both by the influence of external factors and by defects, malfunctions and others. The noise $E(t)$ caused by the influence of external factors is not correlated with the useful signal. For example, an increase in pipeline pressure, a change in line voltage, a change in raw material flow rate, etc. The noise $E(t)$ that resulted from defects is already correlated with $X(t)$. Noise caused by changes in external factors at the initial stage does not cause malfunction formation, but constant and

repeated influence of this noise can lead to a decrease in the degree of reliability of the operational state and eventually lead to the formation of a defect [11]. Thus, $E(t) = E_1(t)$ under the influence of external factors, $E(t) = E_2(t)$ in the formation of malfunctions, and correlation coefficients $r_{XE_1} = 0$, $r_{XE_2} \neq 0$.

To determine the degree of reliability of the operational state of technical facilities, it is necessary to determine the probability $P(\alpha \leq X(t) \leq \beta)$, with which the useful signal $X(t)$ falls within the interval $[\alpha, \beta]$ that allows the facility to carry out the assigned functions. If the probability is high, then the degree of reliability of the operational state of the technical facility is also high. As the probability decreases, the degree of reliability of the operational state also decreases. Thus, the problem comes down to determining the probability $P(\alpha \leq X(t) \leq \beta)$.

3 Algorithms for Calculating the Probability of the Useful Signal Falling Within the Intervals of Values of Operational States

Suppose the signal $G(t) = X(t) + E(t)$ comes the sensor. Signals $G(t)$, $X(t)$, $E(t)$ are stationary ergodic with normal distribution law. Mathematical expectation m_G , variance D_G , standard deviation σ_G , normal distribution density function $N(g)$, probability $P(\alpha \leq G(t) \leq \beta)$ of falling within the given interval $[\alpha, \beta]$ for the signal $G(t)$ can be calculated by discrete samples $G(i\Delta t)$ [10]:

$$D_G = \frac{\sum_{i=1}^N (G(i\Delta t) - m_G)^2}{N} = \frac{\sum_{i=1}^N (\overset{\circ}{G}(i\Delta t))^2}{N} \quad (1)$$

$$\sigma_G = \sqrt{D_G}, \quad (2)$$

$$m_G = \sum_{i=1}^N G(i\Delta t), \quad (3)$$

$$N(g) = \frac{1}{\sigma_G \sqrt{2\pi}} e^{-\frac{(g-m_G)^2}{2(\sigma_G)^2}}, \quad (4)$$

$$P(\alpha \leq G(t) \leq \beta) = \int_{\alpha}^{\beta} N(g) dg \quad (5)$$

However, it is not possible to calculate the characteristics (1–5), i.e., mathematical expectation m_X , variance D_X , standard deviation σ_X , normal distribution density function $N(x)$, probability $P(\alpha \leq X(t) \leq \beta)$ of falling within the given interval

$[\alpha, \beta]$, for the signal $X(t)$ because the samples $X(i\Delta t)$ cannot be isolated from the samples $G(i\Delta t)$ [11]:

$$D_X = \frac{\sum_{i=1}^N (X(i\Delta t) - m_X)^2}{N}, \quad (6)$$

$$\sigma_X = \sqrt{D_X}, \quad (7)$$

$$m_X = \sum_{i=1}^N X(i\Delta t), \quad (8)$$

$$N(x) = \frac{1}{\sigma_X \sqrt{2\pi}} e^{-\frac{(x-m_X)^2}{2(\sigma_X)^2}}, \quad (9)$$

$$P(\alpha \leq X(t) \leq \beta) = \int_{\alpha}^{\beta} N(x) dx. \quad (10)$$

At the same time, it is by the probability (10) of the useful signal falling within the interval $[\alpha, \beta]$, which allows to carry out the functions assigned to the technical facility, it is possible to assess the degree of reliability of the operational state, and to determine the critical value beyond which the facility goes into a non-operational state. Therefore, there is a need to calculate characteristics (6–9) for $X(t)$.

It is shown in [10, 11] that the characteristics, i.e., variance D_E^* , standard deviation σ_E^* , expectation $m_E = 0$, normal distribution density function $N^*(\varepsilon)$, probability $P(\alpha \leq E(t) \leq \beta)$ of falling within the given interval $[\alpha, \beta]$ for the noise $E(t)$ not correlated with $X(t)$, can be calculated using the formulas:

$$D_E^* = \frac{1}{N} \sum_{i=1}^N \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}(i\Delta t) - 2 \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}((i+1)\Delta t) + \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}((i+2)\Delta t) \quad (11)$$

$$\sigma_E^* = \sqrt{\frac{1}{N} \sum_{i=1}^N \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}(i\Delta t) - 2 \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}((i+1)\Delta t) + \overset{\circ}{G}(i\Delta t) \overset{\circ}{G}((i+2)\Delta t)} \quad (12)$$

$$N^*(\varepsilon) = \frac{1}{\sigma_E^* \sqrt{2\pi}} e^{-\frac{\varepsilon^2}{2(\sigma_E^*)^2}}, \quad (13)$$

$$P^*(\alpha \leq \varepsilon(t) \leq \beta) = \int_{\alpha}^{\beta} N^*(\varepsilon) d\varepsilon. \quad (14)$$

Then, given that $r_{XE} = 0$, we compute expressions for the characteristics of $X(t)$. Taking into account (1, 2, 6, 7, 11 and 12), as well as.

$$D_G = D_X + D_E,$$

we obtain expressions for the variance of $X(t)$ and standard deviation:

$$D_X^* = D_G - D_E^*, \quad (15)$$

$$\sigma_X^* = \sqrt{D_X^*}. \quad (16)$$

Since $m_E = 0$, given the expression $m_G = m_X + m_E$, we have.

$$m_X^* = m_G. \quad (17)$$

Thus, knowing the values of σ_X^* and m_X^* from expressions (16, 17), we can determine the distribution density function of the useful signal [12]:

$$N^*(x) = \frac{1}{\sigma_X^* \sqrt{2\pi}} e^{-\frac{(x-m_X^*)^2}{2(\sigma_X^*)^2}}. \quad (18)$$

Then the probability with which $X(t)$ falls within the interval that allows the facility to carry out the assigned functions will be determined by the expression:

$$P(\alpha \leq X(t) \leq \beta) = \int_{\alpha}^{\beta} N^*(x) dx. \quad (19)$$

4 Technologies for Creating a Database for Assessment of the Reliability of the Operational State of Facilities

Suppose that at the initial instant t_0 the facility is in the operational condition, i.e. it meets all the requirements necessary for normal operation. In this case, the values of all parameters $X_1(t), X_2(t), \dots, X_n(t)$ characterizing the reliability of the technical state are within the specified limits:

$$X_1 \min \leq X_1 \leq X_1 \max, \dots, X_n \min \leq X_n \leq X_n \max. \quad (20)$$

For $X_1(t), X_2(t), \dots, X_n(t)$ we calculate:

- (1) standard deviations $\sigma_{X_{1,t0}}^*, \sigma_{X_{2,t0}}^*, \dots, \sigma_{X_{n,t0}}^*$ by formula (16);
- (2) mathematical expectations $m_{X_{1,t0}}^*, m_{X_{2,t0}}^*, \dots, m_{X_{n,t0}}^*$ by formula (17);
- (3) distribution density functions $N^*(x_1)_{t0}, N^*(x_2)_{t0}, \dots, N^*(x_n)_{t0}$ by (18);
- (4) probabilities with which the parameters fall within the intervals (19) allowing to carry out the functions assigned to the facility:

$$\begin{aligned} P_{1,t0}(X_1 \min \leq X_1 \leq X_1 \max) &= \int_{X_1 \min}^{X_1 \max} N^*(x_1)_{t0} dx_1, \dots, \\ P_{n,t0}(X_n \min \leq X_n \leq X_n \max) &= \int_{X_n \min}^{X_n \max} N^*(x_n)_{t0} dx_n. \end{aligned} \quad (21)$$

Probabilities (21) are reference probabilities. This means that the probabilities calculated at subsequent time instants $t_1, t_2, t_3, \dots$, must be compared with these values. This comparison allows to assess how much the technical state of the facility has improved or deteriorated compared to the time instant t_0 .

At given time intervals, the calculations according to formulas (16–21) should be repeated. If nothing is known about the apparent change of external factors, the intervals $\Delta t = t_1 - t_0 = t_2 - t_1 = t_3 - t_2 = \dots$, can be the same. But, if the change of external factors is obvious, it is necessary to perform extra calculations.

Thus, at time instant t_1 the following characteristics are re-calculated standard deviations $\sigma_{X_{1,t1}}^*, \dots, \sigma_{X_{n,t1}}^*$; mathematical expectations $m_{X_{1,t1}}^*, \dots, m_{X_{n,t1}}^*$; distribution density functions $N^*(x_1)_{t1}, \dots, N^*(x_n)_{t1}$; probabilities.

$$\begin{aligned} P_{1,t1}(X_1 \min \leq X_1 \leq X_1 \max) &= \int_{X_1 \min}^{X_1 \max} N^*(x_1)_{t1} dx_1, \dots, \\ P_{n,t1}(X_n \min \leq X_n \leq X_n \max) &= \int_{X_n \min}^{X_n \max} N^*(x_n)_{t1} dx_n. \end{aligned} \quad (22)$$

This process continues for each t_i -th time instant. Each time after the probabilities are calculated, the reliability degree of the technical facility is assessed as follows:

1. If for all parameters the probability values are

$$\begin{aligned} P_{1,ti}(X_1 \min \leq X_1 \leq X_1 \max) &\approx 1, \dots, \\ P_{n,ti}(X_n \min \leq X_n \leq X_n \max) &\approx 1, \end{aligned} \quad (23)$$

it means that the reliability of the technical state of the facility is high, and the values of all parameters are within the limits that allow to perform the necessary functions.

2. If at least for one of the parameters $X_k(t)$ it is true that

$$0.5 \leq P_{k,ti}(X_k \min \leq X_k \leq X_k \max) < 1, \quad (24)$$

it means that the degree of reliability of the operational state has changed, and the facility is no longer able to perform the assigned functions properly. At the same time, however, the fail-safe operation of the facility is preserved, i.e. the facility can maintain its operational state for some time.

3. If at least for one of the parameters $X_k(t)$ it is true that

$$0.2 \leq P_{k,ti}(X_k \min \leq X_k \leq X_k \max) \leq 0.5, \quad (25)$$

it means that the degree of reliability of the operational state is low, and the facility will not be able to maintain durability, i.e. operational state before the limit state occurs.

4. If at least for one of the parameters $X_k(t)$ it is true that

$$0.01 \leq P_{k,ti}(X_k \min \leq X_k \leq X_k \max) < 0.2, \quad (26)$$

it means that the degree of reliability of operational state is quite low, but the facility is in repairable condition, i.e. it can return to operational state after repair.

5. If conditions (23) are violated for two or more parameters, then the facility loses the property of preservability, i.e. it cannot keep the values of the parameters characterizing the ability to perform the required functions within the specified limits. This eventually leads to the facility going into a non-operational state.

5 Conclusion

The developed algorithms for calculating the probabilistic characteristics of the useful component of the noisy signal can be widely used in information systems. These algorithms make it possible to determine the probabilities of the useful component of the noisy signal at the change of external influences falling within the intervals that allow to carry out the functions assigned to the facility. Creating a database of probabilistic characteristics of the useful signal makes it possible to determine the early period of change in the degree of reliability of the operational state of the technical facility and prevent its transition to a non-operational state.

References

1. Alessandro Birolini, *Reliability Engineering: Theory and Practice* (Springer, 2018), p. 651. https://www.thriftbooks.com/w/reliability-engineering-theory-and-practice_alessandro-birolini/2334384/all-editions/
2. Alessandro Birolini, *Quality and Reliability of Technical Systems: Theory, Practice, Management* (Springer, 2012), p. 502. https://www.thriftbooks.com/w/quality-and-reliability-of-technical-systems-theory-practice-management_alessandro-birolini/10726566/#edition=10004573&idq=14814653
3. J. Žurek, J. Małachowski, J. Ziółkowski, J. Szkutnik-Rogoż, Reliability analysis of technical means of transport. *Appl. Sci.* **10**(9), 3016 (2020). <https://doi.org/10.3390/app10093016>
4. T. Nowakowski, M. Młyńczak, A. Jodejko-Pietruczuk, S. Werbińska-Wojciechowska, *Safety and Reliability: Methodology and Applications*, 1st Edn (CRC Press, 2014), p. 408. <https://www.amazon.com/Safety-Reliability-Applications-Tomasz-Nowakowski/dp/1138026816>
5. S. Duer, Examination of the reliability of a technical object after its regeneration in a maintenance system with an artificial neural network. *Neural Comput. Applic.* **21**, 523–534 (2012). <https://doi.org/10.1007/s00521-011-0723-2>
6. A. Breznická, M. Kohutiar, M. Krbaťa, M. Eckert, P. Mikuš, Reliability analysis during the life cycle of a technical system and the monitoring of reliability properties. *Systems* **11**(12), 556 (2023). <https://doi.org/10.3390/systems11120556>
7. Z. Kasprzyk, Methodology of maintainability and reliability prediction of technical objects. *Transp. System. Telemat.* **3**(2), 18–23 (2010). <https://bibliotekanauki.pl/articles/393275>
8. T.A. Aliev, *Noise Control of the Beginning and Development Dynamics of Accidents* (Springer, 2019), p. 201. <https://doi.org/10.1007/978-3-030-12512-7>
9. T.A. Aliev, N.F. Musaeva, M.T. Suleymanova, B.I. Gazizade, Analytical representation of the density function of normal distribution of noise. *J. Autom. Inf. Sci.* **47**(8), no 4, 24–40 (2015). <https://doi.org/10.1615/JAutomatInfScien.v47.i8.30>
10. T.A. Aliev, N.F. Musaeva, M.T. Suleymanova, Algorithms for indicating the beginning of accidents based on the estimate of the density distribution function of the noise of technological parameters. *Autom. Control. Comput. Sci.* **52**(3), 231–242 (2018). <https://doi.org/10.3103/S0146411618030021>
11. T.A. Aliev, N.F. Musaeva, Technologies for early monitoring of technical objects using the estimates of noise distribution density. *J. Autom. Inf. Sci.* **51**(9), 12–23 (2019). <https://doi.org/10.1615/JAutomatInfScien.v51.i9.20>
12. T.A. Aliev, N.F. Musaeva, B.I. Gazizade, Algorithms for calculating high-order moments of the noise of noisy signals. *J. Autom. Inf. Sci.* **50**(6), 1–13 (2018)

GANs Applications in Chemical Engineering



Köksal Erentürk and Saliha Erentürk

Abstract In the field of machine learning, Generative Adversarial Networks (GANs) have arisen as an influential tool demonstrating exceptional capabilities in generating synthetic data that closely comply with real-world samples. The last decade has seen remarkable advances in speech, image and language recognition tools that have been made available to the public through computer and mobile devices' applications in broad area of sciences, including Chemical Engineering. In this paper, existing and future potential applications of GANs in Chemical Engineering have been given in detail. GANs have a strong potential to revolutionize in chemical engineering and its multiple subfields such as process optimization, molecular design, and materials science. Based on the searches on recent literature and case studies, generating realistic molecular structures, fault diagnosis and predicting chemical properties, optimizing process parameters, and designing novel materials are some of the GANs applications in chemical engineering. Challenges and opportunities associated with the integration of GANs into traditional chemical engineering workflows have also been discussed. Additionally, future potential applications of GANs have been examined with examples, in this paper.

Keywords Generative adversarial networks · Artificial intelligence · Chemical engineering

1 Introduction

Generative Adversarial Networks, or GANs for short, are an approach to generative modeling using deep learning methods, such as convolutional neural networks. Generative modeling is an unsupervised learning task in machine learning that involves automatically discovering and learning the regularities or patterns in input data in such a way that the model can be used to generate or output new examples that plausibly could have been drawn from the original dataset. GANs are a class of Deep

K. Erentürk (✉) · S. Erentürk
Atatürk University, Erzurum, Türkiye
e-mail: keren@atauni.edu.tr

Neural Networks that have the ability to generate realistic synthetic data including but not limited to images, text, or even videos. As a result of these revolutionary developments, Generative AI is transforming the field of chemical engineering, enabling researchers and engineers to achieve higher levels of efficiency, reduce costs, and drive innovation in the chemical industry. Generative Adversarial Networks (GANs) is one of the most popular approaches used in Generative AI, providing researchers with a powerful tool for designing novel chemical compounds, optimizing process parameters, and designing innovative materials.

A GAN uses a decoder (or generator) and discriminator to learn the materials data distribution implicitly. In GAN approaches, a key component of crystal structural generative models is the invertibility from material representation (features) to real structure of material since the features generated from the latent vector should eventually be inverted back to the real structure of material in order to confirm the generated material [1, 2]. A generative adversarial network (GAN) employed for crystal structure generation using a coordinate-based crystal representation inspired by point clouds [3]. By conditioning the network with the crystal composition, the proposed model generated materials with a desired chemical composition. As an application in this model, GAN architecture applied to generate new Mg – Mn – O ternary compounds to find potential photoanode materials and discovered 23 new crystal compounds with reasonable stability in an aqueous environment and band gap.

Stacked denoising autoencoders (SDAE) combined with GAN for planetary gearbox fault pattern recognition with limited fault samples and strong noise interference [4]. An auxiliary classifier GAN (ACGAN) model developed to learn from mechanical sensor signals and generate realistic one-dimensional raw data and used a CNN classifier to output the machine fault diagnosis result [5]. A supervised classifier framework with GAN introduced to increase the number of faulty training samples and re-balance the training dataset and performed robust fault diagnosis for air handling units [6]. An enhanced ACGAN built with deep neural network (DNN) to solve imbalanced problems and classify the chemical process faults [7]. A high-efficiency GAN model (HGAN) proposed for chemical process fault diagnosis [8]. HGAN integrates the advantages of Wasserstein GAN and Auxiliary Classifier GAN to promote the generating model training stability and the discriminative model training efficiency with Bayesian optimization.

A vast number of organic molecules are applied in solar cells, such as organic light emitting diodes, conductors, and sensors [9]. Synthesis of new organic and inorganic compounds is a challenge in physics, chemistry and in materials science. A number of machine learning approaches were proposed to facilitate the search for novel stable compositions [10]. There was an attempt to find new compositions using an inorganic crystal structure database, and to estimate the probabilities of new candidates based on compositional similarities. A learning method called CrystalGAN proposed to discover cross-domain relations in real data, and to generate novel structures [11].

Near-infrared (NIR) spectroscopy has been widely used to predict the chemical properties of materials especially gasoline properties that are difficult to measure online during gasoline blending. NIR models should be prepared in advance to apply

this technique successfully. Obtaining a high-accuracy NIR model in practice is hard because abundant labelled samples are difficult to acquire. The application of WGAN presented for the prediction of octane number by NIR spectroscopy. It is observed that when labelled data are insufficient during gasoline blending, the proposed method can help to establish an initial NIR model quickly [12].

2 GANs Applications in Chemical Engineering

Detailed representations of three different GANs applications have been given in detail.

2.1 Generative Adversarial Networks (GANs)

In 2014, a groundbreaking paper by Ian Goodfellow et al. introduced the world to Generative Adversarial Networks (GANs). As depicted in Fig. 1, this ingenious technique revolutionized the field of artificial intelligence by creating a system that pits two neural networks against each other. One network, the generator, strives to create ever-more realistic data, be it images, music, or even 3D models. The other network, the watchful discriminator, aims to distinguish the generated creations from real-world samples.

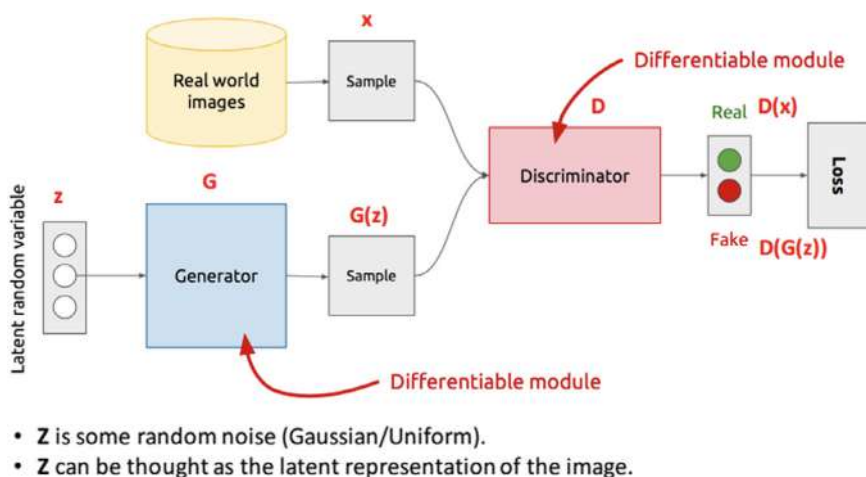


Fig. 1 GANs architecture

2.2 Generative Adversarial Networks for Crystal Structure Prediction

One of the proposed GANs model [3] for Crystal Structure Prediction consists of three network components: a generator, a critic, and a classifier as shown in Fig. 2. The generator takes the random Gaussian noise vector (Z) and one-hot encoded composition vector (C_{gen}) as the input to generate new 2D-representations. The one-hot encoded composition vector is used as a condition to generate materials with target composition. The critic computes the Wasserstein distance which represents dissimilarity between the true and trained data distributions, and by reducing this distance the generator would generate more realistic materials. The critic network is composed of three-shared multilayers perceptions (MLPs) followed by average pooling layers to ensure the permutation invariance under the reordering of points in the 2D-representation. It is noted that the permutation invariance under the reordering of input is satisfied by using shared weight parameters and average pooling since the averaged value is unchanged under the change of orders. The classifier network, which outputs the composition vector from the input 2D-representation, is used to ensure that the generated new materials meet the given composition condition. The loss of the classifier is back-propagated to the generator only if the generated 2D-representation ($\tilde{x}$) is taken as input.

Z , C_{gen} , and C_{real} denote a random input noise, user-desired composition condition, and composition of real material, respectively. The variables $\tilde{x}$ and x denote the feature (representation) of generated and real materials, respectively. $\hat{C}_{gen}$ and $\hat{C}_{real}$ denote the predicted composition of the generated and real features, respectively. $D(x)$ is the critic function also known as the critic network.

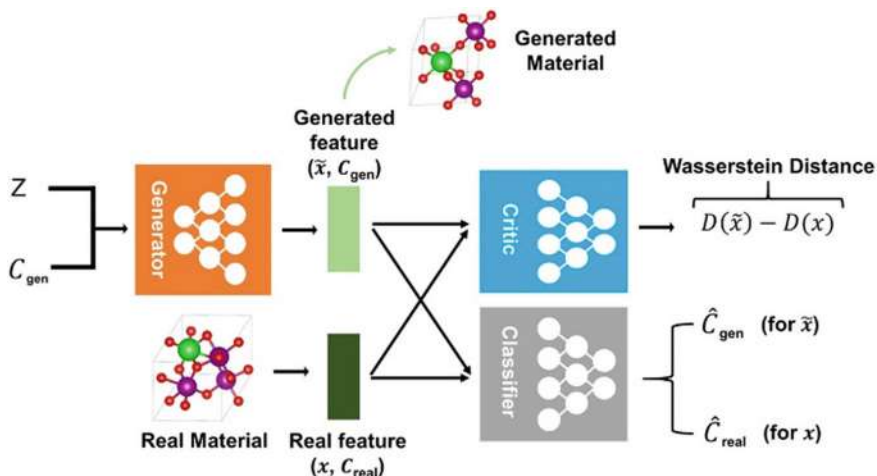


Fig. 2 GANs architecture for crystal structure prediction [3]

As an application in this model, GAN architecture applied to generate new Mg – Mn – O ternary compounds to find potential photoanode materials and discovered 23 new crystal compounds with reasonable stability in an aqueous environment and band gap.

2.3 Application of GANs in Chemical Process Fault Diagnosis

Another one of the proposed GANs model used to fault diagnosis in a chemical process [8]. As given in Fig. 3, the proposed HGAN model incorporates the advantages of WGAN and ACGAN, which utilizes Wasserstein distance and gradient penalty to empower the generator training process, and adopts the Bayesian optimization strategy to enhance the discriminator performance.

At the offline stage, raw data of the chemical process is collected including operation parameters and process parameters, which is subsequently divided into different sets for training and testing. After normalization and other necessary pre-processing steps, the training dataset is used to train the HGAN model and acquire the optimized parameters of the generator and the discriminator. The benchmark Tennessee Eastman process (TEP) simulator is applied to evaluate the performance of the proposed model. The experiments are conducted with small size of training samples under data balance and data imbalance conditions separately. The diagnosis results of HGAN are compared with a traditional statistical model.

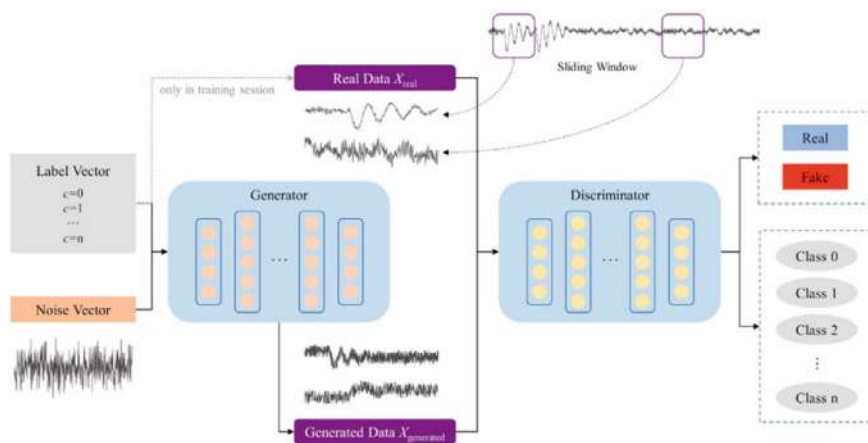


Fig. 3 ACGAN (HAGN) model structure for fault diagnosis

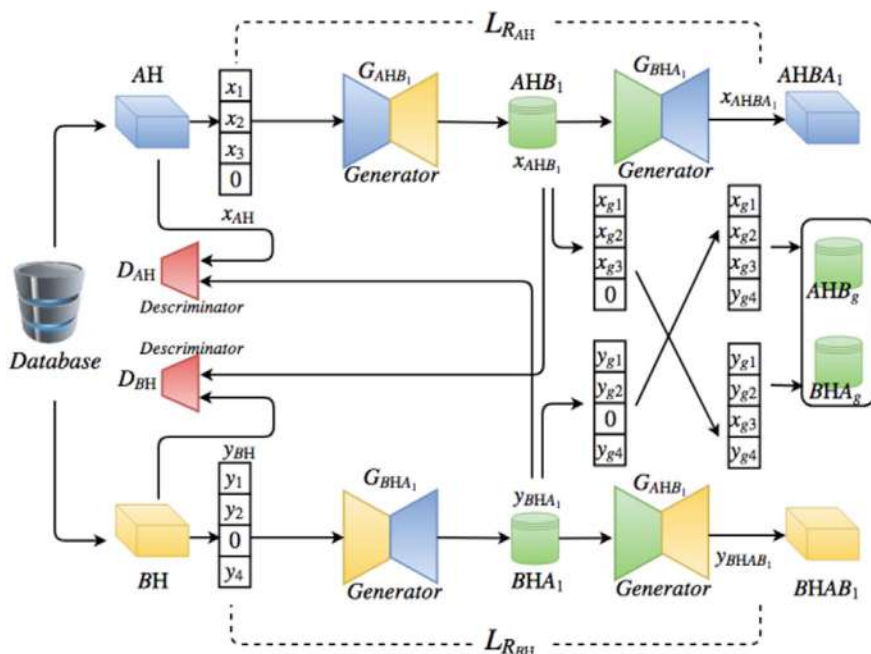


Fig. 4 CrystalGAN model architecture for discovering crystallographic structures

2.4 Learning to Discover Crystallographic Structures with GANs

As the last example, a CrystalGAN [11] which generates new chemically stable crystallographic structures with increased domain complexity will be given. The basic architecture of the CrystalGAN presented in Fig. 4.

In this work [11], authors focus on applications of hydrogen storage, and in particular, the problem to investigate novel chemical compositions with stable crystals. Based on the CrystalGAN model architecture for discovering crystallographic structures, it is noted that although the CrystalGAN was developed and tested for applications in materials science, it is a general method where the constraints can be easily adapted to any scientific problem.

3 Conclusions

Generative Adversarial Networks (GANs) have emerged as a powerful tool for Chemical Engineering, offering a glimpse into a generative future. Their ability to create realistic and intricate data holds immense potential for accelerating material

discovery, optimizing processes, and empowering data-driven approaches. While challenges like data scarcity and model interpretability remain, advancements in machine learning and computational resources promise even more sophisticated applications in the coming years. By embracing GAN technology, Chemical Engineering can unlock entirely new avenues for innovation, leading to a more sustainable and efficient future for the industry.

References

1. E. Putin, A. Asadulaev, Y. Ivanenkov, V. Aladinskiy, B. Sanchez-Lengeling, A. Aspuru-Guzik, A. Zhavoronkov, Reinforced adversarial neural computer for de novo molecular design. *J. Chem. Inf. Model.* **58**(6), 1194 (2018)
2. A. Zhavoronkov, Y.A. Ivanenkov, A. Aliper, M.S. Veselov, V.A. Aladinskiy, A.V. Aladinskaya, V.A. Terentiev, D. Polykovskiy, M.D. Kuznetsov, A. Asadulaev, Deep learning enables rapid identification of potent DDR1 kinase inhibitors. *Nat. Biotechnol.* **37**(9), 1038 (2019)
3. S. Kim, J. Noh, G.H. Gu, A. Aspuru-Guzik, Y. Jung, Generative adversarial networks for crystal structure prediction. *ACS Cent. Sci.* **6**(8), 1412–1420 (2020)
4. Z. Wang, J. Wang, Y. Wang, An intelligent diagnosis scheme based on generative adversarial learning deep neural networks and its application to planetary gearbox fault pattern recognition. *Neurocomputing* **310**, 213–222 (2018)
5. S. Shao, P. Wang, R. Yan, Generative adversarial networks for data augmentation in machine fault diagnosis. *Comput. Ind.* **106**, 85–93 (2019)
6. K. Yan, Unsupervised learning for fault detection and diagnosis of air handling units. *Energy Build.* **210**, 109689 (2020)
7. P. Peng, Y. Wang, W. Zhang, Y. Zhang, H. Zhang, Imbalanced process fault diagnosis using enhanced auxiliary classifier GAN,” in *Chinese Automation Congress* (2020)
8. R. Qin, J. Zhao, High-efficiency generative adversarial network model for chemical process fault diagnosis. *IFAC Pap. Line* **55**(7), 732–737 (2022)
9. X. Yang, J. Zhang, K. Yoshizoe, K. Terayama, K. Tsuda, ChemTS: an efficient Python library for de novo molecular generation. *Sci. Technol. Adv. Mater.* **18**(1), 972–976 (2017)
10. K.T. Butler, D.W. Davies, H. Cartwright, O. Isayev, A. Walsh, Machine learning for molecular and materials science. *Nature* **559**, 547–555 (2018)
11. A. Noura, N. Sokolovska, J.C. Crivello, CrystalGAN: learning to discover crystallographic structures with generative adversarial networks, in *AAAI Spring Symposium: Combining Machine Learning with Knowledge Engineering 2019* (Palo Alto, CA, 2019)
12. K. He, J. Liu, Z. Li, Application of generative adversarial network for the prediction of gasoline properties. *Chem. Eng. Trans.* **81**, 907–912 (2020)

The Method of Group Temperature Measurements of Remote Atmospheric Objects



Huseynova Rena and Aliyeva Gunel

Abstract The article is devoted to the developed methodology for conducting group temperature measurements of remote objects. The problem of measuring the total temperature of a group of non-identical objects performing a group flight is formulated and solved using the variation optimization method. Examples of such objects include flights of aircraft in a group, ground scenes involving a group of objects of interest to us, temperature diagnostics of various points of buildings, traffic control on highways, etc. The condition is determined, under which the total value of the measured temperatures of all elements in the group reaches an extreme value. The regression relationship function between the emissivity of the group elements and the atmospheric transmittance has been specifically calculated. The problem of optimal placement of group elements with different emissivity values in the atmosphere on routes with different atmospheric transmittance is practically solved. It is noted that many different values of atmospheric transmission can be obtained by measuring at different wavelengths. The proposed method can be used for remote control of flight or operation of a group of objects with different values of the emissivity coefficient with a special order of selection of the wavelengths of measurements based on the condition of constancy of the sum of the optical thicknesses of the atmosphere and the wavelengths used for measurements. It is shown that the property of the total temperature extremum is achieved when the value of the sum of the emissivity coefficient of all flying objects is constant.

Keywords Group temperature measurements · Remote objects · Optimization

H. Rena (✉)

Azerbaijan University of Architecture and Construction, Baku, Azerbaijan

e-mail: renahuseynova55@gmail.com

A. Gunel

Azerbaijan State Oil and Industry University, National Aerospace Agency, Baku, Azerbaijan

e-mail: gunel.aliyeva.va@asoiu.edu.az

1 Introduction

It is well known that any object with a temperature above zero Kelvin emits infrared radiation [1]. At the same time, infrared radiation makes it possible to study both moving and non-moving objects emitting in the thermal range [2]. Infrared temperature measurement of objects is one of the methods of remote sensing of the surrounding world and is widely used for both civilian and military purposes [3]. Infrared temperature measurements are, in principle, non-contact measurements and can be used to study objects at both close and long distances. Temperature infrared measurements at close distances, or proximal thermal measurements, are used in medicine [4–7], veterinary medicine [8–11], in the study of technological processes [12–14], in the construction of buildings [15–17], in non-destructive testing [8], etc. At the same time, the objects of interest to us often exist in the form of a certain group of thermal radiators and a general qualitative characteristic of the functioning of such groups can be given using such an indicator as the total temperature of the elements of such groups. Examples of such objects include flights of aircraft in a group, ground scenes involving a group of objects of interest to us, temperature diagnostics of various points of buildings, traffic control on highways, control of group flights of birds, etc. Obviously, such measurements must be calibrated in one way or another, and the quality of such measurements is often determined by the signal-to-noise ratio in the resulting result. Therefore, the researcher is always interested in getting the best result with low total noise. This property of remote measurements actualizes the formation of various optimization tasks related to remote temperature measurements.

2 Materials and Methods

Next, we will specifically focus on conducting group temperature measurements and for this reason we will provide some basic information about the emissions and re-emissions of distant celestial objects. Borrowed from the source [18]. When considering the case of infrared temperature measurement of one object using an infrared camera (Fig. 1), the total radiation at the camera input can be estimated using the following expression:

$$W_{tot} = E_{obj} + E_{reft} + E_{atm} \quad (1)$$

where E_{obj} is the radiation of the object of interest to us; E_{reft} is the re-reflected radiation of the environment; E_{atm} is the radiation of the atmosphere.

These components are determined by the following expressions:

$$E_{obj} = \varepsilon_{obj} \cdot \tau_{atm} \cdot \sigma (T_{obj})^4 \quad (2)$$

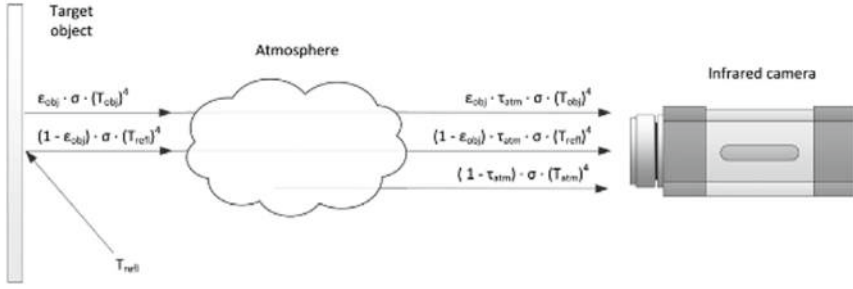


Fig. 1 The flow of infrared radiation from a target object through the atmosphere to an infrared camera

where ϵ_{obj} is the radiance of the object; σ is the Stefan-Boltzmann constant; τ_{atm} is the transmission of the atmosphere; T_{obj} is the temperature of the atmosphere.

$$E_{refl} = \rho_{obj} \cdot \tau_{atm} \cdot \sigma (T_{refl})^4 = (1 - \epsilon_{obj}) \cdot \tau_{atm} \cdot \sigma (T_{refl})^4 \quad (3)$$

where ρ_{obj} is the spectral reflection coefficient; T_{refl} is the reflected temperature

$$E_{atm} = \epsilon_{atm} \cdot \sigma (T_{atm})^4 = (1 - \tau_{atm}) \sigma (T_{atm})^4 \quad (4)$$

where ϵ_{atm} is the radiance of the atmosphere; T_{atm} is the temperature of the atmosphere.

Taking into account the expressions (2), (3), (4) we get [7]:

$$W_{tot} = \epsilon_{obj} \tau_{atm} \sigma (T_{obj})^4 + (1 - \epsilon_{obj}) \tau_{atm} \sigma (T_{refl})^4 + (1 - \tau_{atm}) \sigma (T_{atm})^4 \quad (5)$$

From the expression (5) we will get

$$T_{obj} = \sqrt[4]{\frac{W_{tot} - (1 - \epsilon_{obj}) \tau_{atm} \sigma (T_{refl})^4 - (1 - \tau_{atm}) \sigma (T_{atm})^4}{\epsilon_{obj} \tau_{atm} \sigma}} \quad (6)$$

for further research, we will rewrite formula (6) in the form:

$$T_{obj} = \sqrt[4]{\frac{W_{tot} + \epsilon_{obj} (\tau_{atm} \sigma T_{refl}^4) - (\tau_{atm} \sigma T_{refl}^4) - \sigma (T_{atm})^4 + \tau_{atm} \sigma T_{atm}^4}{\epsilon_{obj} \tau_{atm} \sigma}} \quad (7)$$

As a result of further transformations, we will rewrite formula (7) as

$$T_{obj} = \sqrt[4]{\frac{W_{tot} - \tau_{atm} (\sigma T_{refl}^4 - \sigma T_{atm}^4) - \sigma T_{atm}^4}{\epsilon_{obj} \tau_{atm} \sigma}} + T_{refl} \quad (8)$$

To simplify further writing, we introduce the following notation:

$$\begin{aligned} a_1 &= W_{tot}; a_2 = (\sigma T_{rel}^4 - \sigma T_{atm}^4); a_3 = \sigma T_{atm}^4; \\ a_4 &= \sigma; a_5 = T_{rel}; a_6 = a_1 - a_3 \end{aligned} \quad (9)$$

In this case, formula (8) takes the following form

$$T_{obj}^4 = \frac{\frac{a_6}{\tau_{atm}} - a_2}{a_4 \varepsilon_{obj}} + a_5 \quad (10)$$

Let's introduce the following communication function for consideration:

$$\varepsilon_{obj} = \psi(\tau_{atm}) \quad (11)$$

Next, we consider that it is possible to form an ordered set

$$T = \{\tau_{atm,i}\}; i = \overline{1, n} \quad (12)$$

by two methods:

By making measurements at such wavelengths λ_i ; in which the condition is fulfilled

$$\tau_{atm,i} = \tau_{atm,i-1} + \Delta \tau_{atm}; i = \overline{1, n}; \tau_{atm,0} = 0 \quad (13)$$

Having measured objects located from the camera at such distances R_i , under which conditions (12), (13) are fulfilled.

In this case, you can make up the amount

$$\sum_{i=1}^n T_{obj,i}^4 = \sum_{i=1}^n \frac{a_6}{a_4 \psi(\tau_{atm})} + a_5 \quad (14)$$

Considering the function (11) conditionally continuous, we impose the following requirement on this function

$$\int_0^{\tau_{atm,max}} \psi(\tau_{atm}) d\tau_{atm} = C; C = const \quad (15)$$

Also, considering the expression (14) conditionally continuous, we write

$$F = \int_0^{\tau_{atm,max}} T_{obj}^4 d\tau_{atm} = \int_0^{\tau_{atm,max}} \left[\frac{a_6}{a_4 \psi(\tau_{atm})} + a_5 \right] d\tau_{atm} \quad (16)$$

Taking into account (15), (16), we compose a variational optimization problem for calculating the optimal function $\psi(\tau_{atm})$, in which the functional F , which is the target functional, reaches an extremum, where

$$F = \int_0^{\tau_{atm, \max}} \left[\frac{a_6}{a_4 \psi(\tau_{atm})} + a_5 \right] d\tau_{atm} + \gamma \left[\int_0^{\tau_{atm, \max}} \psi(\tau_{atm}) d\tau_{atm} - C \right] \quad (17)$$

where γ is the Lagrange multiplier.

According to [18], the solution of problem (17) must satisfy the condition

$$\frac{d \left\{ \frac{a_6}{a_4 \psi(\tau_{atm})} + a_5 + \gamma \psi(\tau_{atm}) \right\}}{d \psi(\tau_{atm})} = 0 \quad (18)$$

where

$$a_7 = \frac{a_6}{a_4}; a_8 = \frac{a_2}{a_4} \quad (19)$$

From condition (18) we obtain

$$-\frac{\left(\frac{a_7}{\tau_{atm}} - a_8 \right)}{\psi^2(\tau_{atm})} + \gamma = 0 \quad (20)$$

From expression (20) we find

$$\psi(\tau_{atm}) = \sqrt{\frac{\left(\frac{a_7}{\tau_{atm}} - a_8 \right)}{\gamma}} \quad (21)$$

Taking into account expressions (15) and (21), we obtain

$$\int_0^{\tau_{atm, \max}} \sqrt{\frac{\left(\frac{a_7}{\tau_{atm}} - a_8 \right)}{\gamma}} d\tau_{atm} = C \quad (22)$$

The Lagrange multiplier calculated by expression (22) is denoted as γ_0 . Therefore, under the condition

$$\psi(\tau_{atm}) = \sqrt{\frac{\left(\frac{a_7}{\tau_{atm}} - a_8 \right)}{\gamma_0}} \quad (23)$$

the functional F reaches an extremum.

At the same time, if $\left(\frac{a_7}{\tau_{atm}} - a_8\right) > 0$ F reaches a minimum if $\left(\frac{a_7}{\tau_{atm}} - a_8\right) < 0$ F reaches the maximum.

3 Discussion

Thus, the optimization problem of measuring the total temperature of a group of non-identical objects performing a group flight is formulated and solved. The condition is determined, under which the total value of the measured temperatures of all elements in the group reaches an extreme value. The relationship function between the emissivity of the group elements and the atmospheric transmittance has been specifically calculated. The problem of placing group elements with different emissivity values on routes with different atmospheric transmittance is physically solved. It is noted that many different values of atmospheric transmission can be obtained by measuring at different wavelengths.

4 Conclusion

An extreme method for measuring the total temperature in the IR range of a group of unequal aircraft is proposed, the radiance of which is determined in a functional relationship from the transmission of the atmosphere measured at different wavelengths. It is shown that the property of the total temperature extremum is achieved when the value of the sum of the emissivity coefficient of all flying objects is constant.

The proposed method can be used for remote control of flight or operation of a group of objects with different values of the emissivity coefficient with a special order of selection of the wavelengths of measurements based on the condition of constancy of the sum of the optical thicknesses of the atmosphere and the wavelengths used for measurements.

References

1. M.F.Modest, in *Radiative Heat Transfer* (Academic press, Waltham, MA, USA, 2013)
2. R. Gade, T.B. Moeslund, Thermal cameras and applications: a survey. *Mach. Vision Appl.* **25**, 245–262 (2014)
3. C. Meola, *Infrared thermography: recent advances and future trends* (Bentham Science, New York, NY, USA, 2012)
4. Lahiri B., Bagavathiappan S., Jayakumar T., Philip J. Medical applications of infrared thermography: a review// *Infrared Phys. Technol.* **55**. 221–235, (2012).
5. A.G. Becerra, J.E. Olguín-Tiznado, J.L. García Alcaraz, C. Camargo Wilson, B.R. García-Rivera, R. Vardasca et al., Infrared thermal imaging monitoring on hands when performing

- repetitive tasks: an experimental study. PLoS ONE **16**, e0250733 (2021). <https://doi.org/10.1371/journal.pone.0250733>
6. I. Cruz-Vega, D. Hernandez-Contreras, H. Peregrina-Barreto, J.D.J. RangelMagdaleno, J. Ramirez-Cortes, MDeep learning classification for diabetic foot thermograms. *Sensors* **20**, 1762 (2020). <https://doi.org/10.3390/s20061762>
 7. A. Kirimat, O. Krejcar, A. Selamat, E. Herrera-Viedma, FLIR vs SEEK thermal cameras in biomedicine: comparative diagnosis through infrared thermography. *BMC Bioinformatics* **21**, 1–10 (2020). <https://doi.org/10.1186/s12859-020-3355-7>
 8. Purohit R., Turner T. A., Pascoe D. D. Use of infrared imaging in veterinary medicine// Medical infrared imaging. Diakides N.A., Bronzino J.D. CRC Press: Boca Raton, FL. USA, (2008).
 9. Dórea FC, Vial F, Hammar K, Lindberg A, Lambrix P, Blomqvist E, et al. Drivers for the development of an Animal Health Surveillance Ontology (AHSO). *Prev Vet Med.* 166:39–48. <https://doi.org/10.1016/j.prevetmed.2019.03.002>, (2019)
 10. Anagnostopoulos A, Barden M, Tulloch J, Williams K, Griffiths B, Bedford C, et al. A study on the use of thermal imaging as a diagnostic tool for the detection of digital dermatitis in dairy cattle. *J Dairy Sci.* 104:10194– 202. <https://doi.org/10.3168/jds.2021-20178>, (2021)
 11. Bleul U, Hässig M, Kluser F. Screening of febrile cows using a small handheld infrared thermography device. *Tierarztl Prax Ausgabe G Grosstiere Nutztiere.* 49:12–20. <https://doi.org/10.1055/a-1307-9993>, (2021)
 12. R. Bogue, Sensors for condition monitoring: a review of technologies and applications. *Sens. Rev.* **33**, 295–299 (2013)
 13. G.T. Alliot, R.L. Higginson, G.D. Wilcox, Producing a thin coloured film on stainless steels – a review. Part 2: nonelectrochemical and laser processes. *Trans. IMF* **101**(2), 72–78, (2023)
 14. H. Lin, W. Lin, M. Hong, Surface coloring by laser irradiation of solid substrates. *APL Photonics* **4**(5), 51101 (2019)
 15. A.R. Al-Kassir, J. Fernandez, F. Tinaut, F. Castro, Thermographic study of energetic installations. *Appl. Therm. Eng.* **25**, 183–190 (2005)
 16. A. Sivasuriyan, D.S. Vijayan, W. Górski, Ł. Wodzy ´n, M.D. Vaverková, E. Koda, Practical implementation of structural health monitoring in multi-story buildings. *Buildings* **11**, 263 (2021)
 17. J.F. Falorca, J.P.N.D. Miraldes, J.C.G. Lanzinha, New trends in visual inspection of buildings and structures: study for the use of drones. *Open Eng.* **11**, 734–743 (2021)
 18. P. Venegas, J. Guerediaga, L. Vega, J. Molleda, F.G. Bulnes, Review infrared thermography for temperature measurement and non-destructive testing. *Sensors* **14**, 12305–12348 (2014). <https://doi.org/10.3390/s140712305>. OPEN ACCESS sensors

Fuzzy Control Systems

Composite Indicators Design Based on Concordant Information Using Semi-supervised Learning on Aggregated Data



Leonid Lyubchyk , Galyna Grinberg , and Klym Yamkovyi 

Abstract The problem of constructing composite indicators to evaluate complex systems is considered by integrating expert assessments, relative weights of partial indices, and their statistical measurements through a machine-learning approach. A nonlinear learning model for composite indicators is developed using a biased kernel ridge regression method. The issue of achieving optimal concordance between expert and statistical data is tackled with a regularization technique that supports the proposed functional coordination. This approach utilizes statistical measurements of partial indices and expert evaluations derived from aggregated data obtained through preliminary clustering in feature space, thereby reducing model complexity. Within a semi-supervised learning framework, composite indicators are designed for a partially labeled training dataset, where unlabeled data is used to create an additional data graph model-based regularization term or to adjust kernel functions based on the geometric properties of the data cloud, ensuring the smoothing of the resulting model.

Keywords Composite indicator · Concordation · Expert assessment · Graph Laplacian · Kernel learning · Semi-supervised learning · Ridge regression

1 Introduction

Composite indicators (CIs) are quantitative measures of complex systems, defined by partial indices (PIs) reflecting different functional quality aspects. They are often used to create a generalized comparative integral assessment of the efficiency and

L. Lyubchyk (✉) · G. Grinberg · K. Yamkovyi
National Technical University “Kharkiv Polytechnic Institute”, Kharkiv, Ukraine
e-mail: leonid.lyubchyk@khpi.edu.ua

G. Grinberg
e-mail: Galyna.Grinberg@khpi.edu.ua

K. Yamkovyi
e-mail: Klym.Yamkovyi@khpi.edu.ua

quality of groups of objects. This approach enables a comprehensive comparative analysis while simultaneously utilizing expert and statistical information [1]. CIs are widely applied in multi-criteria decision-making contexts, such as ranking and selecting alternatives for designing complex systems, project management, investment planning, policy analysis, public communication, and other practical uses. The specified approach is also utilized in assessing perfection and forecasting trends in the sustainable development of large technical, economic, and social systems. For example, the popular CI, known as the “Leading Economic Index” (LEI), is used to predict the direction of global economic movements. Businesses and investors use this CI to help plan their activities around the expected performance of the economy and protect themselves from economic downturns. The known CI “Human Development Index” (HDI) was created to emphasize that people and their capabilities should be the ultimate criteria for assessing the development of a country, not economic growth alone. A feature of the construction of the CI is the need to use heterogeneous information and involve experts since the problem is not fully formalized.

Numerous methods for constructing CIs have been developed, ranging from heuristic and statistical approaches to modern machine-learning techniques [2–4]. Constructing a CI involves aggregating a set of partial indexes (PIs) that characterize the individual properties or quality indexes of the objects being assessed or compared into a generalized integral index. This index reflects the views and preferences of experts involved in creating a generalized assessment while considering the statistical characteristics of the measured PIs [5]. The corresponding mathematical problem entails creating a CI model as a scalar function that performs convolutions of variables describing the measured values of the PIs.

The usual procedure for developing a CI includes the following steps:

- Defining the phenomenon to be measured. The definition of the concept should give a clear sense of what is being measured by the composite index.
- Selecting a group of partial individual indexes. It should be selected according to relevance, analytical soundness, timeliness, accessibility, etc.
- Normalizing the individual indicators. This step aims to make the indicators comparable. Normalization is required before any data aggregation as the indicators in a data set often have different measurement units.
- Aggregating the normalized individual indicators. It combines all the components to form one or more composite indices (as scalar mathematical functions).

A relatively straightforward linear CI model is often used as a linear combination of PIs with specific weighting coefficients that define their influence on the generalized complex assessment [6]. The traditional approach views these weighting coefficients as indicators of the relative importance of individual PIs. In this context, various modifications of the hierarchical process analysis method are applied to identify them, using paired comparison algorithms based on expert opinions [7]. Alternatively, statistical information regarding PI measurements can be utilized, in this case, the principal component analysis (PCA) method is typically used to estimate the weighting coefficients [8]. In developing this approach, using information

entropy to weight partial indicators demonstrates the potential to enhance discriminatory ability [9]. It should be noted that the last approach, despite its popularity, is often criticized because justifying the influence of the statistical dispersion of individual PIs on the overall quality of the system under study is usually difficult. Due to socioeconomic variables' heterogeneous and correlated nature, relative nominal weights are unlikely to align with relative main effects. This leads to the paradox of linear aggregation: weights are treated as if they are importance coefficients when, in fact, they act as trade-off coefficients. Linear aggregation rules have faced criticism because strengths in some dimensions balance weaknesses in others; this quality is called "compensatory" [10]. PCA is an empiricist method that relies on observed correlations while overlooking the polarity of individual indicators. The PCA-based index is often elitist, favoring highly intercorrelated indicators and neglecting others, regardless of their potential contextual significance. As a result, substantiating the impact of the statistical variance of individual performance indicators on the overall quality of the system being studied becomes challenging [11].

In this regard, the following ways of improving the procedures for constructing the CI can be identified:

- The simultaneous use of expert and statistical information is preferred when constructing a CI model. This process involves integrating expert assessments of indicators for a set of objects with the measured values of related PI. Consequently, constructing a CI model is analogous to constructing a regression model; for expert assessments of PI on a point scale, the ordinal regression method can be applied. This introduces the challenge of aligning expert assessments of CI, PI weights, and statistical measurements of PI when developing a generalized model, a procedure known as their concordance.
- A transition to nonlinear CI models is necessary since a linear model does not always adequately represent the true structure of expert preferences, which can be quite complex. The intricate dependence of CI on PI is not predetermined, requiring a highly complex model to accurately capture the actual structure of preferences. This creates the challenge of constructing nonlinear, nonparametric CI models that remain limited in complexity.

In this regard, a promising approach to constructing CI models involves utilizing direct expert assessments of known related PIs. Thus, developing a CI model essentially reduces to creating a regression model; ordinal regression can be applied for expert assessments of CI on a scoring scale [9], where variations of specific indicators serve as explanatory variables and expert assessments of CI act as the response variable. The complexity of implementing the regression approach arises from the intricate and a priori unknown nature of the CI model's structure, which reflects the experts' preferences. Since simple linear models often lack the necessary predictive power, there is a need to develop sufficiently complex nonlinear models that represent expert preferences and can be implemented within the machine learning framework [12].

This problem can be understood within a specific concept in machine learning, particularly preference learning, as the CI model should reflect experts' preferences.

In this case, a nonparametric approach implemented in a kernel-based learning scheme is preferred for typical scenarios with a small training sample due to the complex nature of the expert preference function, which has an unknown structure. The corresponding quasi-linear model is an additive combination of kernel functions linked to the points in the training sample, allowing for an effective approximation of the complex latent dependencies of expert preferences. Expert estimates regarding the relative significance of the weights of PI can provide a priori information to regularize the problem. Such an approach to constructing a regularizing functional in the problem of CI model training can be considered an optimal concordation of expert estimates of CI values and PI weights.

Within the framework of this approach, effective computational procedures for developing nonlinear CI models using kernel learning and the optimal concordance of various expert-statistical information were proposed [12, 13]. The continued advancement of methods for constructing CI based on machine learning techniques is associated with considering the actual structure and size of the training dataset. This indicates a limited capacity to leverage expert assessments, as involving experts incurs certain costs, underscoring the need to limit their participation. The development of CI learning models employing a dataset that partially includes expert information based on semi-supervised learning methods was explored in [14, 15].

In this context, another crucial task arises, which pertains to limiting the complexity of the CI model. Under the kernel learning framework, the number of kernel model components is determined by the size of the training sample, which can lead to considerable model complexity. This may result in overfitting effects, significant irregularities in the model due to the impact of PI measurement errors, and the emergence of false effects within its structure. An additional objective arises to reduce the model's dimensionality and smooth it out, which can be achieved through preliminary data aggregation based on clustering in the feature space, representing the results of PI measurements, this approach embodies the concept of learning on clusters [16, 17].

This work presents a unified approach to constructing nonlinear CI models using kernel semi-supervised learning methods with aggregated data. Within this framework, a method is developed to optimally concordat the expert assessment of CI with expert and/or statistical information regarding PI measurements and their relative weights. This is incorporated into the regularization functional within the biased ridge regression scheme. This approach utilizes statistical measurements of PI and expert assessments of objects derived from aggregated data obtained through preliminary clustering in feature space, ultimately reducing model complexity. Regularized clustering is used, which considers not only the distance between the points in the training sample within feature space but also the additional requirement for the proximity of expert assessments of the CI within clusters. The unlabeled dataset offers additional regularization through a training dataset graph model, which smooths the CI model by incorporating an extra data graph-based regularization term or modifying kernel functions based on the geometric properties of the data cloud. This approach facilitates the reduction of complexity in the resulting CI model and smooths out the undesirable effects of spatial fluctuations, partially eliminating them.

2 Problem Statement

The problem of a CI design for a set of n objects is considered, where a PI vector characterizes each of the objects $\mathbf{x} = [x^1 \ x^2 \ \dots \ x^p]^T$ and p is the number of PIs. The unknown scalar function $J(\mathbf{x})$ describes the CI model as subject to restoration based on available information.

We accept the CI model in a quasilinear parametrized form $J(\mathbf{x}, \mathbf{w}) = \varphi^T(\mathbf{x})\mathbf{w}$, where $\varphi(\mathbf{x}) = [\varphi_1(\mathbf{x}) \ \varphi_2(\mathbf{x}) \ \dots \ \varphi_m(\mathbf{x})]^T$ is defined as the vector of basis coordinate functions and $\mathbf{w} = [w_1 \ w_2 \ \dots \ w_m]^T$ is the CI model parameters vector. In the specific case of a linear CI model $J(\mathbf{x}, \mathbf{w}) = \mathbf{x}^T \bar{\mathbf{w}}$, a components of $\bar{\mathbf{w}} = [\bar{w}_1 \ \bar{w}_2 \ \dots \ \bar{w}_p]^T$ are treated as relative PI weights. In the problem under consideration, a reasonable choice of a set of basis coordinate functions is significantly complicated due to the lack of reliable information about the complex structure of the desired model $J(\mathbf{x}, \mathbf{w})$.

It is assumed, that the following expert and statistical information is available:

- Statistical information in the form of a data matrix of measurements of PI vectors $\mathbf{X}_n = [x_{ij}^j]_{i,j=1}^{n,p} = [\mathbf{x}_1 \ \mathbf{x}_2 \ \dots \ \mathbf{x}_n]^T$ for a set of n objects;
- Expert information, specified as a vector $\mathbf{y}_n = [y_1 \ y_2 \ \dots \ y_n]^T$ of the expert scores y_i , $i = \overline{1, n}$, for each object's CI from the considered set;
- Information concerning vector of relative PI weights $\bar{\mathbf{w}} = [\bar{w}_1 \ \bar{w}_2 \ \dots \ \bar{w}_p]^T$ obtained by expert evaluation using either the pairwise comparisons method or statistical estimation using the PCA approach.

To build a reduced-dimensionality model of restricted complexity, we will distinguish two cases depending on the method used to form the training dataset.

Partial expert information. It is assumed that, within the structure of the training dataset, a group of observations is selected, where each object with a measured vector of PIs is linked to the corresponding value of an expertly estimated CI (labeled data). Additionally, there exists a group of observations for which such expert estimates are unavailable (unlabeled data). This scenario is more realistic, as involving experts incurs certain material costs, clarifying the rationale for limiting it. This case of partially labeled data aligns with the concept of semi-supervised learning, and the problem of CI reconstruction is reduced to estimating the model $J(\mathbf{x}, \mathbf{w})$ parameters $\mathbf{w}$ using available training datasets $\{\mathbf{x}_i, y_i\}_{i=1}^{i=l}, \{\mathbf{x}_i\}_{i=l+1}^{i=n}$.

Aggregated data. In this case, the clustering of objects in the feature space has been completed $I_k = \{i : \mathbf{x}_i \in C_k\}$, $i = \overline{1, n}$, $k = \overline{1, q}$, where I_k is a set of indices of objects belonging to the k -th cluster, q defined the total number of clusters.

Considering the ultimate goal of forming an aggregated dataset for constructing a reduced-order smoothed CI model, we use a regularized modification of the k-means clustering method. Clusters are determined as discrete optimization problem solution

$$\begin{aligned}
& \min_{C_k, \bar{\mathbf{x}}_k} \sum_{k=1}^q \sum_{i \in I_k} \|\mathbf{x}_i - \bar{\mathbf{x}}_k\|^2 + \omega \sum_{k=1}^q R_k(\bar{y}_k, C_k), \\
& R_k(\bar{y}_k, C_k) = \sum_{i \in I_k} (y(\mathbf{x}_i) - \bar{y}_k)^2, \quad I_k = \{i : \mathbf{x}_i \in C_k\},
\end{aligned} \tag{1}$$

where $\bar{\mathbf{x}}_k$ are the centroid of the C_k cluster in feature space, $\bar{y}_k$ is the average of the CI expert scores assigned to objects in each cluster and $\omega > 0$ is the regularization coefficient. This choice of regularizing functional ensures that clusters are formed by minimizing the distance of their elements from the centroid while introducing an additional penalty that depends on the spread of the CI estimates for all elements of the corresponding cluster. Consequently, the proposed formation of an aggregated dataset further smooths the reconstructed CI model.

Next, we use an aggregated dataset $\{\bar{\mathbf{x}}_k, \bar{y}_k\}_{k=1}^q$ to construct a kernel-based reduced-order model of the CI. When implementing a semi-supervised learning technique, except for the marked matrix of aggregated PI measurements $\bar{\mathbf{X}}_q$, the complete set of PI measurements for all objects $\mathbf{X}_n = \{\mathbf{x}_i\}_{i=1}^{i=n}$ can be used as an unlabeled dataset.

The measurement models, which used the full (a) $\{\mathbf{x}_i, y_i\}_{i=1}^{i=n}$ and partially-labeled (b) $\{\mathbf{x}_i, y_i\}_{i=1}^{i=l}$ training datasets are the following:

$$\begin{aligned}
(a) \quad & y_i = \varphi^T(\mathbf{x}_i) \mathbf{w} + \varepsilon_i, \quad i = \overline{1, n}, \\
& \mathbf{y}_n = \Phi_n \mathbf{w} + \mathbf{\varepsilon}_n, \quad \Phi_n^T = [\varphi(\mathbf{x}_1) \ \varphi(\mathbf{x}_2) \ \dots \ \varphi(\mathbf{x}_n)], \\
(b) \quad & y_i = \varphi^T(\mathbf{x}_i) \mathbf{w} + \varepsilon_i, \quad i = \overline{1, l}, \\
& \mathbf{y}_l = \Phi_l \mathbf{w} + \mathbf{\varepsilon}_l, \quad \Phi_l^T = [\varphi(\mathbf{x}_1) \ \varphi(\mathbf{x}_2) \ \dots \ \varphi(\mathbf{x}_l)],
\end{aligned} \tag{2}$$

where ε_i are the errors that model inaccuracies in the expert assessments of CIs.

Similarly for the aggregated dataset formed based on clustered data

$$\begin{aligned}
& \bar{y}_k = \varphi^T(\bar{\mathbf{x}}_k) \mathbf{w} + \bar{\varepsilon}_k, \quad k = \overline{1, q}, \\
& \bar{\mathbf{y}}_q = \bar{\Phi}_q \mathbf{w} + \bar{\mathbf{\varepsilon}}_q, \quad \bar{\Phi}_q^T = [\varphi(\bar{\mathbf{x}}_1) \ \varphi(\bar{\mathbf{x}}_2) \ \dots \ \varphi(\bar{\mathbf{x}}_q)].
\end{aligned} \tag{3}$$

3 Materials and Methods

3.1 Optimal Concordat of Expert Information

To avoid the difficulties related to the choice of coordinate functions $\varphi(\mathbf{x})$ and to restrict the complexity of the CI model, we use a kernel-based machine learning technique, introducing into consideration Reproducing Kernel Hilbert Space (RKHS), induced by a positive semi-definite kernels $\kappa(\mathbf{x}, \mathbf{z})$ [18].

Herein, coordinate functions $\{\varphi(\bar{\mathbf{x}}_i)\}_{i=1}^m$ are taken hereby that its dot products $\varphi^T(\mathbf{x})\varphi(\mathbf{z}) = \kappa(\mathbf{x}, \mathbf{z})$. We will use the Gaussian kernel $\kappa(\mathbf{x}, \mathbf{z}) = \exp\{-\|\mathbf{x} - \mathbf{z}\|^2/2\sigma^2\}$ with the bandwidth parameter $\sigma > 0$.

Let's consider constructing a reduced-order kernel-based CI model using an aggregated dataset $\{\bar{\mathbf{x}}_k, \bar{y}_k\}_{k=1}^q$ through the measurement Eq. (3). Define the optimal estimates of the CI model parameters as a solution to the equality-type constrained optimization problem, derived from the kernel SVM scheme [18]:

$$\begin{aligned} I(\mathbf{w}) \rightarrow \min_{\boldsymbol{\varepsilon}_q, \boldsymbol{\eta}_q, \mathbf{w}, \mathbf{w}_0}, \quad 2I(\mathbf{w}) &= \|\boldsymbol{\varepsilon}_q\|^2 + \alpha \|\mathbf{w} - \mathbf{w}_0\|^2 + \beta \|\mathbf{w}_0\|^2 + \|\boldsymbol{\eta}_q\|^2, \\ \bar{\mathbf{y}}_q - \bar{\Phi}_q \mathbf{w} &= \bar{\boldsymbol{\varepsilon}}_q, \quad \bar{\mathbf{X}}_q \bar{\mathbf{w}} - \bar{\Phi}_q \mathbf{w}_0 = \boldsymbol{\eta}_q, \end{aligned} \quad (4)$$

where $\alpha > 0$, $\beta > 0$ —regularization parameters.

Expert-estimated PI weights $\bar{\mathbf{w}}$ are used as a priori information for regularization purposes. It becomes necessary to express the vector of a *prior* value of the model parameters $\mathbf{w}_0$ through the initial data $\{\bar{\mathbf{X}}_q, \bar{\mathbf{w}}\}$, which implements, in fact, the optimal concordant of expert-statistical information in the construction of CI. Thus, the proposed approach implements an expert-statistical information concordat using a biased kernel ridge regression. The central concept is to choose a *prior* vector $\mathbf{w}_0$ of model parameters in (2) based on the proximity condition of the predicted vector of the CI's *prior* values $\bar{\Phi}_q \mathbf{w}_0$ to the estimates obtained from its linear model $\bar{\mathbf{X}}_q \bar{\mathbf{w}}$, which is considered as the first approximation.

Let us introduce the Lagrange function

$$\begin{aligned} L(\mathbf{w}, \mathbf{w}_0, \boldsymbol{\lambda}, \boldsymbol{\mu}) &= 2(\|\boldsymbol{\varepsilon}_q\|^2 + \alpha \|\mathbf{w} - \mathbf{w}_0\|^2 + \beta \|\mathbf{w}_0\|^2 + \|\boldsymbol{\eta}_q\|^2) \\ &+ \boldsymbol{\lambda}^T (\bar{\mathbf{y}}_q - \bar{\Phi}_q (\bar{\mathbf{X}}_q \mathbf{w}) - \bar{\boldsymbol{\varepsilon}}_q) + \boldsymbol{\mu}^T (\bar{\mathbf{X}}_q \bar{\mathbf{w}} - \bar{\Phi}_q (\bar{\mathbf{X}}_q \mathbf{w}_0) - \boldsymbol{\eta}_q), \end{aligned} \quad (5)$$

where $\boldsymbol{\lambda}$, $\boldsymbol{\mu}$ —vectors of Lagrange multipliers of appropriate dimensions.

The necessary extremum conditions for the constrained optimization problem (2) derived from the Lagrange function (3) give the normal equations relating to the parameters of the CI model and the Lagrange multipliers.

$$\begin{bmatrix} \alpha \mathbf{I}_m & -\alpha \mathbf{I}_m \\ -\alpha \mathbf{I}_m & (\alpha + \beta) \mathbf{I}_m \end{bmatrix} \begin{bmatrix} \mathbf{w} \\ \mathbf{w}_0 \end{bmatrix} = \begin{bmatrix} \bar{\Phi}_q^T \boldsymbol{\lambda} \\ \bar{\Phi}_q^T \boldsymbol{\mu} \end{bmatrix}, \quad (6)$$

wherein $\boldsymbol{\lambda} = \boldsymbol{\varepsilon}_q$, $\boldsymbol{\mu} = \boldsymbol{\eta}_q$, and $\mathbf{I}_m$ is the identity matrix of the corresponding dimension.

Using the Frobenius formula to find $\mathbf{w}$, $\mathbf{w}_0$ and substituting their values into equations of constraints from (2), we obtain

$$\begin{bmatrix} \gamma^{-1} \mathbf{K}_{q,q} + \mathbf{I}_n & \beta^{-1} \mathbf{K}_{q,q} \\ \beta^{-1} \mathbf{K}_{q,q} & \beta^{-1} \mathbf{K}_{q,q} + \mathbf{I}_n \end{bmatrix} \begin{bmatrix} \boldsymbol{\varepsilon}_n \\ \boldsymbol{\eta}_n \end{bmatrix} = \begin{bmatrix} \mathbf{y}_n \\ \mathbf{X}_q \bar{\mathbf{w}} \end{bmatrix}, \quad (7)$$

where kernel matrix $\mathbf{K}_{q,q} = \overline{\Phi}_q \overline{\Phi}_q^T$, and $\gamma = (\alpha + \beta)/\alpha\beta$.

Whence follows the expression for the optimal CI kernel model parameters [13]

$$\mathbf{w}_n^* = [\gamma^{-1} \Phi^T(\mathbf{X}_q) \quad \beta^{-1} \Phi^T(\mathbf{X}_q)] \cdot \begin{bmatrix} \mathbf{K}_{q,q}(\gamma^{-1}) & \beta^{-1} \mathbf{K}_{q,q} \\ \beta^{-1} \mathbf{K}_{q,q} & \mathbf{K}_{q,q}(\beta^{-1}) \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{y}_q \\ \mathbf{X}_q \overline{\mathbf{w}} \end{bmatrix}, \quad (8)$$

where the inverse block matrix in (6) taking into account the designations $\mathbf{K}_{q,q}(\otimes^{-1}) = \otimes^{-1} \mathbf{K}_{q,q} + \mathbf{I}_q$ is defined as

$$\begin{aligned} & \begin{bmatrix} \mathbf{K}_{q,q}(\gamma^{-1}) & \beta^{-1} \mathbf{K}_{q,q} \\ \beta^{-1} \mathbf{K}_{n,n} & \mathbf{K}_{q,q}(\beta^{-1}) \end{bmatrix}^{-1} \\ &= \begin{bmatrix} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) + \beta^{-2} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) \mathbf{K}_{q,q} \Lambda^{-1} \mathbf{K}_{q,q} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) & -\beta^{-1} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) \mathbf{K}_{q,q} \Lambda^{-1} \\ -\beta^{-1} \Lambda^{-1} \mathbf{K}_{q,q} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) & \Lambda^{-1} \end{bmatrix} \end{aligned} \quad (9)$$

wherein according to the Sherman–Morrison–Woodbury identity

$$\begin{aligned} \Lambda &= \mathbf{K}_{q,q}(\beta^{-1}) - \beta^{-2} \mathbf{K}_{q,q} \mathbf{K}_{q,q}^{-1}(\gamma^{-1}) \mathbf{K}_{q,q}, \\ \Lambda^{-1} &= \mathbf{K}_{q,q}^{-1}(\beta^{-1}) - \beta^{-2} \mathbf{K}_{q,q}^{-1}(\beta^{-1}) \mathbf{K}_{q,q} \mathbf{R}^{-1} \mathbf{K}_{n,n} \mathbf{K}_{q,q}^{-1}(\beta^{-1}), \\ \mathbf{R} &= \mathbf{K}_{q,q}(\gamma^{-1}) + \beta^{-2} \mathbf{K}_{q,q} \mathbf{K}_{q,q}^{-1}(\beta^{-1}) \mathbf{K}_{q,q}. \end{aligned} \quad (10)$$

From the evident relation $\varphi^T(\mathbf{x}) \Phi^T = \kappa^T(\mathbf{x})$, where the kernel vector is defined as $\kappa(\mathbf{x}) = [\kappa(\mathbf{x}, \bar{\mathbf{x}}_1), \kappa(\mathbf{x}, \bar{\mathbf{x}}_2) \dots \kappa(\mathbf{x}, \bar{\mathbf{x}}_q)]^T$ and specified at points in the feature space that coincide with the cluster centroids, we derive the final expression for the reduced-order kernel-based CI model

$$\hat{J}(\mathbf{x}) = \kappa^T(\mathbf{x})(\gamma^{-1} \mathbf{A} \bar{\mathbf{y}}_q + \beta^{-1} \mathbf{B} \bar{\mathbf{X}}_q \bar{\mathbf{w}}), \quad (11)$$

where numerical matrices, whose elements depend only on the kernel matrix $\mathbf{K}_{q,q} = [\kappa(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j)]_{i,j=1}^{q,q}$, defined on the cluster centroids.

Thus, the derived CI model defines the general representation of the kernel model in accordance with the representation theorem and illustrates the optimal alignment of various expert information, taking into account the PI's statistical measurements.

3.2 Semi-supervised CI Model Learning

Consider constructing a CI model on a partially labeled dataset where expert assessments are limited to a relatively small number of objects, while the remaining PI measurements are utilized to smooth the resulting kernel model. We will use the

training dataset $\{\mathbf{x}_i, y_i\}_{i=1}^{i=l}, \{\mathbf{x}_i\}_{i=l+1}^{i=n}$, an unlabeled part provides additional regularization using the graph data model by smoothing of desired CI model on the data cloud, taking into account the distance between its elements [19, 20].

Let's use the following graph data model, it is assumed that the adjacent data graph consists of nodes that are all observed data points $\mathbf{X}_n$, corresponding to the measured values of PIs. The matrix $\mathbf{V} = \|v_{ij}\|$ determines edge weights v_{ij} in the data adjacency graph, wherein the edge, connecting two data graph nodes $\mathbf{x}_i$ and $\mathbf{x}_j$, is weighted by kernel function $v_{ij} = \kappa(\mathbf{x}_i, \mathbf{x}_j)$ defined over graph nodes.

To obtain estimates of the model parameters, we introduce functional that includes two types of regularizer, namely, to consider a priori information about the weights of the PI components and to smooth the model on the data cloud.

Wherein estimates of CI model parameters were obtained by the minimization of regularized functional

$$2R_c(\mathbf{y}_l, \mathbf{w}) = \|\mathbf{y}_l - \Phi_l^T \cdot \mathbf{w}\|^2 + \alpha \cdot \Omega_1(\mathbf{w}, \mathbf{w}_0) + \beta \cdot \Omega_2(\mathbf{J}_n, \mathbf{X}_n), \quad (12)$$

$$\Phi_l = (\varphi(\mathbf{x}_1) \varphi(\mathbf{x}_2) \dots \varphi(\mathbf{x}_l)), \quad \mathbf{J}_n = (J(\mathbf{x}_1), J(\mathbf{x}_2), \dots, J(\mathbf{x}_n))^T.$$

where $\Omega_1(\mathbf{w}, \mathbf{w}_0) = (\mathbf{w} - \mathbf{w}_0)^T \cdot (\mathbf{w} - \mathbf{w}_0)$ is a Tikhonov regulariser, taking into account a priori expert information about the relative significance weights of PIs $\mathbf{w}_0$, and $\Omega_2(\mathbf{w}, \mathbf{X}_n) = \sum_{i,j=1}^n v_{ij} (J(\mathbf{x}_i) - J(\mathbf{x}_j))^2$ is *data graph* model-based regularization term [20], smoothing the evaluated CI model $J(\mathbf{x})$ on the full training dataset.

The manifold regularization term is chosen [18] as a data-dependent semi-norm $\Omega_2(\mathbf{J}_n, \mathbf{L}) = \mathbf{J}_n^T \cdot \mathbf{L} \cdot \mathbf{J}_n$, where $\mathbf{L}$ is graph Laplacian matrix given by $\mathbf{L} = \mathbf{D} - \mathbf{V}$, with edge weights matrix $\mathbf{V} = \|v_{ij}\|$ and diagonal matrix $\mathbf{D}$, for which its (i, i) element is equal to the sum of the i -th column of the matrix $\mathbf{V}$.

Under a semi-supervised kernel-based framework, we use data-dependent kernel functions corresponding to semi-norm $\Omega_2(\mathbf{J}_n, \mathbf{L})$. Explicit form of transformed reproducing kernel $\tilde{\kappa}(\mathbf{x}, \mathbf{x}')$ derived in [21, 22] are given as:

$$\tilde{\kappa}(\mathbf{x}, \mathbf{x}') = \kappa(\mathbf{x}, \mathbf{x}') + \kappa_n^T(\mathbf{x}) \cdot (\mathbf{I}_n + \mathbf{L} \mathbf{K}_{n,n})^{-1} \mathbf{L} \cdot \mathbf{k}_n(\mathbf{x}), \quad (13)$$

$$\kappa_n^T(\mathbf{x}) = (\kappa(\mathbf{x}, \mathbf{x}_1), \kappa(\mathbf{x}, \mathbf{x}_2), \dots, \kappa(\mathbf{x}, \mathbf{x}_n)).$$

where $\mathbf{K}_{n,n} = \{\kappa(\mathbf{x}_i, \mathbf{x}_j)\}_{i,j=1}^{n,n}$ is the kernel matrix.

Transformation (11) describes initial kernel deformation along finite-dimensional space given by data. In the presence of unlabeled data, such a choice of kernel functions implements the smoothness assumption concerning training data $\mathbf{X}_n$ geometric structure. The resulting transformed kernels can be used to modify the indicator model estimates obtained previously within the supervised kernel learning framework.

Then, model parameters estimated $\hat{\mathbf{w}}_l$ under can be represented as:

$$\begin{aligned}\hat{\mathbf{w}}_l &= \Phi_l^T A_\gamma^{-1} \mathbf{y}_l + \Phi_l^T (\mathbf{I}_l - A_l^{-1}(\alpha) \tilde{\mathbf{K}}_{l,l}) A_l^{-1}(\gamma) \mathbf{X}_l \hat{\mathbf{w}}_n, \\ A(\gamma) &= \gamma \mathbf{I}_l + \tilde{\mathbf{K}}_{l,l}, A(\alpha) = \alpha \mathbf{I}_l + \tilde{\mathbf{K}}_{l,l},\end{aligned}\quad (14)$$

where $\tilde{\mathbf{K}}_{l,l} = \{\tilde{k}(\mathbf{x}_i, \mathbf{x}_j)\}_{i,j=1}^{l,l}$ is a transformed kernel matrix, built on the labeled part of the training dataset, $\alpha > 0$, $\gamma > 0$ are regularization parameters and $\mathbf{w}_0$, as before, is determined based on prior information.

Considering that $\varphi^T(\mathbf{x}) \cdot \Phi_l^T = (\tilde{k}(\mathbf{x}, \mathbf{x}_1), \dots, \tilde{k}(\mathbf{x}, \mathbf{x}_l))$, smoothed CI model can be represented as a linear combination of modified kernel functions

$$\hat{\mathbf{f}}(\mathbf{x}) = \tilde{\mathbf{k}}_l^T(\mathbf{x}) \cdot \mathbf{g}_l, \quad \mathbf{g}_l = A_\gamma^{-1} \mathbf{y}_l + (\mathbf{I}_l - A_l^{-1}(\alpha) \cdot \tilde{\mathbf{K}}_l) A_l^{-1}(\gamma) \mathbf{X}_l \mathbf{w}_0 \quad (15)$$

with data-dependent model weight coefficients $\mathbf{g}_l = (g_1, g_2, \dots, g_l)^T$.

4 Conclusion

The proposed approach's advantage lies in restoring a complex structure of experts' preferences within the CI model-building process while utilizing a model of restricted complexity. Simultaneously, the modified regularizer enhances the stability of computational procedures and the optimal alignment of heterogeneous expert and statistical information. This optimal alignment of expert CI assessments and PI relative weights facilitates the integration of prior expert data into the structure of the obtained estimates, allowing effective regularization to occur concurrently.

For the practical implementation of the kernel-based learning approach, limiting complexity is beneficial to ensure the smoothness of the reconstructed CI model and prevent spatial oscillation effects. Future research may focus on developing methods for training the CI model using dynamic aggregated information, such as on-line clustering preliminary measured PI data to reduce model dimensions. Regularized "soft" clustering shows promise, as clustering in PI space improves by minimizing the spread of PI estimates within corresponding clusters. Another significant area of exploration involves devising strategies for aggregating partial indicators to reduce model complexity while addressing the imbalance among indicators. Additionally, combining semi-supervised clustering with model-building is of interest, as it may enable the optimal use of information in the training dataset.

References

1. M. Nardo et al., *Handbook on Constructing Composite Indicators: Methodology and User Guide* (OECD Publishing, Paris, France, JRC European Commission, 2008)
2. El. Gibari, S. Gómez, F. Ruiz, Building composite indicators using multicriteria methods: a review. *J. Bus. Econ.* **89**, 1–24 (2019)

3. E. Fernandez, M. Martos, Review of some statistical methods for constructing composite indicators. *Stud. Appl. Econ.* **38**(1), 1–15 (2020)
4. A. Sarra, E. Nissi, A. Evangelista et al., A Functional approach for constructing dynamic composite indicators. *Stat. Methods Appl.* **33**, 173–204 (2024)
5. D. Lindén, M. Cinelli, M. Spada, W. Becker, P. Gasser, P. Burgher, A framework based on statistical analysis and stakeholders' preferences to inform weighting in composite indicators. *Env. Modeling & Software* **145**, 105208 (2021)
6. W. Becker, M. Saisana, P. Paruolo, I. Vandecasteele, Weights and importance in composite indicators: Closing the gap. *Ecol. Ind.* **80**, 12–22 (2017)
7. W. Zalatar, E. Clark, Constructing a composite indicator for manufacturing companies using lean metrics and analytic hierarchy process, in *2021 IEEE Intern. Conference on Industrial Engineering and Engineering Management* Singapore (2021).
8. K. Boudt, M. d'Errico, H. Luu, R. Pietrelli, Interpretability of composite indicators based on principal components. *Hindawi J. Probab. Stat.* 4155384 (2022)
9. M. Liborio, R. Karagiannis, A. Diniz, P. Ekel, D. Vieira, L. Ribeiro, The use of information entropy and expert opinion in maximizing the discriminating power of composite indicators. *Entropy* **26**(143) (2024)
10. P. Paruolo, M. Saisana, A. Saltelli, Ratings and rankings: voodoo or science? *J. R. Stat. Soc. Ser. A* **176**(3), 609–634 (2013)
11. M. Mazziotta, A. Pareto, Composite index construction by PCA? No, thanks, in *Conference: 52 Riunione Scientifica SIEDS* (Ancona, Italy, 2015)
12. L. Lyubchyk, G. Grinberg, Preference function reconstruction for multiple criteria decision making based on machine learning approach, in *Recent developments and new directions in Soft Computing* ed. by L. Zadeh et al. (Springer, 2014), pp. 53–63
13. L. Lyubchyk, G. Grinberg, Z. Konokhova, K. Yamkovyi, Composite indicators building based on concordant of expert-statistical information using biased ridge kernel regression, in *2023 IEEE 12th International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS)* (Dortmund, Germany, 2023), pp. 617–620
14. L. Lyubchyk, G. Grinberg, K. Yamkovyi, Integral indicator for complex system building based on semi-supervised learning, in *2018 IEEE First International Conference on System Analysis & Intelligent Computing (SAIC)* (Kyiv, Ukraine, 2018), pp. 1–5
15. E. Jiménez-Fernández, A. Sánchez, M. Ortega-Pérez, Dealing with weighting scheme in composite indicators: an unsupervised distance-machine learning proposal for quantitative data. *Socioecon. Plann. Sci.* **83**, 101339 (2022)
16. A. Verma, O. Angelini, T. Di Matteo, A new set of clusters driven composite development indicators. *EPJ Data Sci.* **9** (2020)
17. F. Mariani, M. Ciommi, Aggregating composite indicators through the geometric mean: a penalization approach. *Computation* **10**(4), 64 (2022)
18. N. Cristianini, J. Shawe-Taylor, *An Introduction to Support Vector Machines and Other Kernel-Based Learning Methods* (Cambridge University Press, Cambridge, 2000)
19. X. Guo, K. Uehara, Graph-based semi-supervised regression and its extensions. *Intern. J. Adv. Comput. Sci. Appl.* **6**(6), 260–269 (2015)
20. D. Cai, X. He, J. Han, Semi-supervised regression using spectral techniques, in *Report No. UIUCDCS-R-2006-2749* (University of Illinois at Urbana-Champaign, 2006)
21. A. Pozdnoukhov, S. Bengio, Semi-supervised kernel methods for regression estimation, in *Proceedings of IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (Toulouse, France, 2006)
22. V. Sindhwani, P. Niyogi, M. Belkin, Beyond the point cloud: from transductive to semi-supervised learning, in *Proceedings of ICML'05* (Bonn, Germany, 2005)

Software of Artificial Intelligence Control System for UAV



Azad Bayramov  and Samir Suleymanov 

Abstract The control of technical system in difficult conditions limits the capabilities of the unmanned aerial vehicle and reduces the efficiency of the entire system. Therefore, the autonomous navigation of an unmanned aerial vehicle using artificial intelligence technology is an urgent task. The aircraft can perform tasks depending on the variability of external factors. The paper presents result of developing an algorithm of software for controlling and navigating an unmanned vehicle using the artificial intelligence technology. An algorithm has been developed using a visual and inertial navigation system to determine the coordinate and time characteristics of an unmanned aerial vehicle in the event of non-functioning of the global satellite navigation system. In this case, the developed algorithm for the functioning of the artificial intelligence of the on-board computer itself determines the coordinate-time characteristics of the unmanned aerial vehicle in autonomous flight mode. The proposed algorithm of software and control architecture based on artificial intelligence are applied to an unmanned aerial vehicle intended to perform a reconnaissance flight.

Keywords Unmanned aerial vehicle · Software · Control system · Artificial intelligence · Algorithm

1 Introduction

It is known that unmanned aerial vehicles (UAVs) cannot perform their tasks optimally due to reliance on human control and limitations of radio communications. Therefore, autonomous UAV navigation plays an important role in achieving an optimal result. Various methods of localization, mapping and detection methods are used to achieve autonomous navigation. Many works [e.g. 1–4] are devoted to these methods.

A. Bayramov (✉) · S. Suleymanov
Republican Seismic Survey Center, Azerbaijan National Academy of Sciences, Baku, Azerbaijan
e-mail: azad.bayramov@yahoo.com

Up to now, the various navigation methods have been proposed: inertial navigation, satellite navigation and observation-based navigation. However, depending on the task, it is necessary to choose the optimal technology for UAV autonomous navigation. One of the effective ways is the Artificial Intelligence (AI) technology, which is widely used in the field of engineering research. AI can detect anomalies and predict the potential of scenarios, respond to changing situations, study complex problems associated with huge amounts of data, and find regularities that a human might ignore. It can analyze and use external factors to improve the maneuvering of the UAV [5–8].

The control technologies are used to increase the autonomy of the UAV to the level of self-management. While there is no universal consensus on how to define or measure an intelligent system, there are several characteristics that an intelligent controller possesses: adaptability, learning ability, non-linearity, autonomous character interpretation, and goal-directedness. In the case of AI-assisted UAV flight, the controller output is determined using input sensors to create an internal representation (recognition) of the environment. In this they differ from UAVs, where pre-established mathematical models are used [4].

2 UAV Integrated Navigation Model

There are many different types of UAVs for military and civilian applications. UAVs are often classified according to characteristics related to shape, flight range, maximum takeoff weight and payload. The payload may include cameras, sensors, mobile phones and bases, cellular assistance stations. The larger the payload, the more equipment, and accessories can be carried.

UAV navigation can be divided into four categories according to application: outdoor navigation, indoor navigation, search and rescue navigation, and wireless network navigation. External navigation includes surveillance, delivery of goods, target tracking, and crowd monitoring. Indoor navigation includes indoor mapping, production automation, and indoor surveillance. In addition, UAV navigation can be categorized by navigation parameters: inertia navigation, signal navigation, and surveillance navigation. For the inertial navigation, UAVs use gyroscopes, accelerometers, and altimeters to control the onboard flight controller. UAVs use GPS modules and a remote radio head in the case of cellular communications for signal-based navigation, and cameras for visual navigation.

The planned path of the UAV during autonomous flight is carried out using various methods of artificial intelligence: optimization-based approaches or learning-based approaches. The UAV flies autonomously [6, 7] using various artificial intelligence methods: optimization-based approaches or learning-based approaches.

At first, the altitude and horizon controllers receive signals from these sensors and the pitch and yaw controllers depending on the trajectory. Then the pitch and yaw controllers control the elevators and maneuvering ailerons of the UAV depending on the signals from these sensors.

Methods and technologies based on mathematical optimization cover traditional mathematical algorithms for solving AI problems. These algorithms make it possible to achieve near-optimal solutions for any given non-deterministic polynomial complex problem. However, these algorithms are quite difficult in terms of characteristics such as time and coordinates [8]. Let considered some of these approaches.

Optimization-based approaches determine the movement of UAV depending on its current position and speed. UAV speed continues to update based on the optimal position vector. The control of UAV flight optimization reaches its optimum position when it reaches its target or at the lowest possible error. In navigation, the optimization treats the UAV as particle and controls their movement in 3D space.

Regarding autonomous UAV navigation, an optimization based solution for multiple movements control has been proposed, avoiding the complexities of 3D space. The multiplicity of control optimization overcomes the problem of premature convergence caused by a single movement. Initially, the UAV navigation problem was formulated as a traveling salesman problem, and then as a multi-part problem when a UAV was looking for the best routes to their destination.

For example, genetic algorithm is an optimization algorithm that starts with a population of randomly generated chromosomes known as a start population. Each gene on a chromosome is a series of numerical numbers. Each chromosome in this study reflects the trajectory of the UAV, which is constrained by the dynamics of the UAV. Genetic operations such as crossbreeding, mutation, selection, insertion and deletion will periodically change the population in each generation. The goal of this procedure is to decrease the fitness function as much as possible by identifying a chromosome with a near minimum fitness value.

Learning-based approaches cover traditional models based on AI algorithms. These algorithms can achieve near-optimal results for any given non-deterministic polynomial-time hard problem with very low complexity.

An integrated navigation system should carry out joint processing of information from a strapdown inertial navigation system, a satellite navigation device and a radio engineering system for measuring altitude and velocity components. At the same time, the main task solved by the integrated navigation system is the development of reliable navigation information for the UAV control system with the required errors and dynamic characteristics.

Inertial sensors for primary information, which are part of the block of sensitive elements of the strapdown inertial navigation system:

- A three-axis accelerometer unit that measures the parameters of the linear motion of an object relative to inertial space (apparent acceleration $\vec{a}_{k1}$);
- Three uniaxial laser gyroscopes that measure the parameters of the angular motion of an object relative to inertial space (angular velocity $\vec{\omega}_1$).

The most significant components that determine instrumental errors are:

- Offset of the zero signal;
- Error of the scale conversion factor;
- Non-orthogonality of the measuring axes of the sensors;
- Casual care.

The inertial navigation system is the main meter of trajectory motion and orientation parameters in the integrated navigation system and provides continuous calculation and output to the user (in the UAV control system) of the required navigation information.

When complexing an inertial navigation system, a satellite navigation device and a radio engineering system for measuring altitude and velocity components, an error model of the integrated navigation system is adopted, including the errors of these systems, as the equations of state of the filter for joint information processing.

3 The UAV Control System Based on AI

Figure 1 show the developed architecture of the control and navigation system based on artificial intelligence for the winged UAV. This algorithm can be applied to any winged UAV.

The control and navigation system consists of separate functional blocks. The control system consists of a radio receiver for receiving commands from the ground station in manual control mode, a Global Navigation Satellite System (GNSS) receiver for receiving and processing data such as flight coordinates, flight (movement) speed and direction, flight altitude, and a telemetry unit for transmitting these and other data to the base. The connection between the blocks of the control system is carried out by the microcontroller board. The microcontroller board initially receives the mission task to be performed by the UAV from the ground station and writes it to the internal memory.

Typically, the 1st scenario (mission) consists of getting the vehicle to the appropriate altitude from the runway, and the last scenario (mission) consists of the necessary parameters to return to base or an alternate landing strip. The 2nd and 3rd... N scenario (missions) consist of the necessary coordinates for the UAV to perform the given task, the height and speed of the mission, and other parameters. Scenario (mission) information can be changed and new tasks can be assigned as long as the device is in the air and in communication with the ground station.

The GNSS module, which receives and processes signals from Global navigation satellite systems such as GPS, GLONASS, GALILEO, BEIDOU, QZSS, SBAS, transmits these data to the microcontroller. The exchange between the GNSS and the microcontroller is based on the NMEA0183 protocol [8].

The control system can operate in manual control, semi-automatic, automatic and fully automatic modes. These algorithms were developed in accordance with these modes. The software was written in the internal permanent memory of the

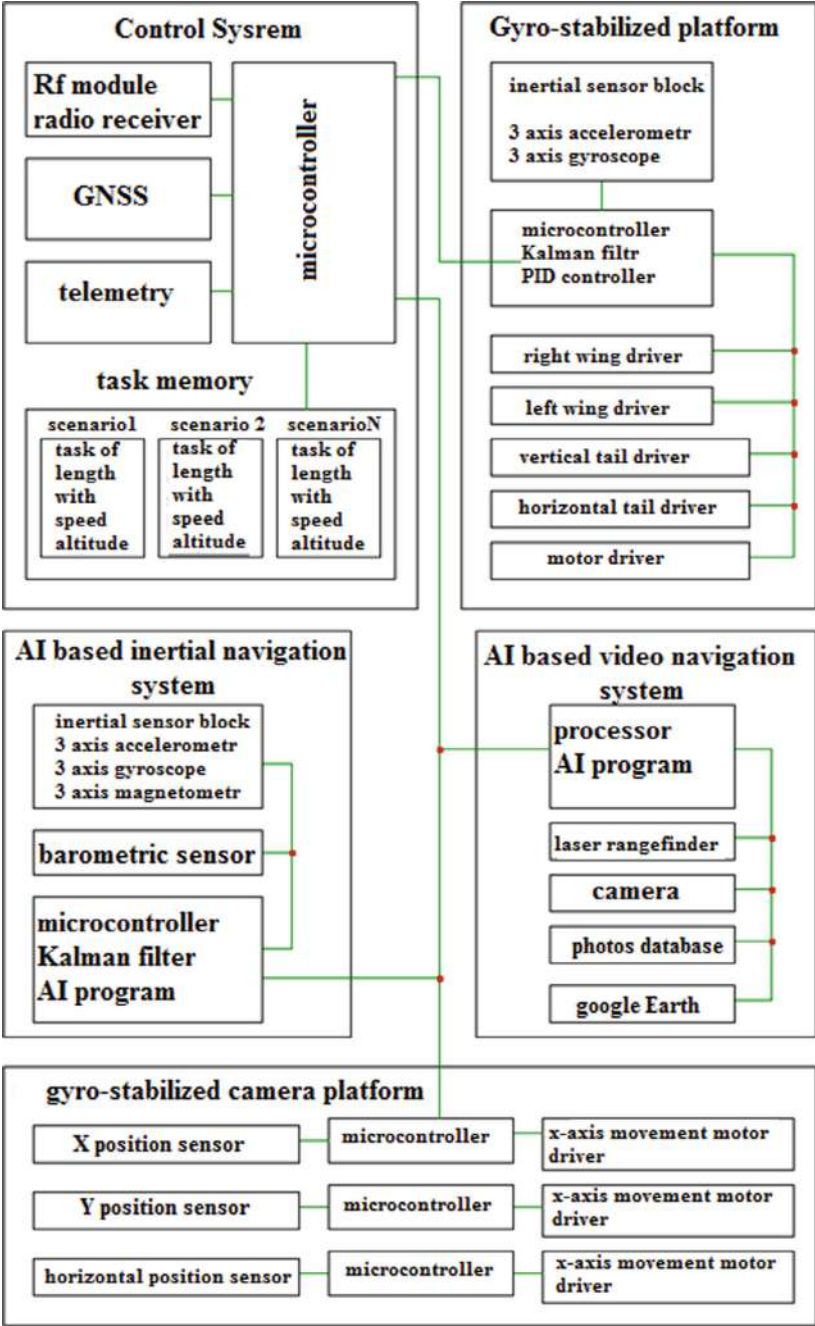


Fig. 1 The architecture of control and navigation system based on artificial intelligence for winged UAV

microcontroller in accordance with the algorithms. In manual control mode, the microcontroller transmits the commands received from the ground station directly to the actuators through the gyro-stabilized platform block.

In semi-automatic mode, the microcontroller of the control system reads the coordinates corresponding to the task from memory, calculates the course and flight (movement) range from its coordinates, and, having gained the required altitude and speed, transmits commands to the actuators to perform the task through block of the gyro stabilized platform. The automatic mode is designed to perform a task in the event of loss of communication with a ground station, and the fully automatic mode is intended to complete a task both in the event of loss of communication with the satellites of the navigation system, and in the absence of communication with the base (in electronic warfare conditions), as in semi-automatic mode.

A gyro-stabilized platform is a hardware and software complex for keeping an UAV in balance in a horizontal plane in the air. The platform consists of a 3-axis accelerometer, a 3-axis gyroscope and a program for processing data received from these devices (based on a Kalman filter), a PID-based program to maintain the stability mechanisms of the UAV, as well as a microcontroller and an actuator. The Kalman filter algorithm works in two stages. During the prediction phase, the Kalman filter extrapolates the values of state variables as well as their uncertainties. At the second stage, based on measurement data (obtained with some error), the extrapolation result is clarified. Due to the step-by-step nature of the algorithm, it can monitor the state of an object in real time using only current measurements and information about the previous state and its uncertainty.

The gyro stabilizer microcontroller sends appropriate commands to the left wing, right wing and horizontal tail drives (motors) to keep the platform stable, i.e. to reset the yaw and pitch angles. In automatic control systems, the PID adjustment algorithm is used to obtain the required accuracy and quality of the switching process. The PID controller generates a control signal, which is the sum of three terms, the first of which is proportional to the difference between the input signal and the feedback signal (mismatch signal), the second to the integral of the mismatch signal, and the third to the derivative of the mismatch signal.

The overall view of developed software algorithm of AI control and navigation system operation for winged UAV is presented in Fig. 2.

To explain the software algorithm, it should be noted that at the first stage, data on flight tasks, frequency characteristics of the satellite radio receiver and telemetry data are transmitted to the control system. Based on satellite signal data, using the Kalman Filter and PID, the UAV coordinates are determined and the corresponding signals are transmitted to electromechanical power drives. If satellite signals disappear, then the AI analysis unit is turned on: based on the data of the inertial and video navigation systems, the main space-time characteristics of the UAV (longitude, latitude, azimuth-course, speed, altitude) are determined. Next, this data is again transmitted to the control system.

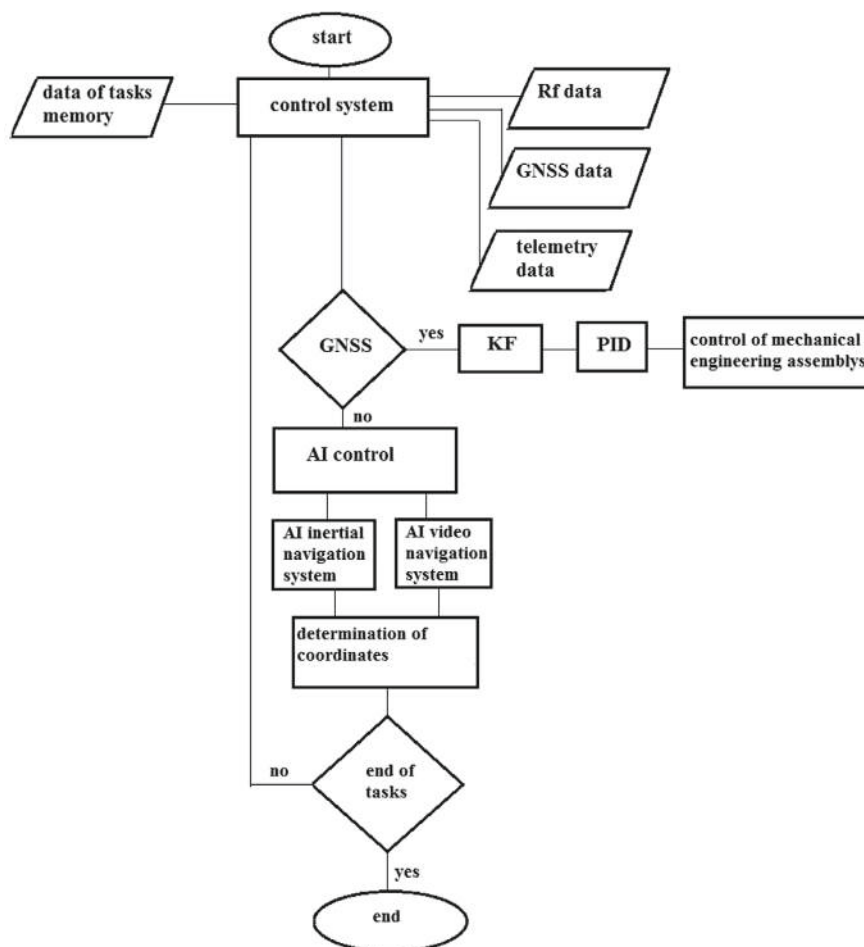


Fig. 2 The algorithm of software of AI control and navigation system for winged UAV

The principled difference of offered software algorithm is in-process UAV flight and performance of task after GNSS signals loss.

4 Conclusion

The paper presents the results of developing an algorithm of software for controlling and navigating an UAV using the artificial intelligence technology. The process of functioning of AI controller of UAV is considered. An algorithm of software has been developed for using AI visual or inertial navigation system to determine the coordinate and time characteristics of UAV in the event of non-functioning of the global

satellite navigation system. In this case, the developed algorithm for the functioning of the AI of the on-board computer itself determines the coordinate-time characteristics of the unmanned aerial vehicle in autonomous flight mode. The proposed algorithm of software and control architecture are applied to an unmanned aerial vehicle intended to perform a reconnaissance flight.

References

1. Y. Zhao, Z. Zheng, Y. Liu, Survey on computational-intelligence based UAV path planning. *Knowl.-Based Syst.* **158**, 54–64 (2018)
2. A.A. Bayramov, E.H. Hashimov, Application SMART for small unmanned aircraft system of systems, in *Handbook of Research on Artificial Intelligence Applications in the Aviation and Aerospace Industries*, eds. T. Shmelova, Y. Sikirda, A. Sterenharz (IGI Global, PA, USA), p. 390, Chapter 8. *Application SMART for Small Unmanned Aircraft System of Systems* (2019)
3. A. Bayramov, E. Hashimov, Y. Nasibov, Unmanned Aerial Vehicle applications for military GIS task solution, vol. 3, Chapter 44, in *Research Anthology on Reliability and Safety in Aviation Systems, Spacecraft, and Air Transport* ed. by D.B.A. Mehdi Khosrow-Pour (IGI-Global Publication, USA, 2021)
4. R. Zhai, Z. Zhou, W. Zhang, S. Sang, P. Li, Control and navigation system for a fixed-wing unmanned aerial vehicle. *Am. Inst. Phys. Adv.* **4** (2014)
5. B. Rahmat, B. Nugroho, Fuzzy and artificial neural networks-based intelligent control systems using python, in *International Seminar of Research Month Science and Technology for People Empowerment* (NST Proceedings, 2018), pp. 152–170
6. M. Sheraz, M. Ahmed, X. Hou, Y. Li, D. Jin, Z. Han, T. Jiang, Artificial intelligence for wireless caching: schemes, performance, and challenges, in *IEEE Communications Surveys & Tutorials*, 1st Quart., (2021) pp. 631–661
7. S. Rezwani, W. Choi, Artificial intelligence approaches for UAV navigation: recent advances and future challenges. *IEEE Access* **10**, 2632–26339 (2022)
8. A.A. Bayramov, S.S. Suleymanov, Artificial intelligence application for unmanned aerial vehicle navigation, in *Modeling, Control and Information Technologies: Proceedings of International Scientific and Practical Conference* (Ministry of Education and Science of Ukraine National University of Water and Environmental Engineering, Rivne, 2023), pp. 21–24

Application Fuzzy PID Controller in UAV



Azad Bayramov , Samir Suleymanov , and Fatali Abdullaev

Abstract A PID controller is one of the most effective methods in automatic control systems. In this paper the operation of PID controller for the developed Unmanned Aerial Vehicle are is analyzed. The PID controller combines proportional, integral and differential components to provide accurate and stable control of the system. It is considered a quadcopter model with only one tilt angle and two left and right motors in two-dimensional space. The Ziegler- Nichols method is used for determining the coefficients P , I and D . The frontal cross-section of a UAV in flight in a state of instability and the angle α of deviation from a stable position are analyzed. The speed given by the PID controller to the motors to bring the platform into equilibrium are obtained. The developed of software algorithm for PID control operation is presented. Also, the offered fuzzy PID controller is presented in the paper. The parameterization using rules and fuzzy membership functions makes it easy to add nonlinearities, logic, and additional input signals to control law. To obtain proportional, integral and derivative control action all together, it is intuitive and convenient to combine PI and PD actions together to form fuzzy PID-type controller.

Keywords PID controller · Unmanned aerial vehicle · Fuzzy method · Software algorithm

1 Introduction

A quadcopter is an unmanned aerial vehicle (UAV) with a simple and compact design and high maneuverability [1]. The high-precision stabilization of this type of platforms in space requires to create of a quality control system. In automatic control systems, the PID controller is one of the most effective control methods. The abbreviation PID stands for proportional–integral–differential. The PID controller ensures the stabilization and maintenance of the set value of the technical system parameters [2–4].

A. Bayramov (✉) · S. Suleymanov · F. Abdullaev
Republican Seismic Survey Center, Azerbaijan National Academy of Sciences, Baku, Azerbaijan
e-mail: azad.bayramov@yahoo.com

In this paper the operation of PID controller for the developed Unmanned Aerial Vehicle are is analyzed. The Ziegler-Nichols method is used for determining the coefficients P for the proportional component, I for the integral component and D for the differential component. The developed of software algorithm for PID control operation and fuzzy PID controller are presented in the paper.

2 PID Controller

The proportional component determines the output signal of the controller in proportion to the difference between the current value of the quantity and the specified value, that is, the greater the difference, the greater the output signal. It allows to react quickly to changes in input data, but when the difference is large, it can cause system variability and instability.

The integral component of PID controller accumulates the control error over time and adjusts the output signal based on this integral. It allows to turn off the static control error and ensure the precise adjustment of the system in the stationary mode. However, in the presence of external influences, the integral component can cause overload and system instability.

The differential component of PID controller sets the output signal of the controller proportional to the rate of change of the value. It allows to quickly react to speed changes and avoid system oscillations when input data changes rapidly. However, the differential component is sensitive to noise and signal variations that can cause instability if not selected properly.

The PID controller combines proportional, integral and differential components to provide accurate and stable control of the system (see Fig. 1).

For simplicity, consider a quadcopter model with only one tilt angle and two left and right motors in two-dimensional space. Let's denote the rotation speeds of the right and left engines by V_{right} , and V_{left} respectively, and the parameter from the control panel to adjust the rotation speed of the engines by Q_{az} .

Figure 2 shows the frontal cross-section of a UAV in flight in a state of instability. α is the angle of deviation from a stable position.

Note, that Q_{az} is the arithmetic average value of the maximum rotation speed expressed as a percentage among the rotation speeds of all engines. Let's assume that the UAV tilts to the left for whatever reason. At this time, the speed given by the

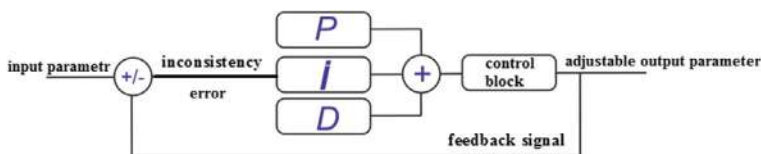
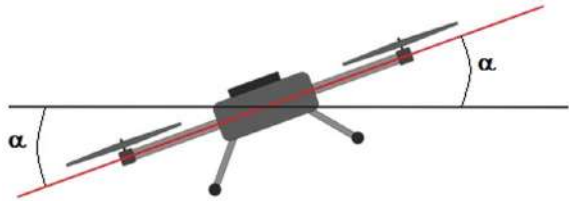


Fig. 1 The chart of PID controller operation

Fig. 2 The frontal cross-section of a UAV in flight in a state of instability



PID controller to the motors to bring the platform into equilibrium is as follows:

$$V_{right} = Q_{az} - F, V_{left} = Q_{az} + F$$

The mismatch factor F produces correspondingly different torques due to the left engine spinning faster than the throttle and the right engine spinning slower. The mismatch coefficient can also take negative values. If we learn to calculate this value in each cycle of the processing cycle, then we can control the quadcopter.

Let's mark the angle at which the quadcopter is at the moment as α_{bank} , and the angle we want to get as α_{target} , and the difference between these two angles as error. The PID controller always tries to minimize the amount of error δ (Fig. 2). The larger the error Δ quantity, the higher the rotation speed of the left motor compared to the right motor. In this case, the discrepancy coefficient can be expressed by the following formula.

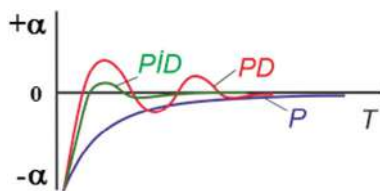
$$F = P \cdot \delta$$

Here, P is the proportionality factor, δ is an error. The higher the value of this coefficient, the faster the response of the quadcopter to return to the required angle. This simple formula means proportional adjustment (blue curve P in Fig. 3). After a few tens of milliseconds (after a few cycles of the processing cycle), the effect of the proportional controller will return the quadcopter to the required (in this case horizontal) position. During all this time, the value of error and F will gradually decrease. Having achieved a certain rate of turn (angular velocity), the UAV will simply roll to the other side, as there is no force acting to keep it in the proper horizontal position (see Fig. 3). It's like a spring that always wants to return to its equilibrium position. For this reason, one more term must be added to the proportional controller, which slows down the quadcopter's rotation and prevents overshoot (rolling in the opposite direction). This term is the rotational speed, and we denote it as $V_{rotation}$. In this case, there is next formula below

$$F = P \cdot \delta + D \cdot V_{rotation}$$

Here, D is the adjustable coefficient and it is called the differentiation coefficient. The greater the value of D , the greater the force that makes the quadcopter stop. It is known that the rate of change of any quantity is the time derivative of this quantity:

Fig. 3 PID controller tries to minimize the amount of error



$$V_{\text{rotation}} = \frac{d\delta}{dt}$$

Then,

$$F = P \cdot \delta + D \cdot \frac{d\delta}{dt}$$

If it is considered the difference between the last value of the quantity error and the previous value is the rate of change of error, then there is below

$$F = P \cdot \delta + D \cdot \frac{-\delta_{\text{before}}}{\Delta t}$$

Here: Δt is an adjustment period, δ before is the before adjustment value.

Let's assume that the inconsistency, that is, the error, is equal to 1 degree. If the adjustment cycle is 0.1 s, the errors will be only 10 degrees. If the tuning chain changes, the sum of the errors will also take different values. In order for the sum of the errors generated in different adjustment periods to be the same, it was necessary to multiply the received total by the adjustment period:

$$\Delta t \sum \delta$$

This formula is the integral of the error in the interval from zero to the present moment. This formula can be written in the following way

$$\int_0^T \delta dt \approx \Delta t \sum \delta$$

Here: T is the time until the present. If substitute this formula into above formula, then the final formula will be like this:

$$F = P \cdot \delta + \dot{I} \cdot \int_0^T \delta dt + D \cdot \frac{d\delta}{dt}$$

Here: $\dot{I}$ is one of the adjusted coefficients, and is called the integration coefficient.

Adjusting the PID controller consists in setting the coefficients P, $\dot{I}$ and D.

There are three different methods for determining the coefficients P, $\dot{I}$ and D. The Ziegler–Nichols method is the most widely used of these methods in practice. The Ziegler–Nichols tuning method is a method of tuning a PID controller. It is performed by setting the I (integral) and D (derivative) gains to 0. The “P” (proportional) gain K_p is then increased (from 0) until it reaches the ultimate gain K_u at which the output of the control loop has stable and consistent oscillations. Ultimate gain K_u and the oscillation period are then used to set the P, $\dot{I}$, and D gains depending on the type of controller used and behavior desired. The algorithm to determine the coefficients using this method have been developed (see Fig. 4).

The algorithm consists of sequentially performing the following operations:

1. Initially all coefficients are reset.
2. With the Q_{az} value from the control panel, the quadcopter is raised to a height of about 2–3 m from the ground.
3. The quadcopter is assigned an α_{target} value.
4. The self-balancing time of the UAV is measured.
5. If the equilibration period is greater than the period t , the value of the proportionality coefficient is gradually increased. The time t takes a different value for different UAVs (0.1 ÷ 0.5 s).
6. At a certain value of the proportionality factor, the quadcopter will receive the α_{target} value. At this time, the ZERO detector will record this period.
7. The quadcopter will oscillate around the α_{target} value, eventually fading. The period of the oscillating movement is calculated as the time between the moments fixed by the ZERO detector (ZERO is a detector subroutine algorithm).

Based on the obtained Pi and T values, we calculate the P, $\dot{I}$, D coefficients using the Ziegler–Nichols method:

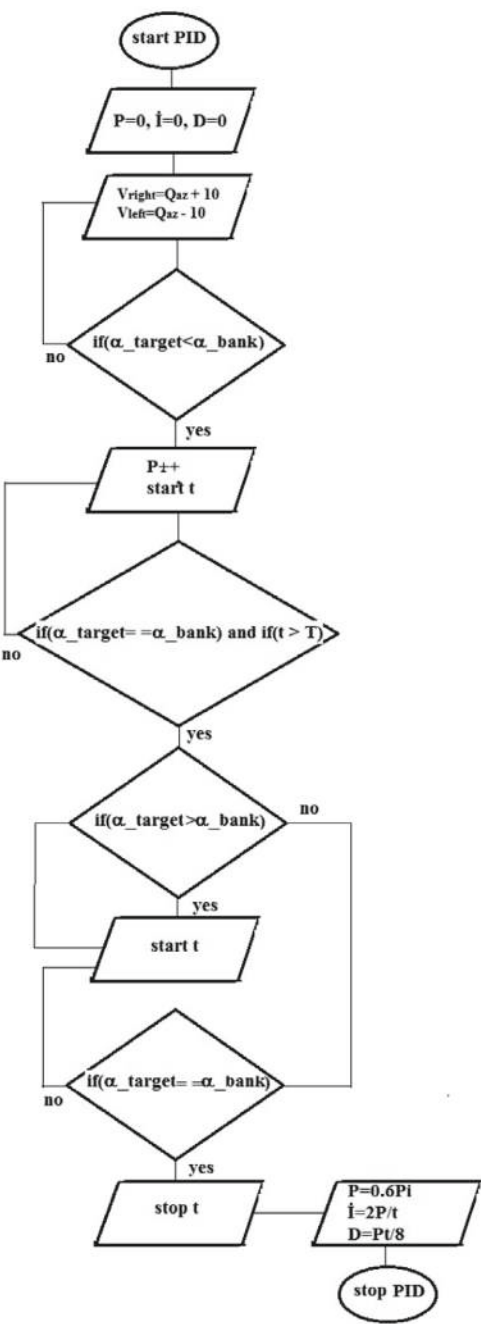
$$P = 0, 6P_i, \dot{I} = 2P/T, D = PT/8$$

Here: The P_i is the proportionality factor of UAV oscillating start moment. T is the period of the oscillating movement. According to these coefficients, the rotation speed of the right and left motors will be as follows

$$V_{\text{right}} = Q_{az} - P \cdot \delta + \dot{I} \cdot \int_0^T \delta \cdot dt + D \cdot \frac{d\delta}{dt}$$

$$V_{\text{right}} = Q_{az} + P \cdot \delta + \dot{I} \cdot \int_0^T \delta \cdot dt + D \cdot \frac{d\delta}{dt}$$

Fig. 4 The developed software algorithm for PID control operation



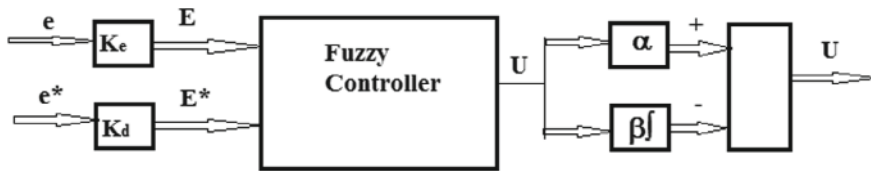


Fig. 5 Fuzzy PID-type controller structure

3 Fuzzy PID Controller

However, the PID controller being linear is not suited for strongly nonlinear systems. Fuzzy control is often mentioned as an alternative to PID control [5, 6]. The parameterization using rules and fuzzy membership functions make it easy to add nonlinearities, logic, and additional input signals to control law. To obtain proportional, integral and derivative control action all together, it is intuitive and convenient to combine PI and PD actions together to form fuzzy.

PID-type controller. Therefore, the formulation of fuzzy PID controller can be achieved by combining fuzzy PI and PD controllers with two distinct rule-bases.

Another and a simpler way of construction a fuzzy PID controller is combining fuzzy PD controller with an integrator and a summation unit at the output (Fig. 5).

This fuzzy PID controller design is presented to control of developed UAV. The fuzzy PID controller, given in Fig. 5, has a simple structure and has only four scaling factors and it is used with a parameter adaptive method that is based on a peak observer. The tuning mechanism adjusts the input and output scaling factors corresponding to the integral and derivative coefficients of the controller.

4 Conclusion

In paper the operation of PID controller for the developed Unmanned Aerial Vehicle are analyzed. It is considered a quadcopter model with only one tilt angle and two left and right motors in two-dimensional space. The Ziegler-Nichols method is used for determining the coefficients P, I and D. The frontal cross-section of a UAV in flight in a state of instability and the angle α of deviation from a stable position are analyzed. The speed formulas given by the PID controller to the motors to bring the platform into equilibrium are obtained. The developed of software algorithm for PID control operation is presented. The offered fuzzy PID controller is presented. To obtain proportional, integral and derivative control action all together, it is intuitive and convenient to combine PI and PD actions together to form fuzzy PID-type controller.

References

1. A.A. Bayramov, E.H. Hashimov, Application SMART for small unmanned aircraft system of systems, in *Handbook of Research on Artificial Intelligence Applications in the Aviation and Aerospace Industries*, ed. by T. Shmelova, Y. Sikirda, A. Sterenharz (IGI Global, PA, USA), p. 390. *Chapter 8, Application SMART for Small Unmanned Aircraft System of Systems* (2019).
2. A.K. Ho, Fundamentals of PID control (PDH Online Course E331 (3 PDH), 2020), p. 20
3. K.J. Astrom, T. Hugglund, PID controllers, 2nd edn (1994). p. 360
4. K.M. Passino, S. Yurkovich, *Fuzzy Control* (Addison-Wesley-Longman, Menlo Park, CA, 1998)
5. G.L. Demidova, D.V. Lukichev, *Controllers Based on Fuzzy Logic in Technical Facility Management Systems* (ITMO University, Sankt-Peterburg, 2017), p. 81
6. D.S. Karpovich, AN. Shumski, V.V. Saroka, Control system of an unmanned aerial vehicle using the theory of fuzzy sets. *Proceeding Belarusian State Technol. Univ.* **6**, 110–116 (2016)

Fuzzy Applications and Analysis System

Data Fusion Is More Complex Than Data Processing: A Proof



Robert Alvarez, Salvador Ruiz, Martine Ceberio, and Vladik Kreinovich

Abstract Empirical data shows that, in general, data fusion takes more computation time than data processing. In this paper, we provide a proof that data fusion is indeed more complex than data processing.

1 Formulation of the Problem

Known empirical fact. Empirical data shows that, in general, data fusion takes more computation time than data processing. In particular, as shown, e.g., in [2], this is true for processing fuzzy data [1, 4, 8, 11, 12, 15].

Challenge. To the best of our knowledge, there has been no theoretical explanation for the above empirical fact. Without such an explanation, we cannot be absolutely sure that this is indeed an universal phenomenon.

What we do in this paper. In this paper, we provide a theoretical explanation for the above empirical phenomenon.

Structure of the papers. To be able to prove our result, we need to gave precise definitions of the corresponding notions. So, in Sect. 2, we recall what is data processing and in Sect. 3, we recall what is data fusion. Finally, in Sect. 4, we provide the desired proof.

R. Alvarez · S. Ruiz · M. Ceberio · V. Kreinovich (✉)
Department of Computer Science, University of Texas at El Paso, El Paso, TX, USA
e-mail: vladik@utep.edu

R. Alvarez
e-mail: rjalvarez@miners.utep.edu

S. Ruiz
e-mail: sruiz13@miners.utep.edu

M. Ceberio
e-mail: mceberio@utep.edu

2 What Is Data Processing: a Brief Reminder

What is data processing. Very often, we are interested in the value of a quantity y which cannot be easily—or not at all—measured directly. For example, today, we cannot directly measure tomorrow’s temperature or the size of the universe.

So what can we do? We cannot measure this quantity y directly, so we can measure it *indirectly*; see, e.g., [13]. What does this mean? It means that:

- we find some auxiliary quantities $x_1, \dots, x_n$ which are related to the quantity y (in which we are interested) by some known relation $y = f(x_1, \dots, x_n)$;
- then, we measure the values $x_1, \dots, x_n$, getting the measurement results $\tilde{x}_1, \dots, \tilde{x}_n$;
- finally, we plug in the measurement results $\tilde{x}_1, \dots, \tilde{x}_n$ into the algorithm f , and produce the estimate $\tilde{y} = f(\tilde{x}_1, \dots, \tilde{x}_n)$ for y .

This is called *data processing*.

Need to take uncertainty into account. In processing data, it is important to take into account that measurements are never absolutely accurate.

For example, if someone has height 170 cm, and you ask the computer to transform it into inches, the computer will give you 13 digits – which doesn’t make any sense. Because, of course, this value is approximate. It’s not really the exact 170.000, it can be 170 plus minus half a centimeter.

Similarly, in many cases, all we know is the upper bound Δ on the absolute value $|\Delta x|$ of the measurement error $\Delta x \stackrel{\text{def}}{=} \tilde{x} - x$: $|\Delta x| \leq \Delta$ (and that’s exactly what happens with height).

In this case, once we get the measurement result $\tilde{x}$, all we know about the actual value x of the measured quantity is that this actual value is somewhere between $\tilde{x} - \Delta$ and $\tilde{x} + \Delta$.

If, for example, the measured value is 1.0, and the upper bound on the measurement error is 0.1, this means the actual value is somewhere in the interval $[0.9, 1.1]$.

In this case, since we don’t know the exact values of x_i , we cannot compute the exact value of y . For different possible values of x_i , we get, in general, different values of y . It is desirable to find the range of all possible values of y for all possible $x_1, \dots, x_n$, i.e., the range

$$[\underline{y}, \overline{y}] = \{f(x_1, \dots, x_n) : x_1 \in [\underline{x}_1, \overline{x}_1], \dots, x_n \in [\underline{x}_n, \overline{x}_n]\}.$$

The problem of computing this range is known as the problem of *interval computation*; see, e.g., [3, 6, 7, 9].

What is the computational complexity of interval computations. In general, the interval computation problem is NP-hard; see, e.g., [5].

However, there are cases when computation is feasible. One of such cases is the case of the so-called *single-use expression* (SUE, for short). A single use expression is an expression in which each variable occurs only once. It can be a complicated expression, such as $x_1 + x_2^{x_3}$, but as long as each variable occurs only once, computing its range is feasible.

3 What Is Data Fusion: a Brief Reminder

What did we overlook when we described data processing. As we have mentioned earlier, to make desired estimates—such as predictions—we measure the corresponding quantities $x_1, \dots, x_n$.

What we did not take into account in the above description is that usually, the quantities x_i are not independent, they are related to each other.

For example, if we measure the temperature in two locations in El Paso, we know that these two temperature cannot differ too much: it's not possible to have 100°F in one location and 80°F in a nearby location.

Taking this dependence into account can help us get more accurate estimates.

Example. For example, suppose that we have two quantities:

- we have the quantity x_1 about which, based on the corresponding measurement result, we know that $x_1 \in [0.9, 1.1]$, and
- we have the quantity x_2 about which, based on the corresponding measurement result, we know that $x_2 \in [0.8, 1.0]$.

Suppose also that we know that the difference $|x_1 - x_2|$ between these two values cannot be larger than 0.01: $|x_1 - x_2| \leq 0.01$. What can we then conclude about the actual values of these two quantities?

- We know that x_1 is from 0.9 to 1.1.
- We also know that x_2 is from 0.8 to 1.0, and that x_2 cannot differ from x_1 by more than 0.01.

So, x_2 cannot be larger than 1.0. Since x_1 can differ from x_2 by no more than 0.01, we can conclude that x_1 cannot be larger than $1.0 + 0.01 = 1.01$.

So, by using information about x_2 , we get a narrower interval for x_1 : the interval $[0.9, 1.01]$.

So what is data fusion. This was an example of what is called data fusion: when we use measurements of related quantities to get we get a better estimate of the given quantity.

4 Our Proof

Problem: reminder. Empirically, data fusion is more time consuming the data processing.

What we do in this section. In this section, we will provide a proof that data fusion is indeed more complex.

How exactly we will prove it. Specifically, we will prove that even if:

- the data processing algorithm $f(x_1, \dots, x_n)$ is described by a single-use expression, and
- all the constraints $g(x_1, \dots, x_n) = a$ and $g(x_1, \dots, x_n) \geq a$ relating the measured quantities are described by single-use expressions $g(x_1, \dots, x_n)$,

data fusion is still NP-hard.

Why this proves that data fusion is more complex than data processing. For data processing, as we have mentioned earlier, if we only have single-use expressions, the problem is feasible. Thus, our result proves that data fusion is indeed more complex.

Main idea behind the proof. To prove that a problem is NP-hard, it is sufficient to prove that some subproblem of this general problem is NP-hard. In line with this idea, we will show that a certain subproblem of the above problem—of computing the range of a SUE expression under SUE constraints—is NP-hard.

Let us start describing the subproblem: what is the data processing algorithm. Let us consider the situation when we have $2n$ auxiliary quantities $x_1, \dots, x_n, x_{n+1}, \dots, x_{2n}$ which are related to the desired quantity y by the following single-use expression

$$y = \frac{1}{n} \cdot \sum_{i=1}^n x_i^2 - \left(\frac{1}{n} \sum_{i=1}^n x_{n+i} \right)^2.$$

It is a single-use expression, because here every variable occurs exactly once.

The corresponding data processing problem is feasible. Since we have a single-use expression, performing data processing would be very straightforward.

Describing the constraints. Let us now assume that the relation between the quantities x_i are described by the SUE formulas $x_i = x_{n+i}$.

The corresponding data fusion problem is NP-hard. Let us show that the resulting data fusion problem is NP-hard.

Indeed, to compute the range of possible values of y under this constraint, we can simply plug in x_i instead of x_{n+i} into the original expression. As a result, we get the following expression:

$$\frac{1}{n} \cdot \sum_{i=1}^n x_i^2 - \left(\frac{1}{n} \sum_{i=1}^n x_i \right)^2.$$

This expression is well known in statistics—it is the sample variance; see, e.g., [14].

And it is known that the problem of computing the range of sample variance under interval uncertainty is actually NP-hard; see, e.g., [10].

Conclusion. So we have the case, when:

- for data processing, we have a feasible algorithm, while
- if we add constraints, even such simple constraints as equalities, we immediately get an NP-hard problem.

This proves theoretically that data fusion is indeed more complex than data processing.

Acknowledgements This work was supported in part by the National Science Foundation grants 1623190 (A Model of Change for Preparing a New Generation for Professional Practice in Computer Science), HRD-1834620 and HRD-2034030 (CAHSI Includes), EAR-2225395 (Center for Collective Impact in Earthquake Science C-CIES), and by the AT&T Fellowship in Information Technology.

It was also supported by a grant from the Hungarian National Research, Development and Innovation Office (NRDI).

References

1. R. Belohlavek, J.W. Dauben, G.J. Klir, *Fuzzy Logic and Mathematics: A Historical Perspective* (Oxford University Press, New York, 2017)
2. D. Dubois, H. Prade, and R. R. Yager, “Computation of intelligent fusion operations based on constraint fuzzy arithmetic”, *Proceedings of the 1998 IEEE International Conference on Fuzzy Systems FUZZ-IEEE’98*, Anchorage, Alaska, May 4–9, 1998, Vol. 1, pp. 767–772
3. L. Jaulin, M. Kiefer, O. Didrit, E. Walter, *Applied Interval Analysis, with Examples in Parameter and State Estimation, Robust Control, and Robotics* (Springer, London, 2012)
4. G. Klir, B. Yuan, *Fuzzy Sets and Fuzzy Logic* (Prentice Hall, Upper Saddle River, New Jersey, 1995)
5. V. Kreinovich, A. Lakeyev, J. Rohn, P. Kahl, *Computational Complexity and Feasibility of Data Processing And Interval Computations* (Kluwer, Dordrecht, 1998)
6. B.J. Kubica, *Interval Methods for Solving Nonlinear Constraint Satisfaction, Optimization, and Similar Problems: from Inequalities Systems to Game Solutions* (Springer, Cham, Switzerland, 2019)
7. G. Mayer, *Interval Analysis and Automatic Result Verification* (de Gruyter, Berlin, 2017)
8. J.M. Mendel, *Explainable Uncertain Rule-Based Fuzzy Systems* (Springer, Cham, Switzerland, 2024)
9. R.E. Moore, R.B. Kearfott, M.J. Cloud, *Introduction to Interval Analysis* (SIAM, Philadelphia, 2009)
10. H.T. Nguyen, V. Kreinovich, B. Wu, G. Xiang, *Computing Statistics under Interval and Fuzzy Uncertainty* (Springer Verlag, Berlin, Heidelberg, 2012)

11. H.T. Nguyen, C.L. Walker, E.A. Walker, *A First Course in Fuzzy Logic* (Chapman and Hall/CRC, Boca Raton, Florida, 2019)
12. V. Novák, I. Perfilieva, J. Močkoř, *Mathematical Principles of Fuzzy Logic* (Kluwer, Boston, Dordrecht, 1999)
13. S.G. Rabinovich, *Measurement Errors and Uncertainty: Theory and Practice* (Springer Verlag, New York, 2005)
14. D.J. Sheskin, *Handbook of Parametric and Nonparametric Statistical Procedures* (Chapman and Hall/CRC, Boca Raton, Florida, 2011)
15. L.A. Zadeh, Fuzzy sets. Inf. Control. **8**, 338–353 (1965)

A Comparative Analysis of AI-Based Cryptocurrencies: Fundamental and Technical Perspectives



Hilala Jafarova 

Abstract The intersection of artificial intelligence (AI) and blockchain technology has led to the emergence of AI-based cryptocurrencies that aim to revolutionize various sectors, from finance to data sharing. This paper presents a comparative analysis of seven prominent AI-based cryptocurrencies: SingularityNET (AGIX), Fetch.AI (FET), Ocean Protocol (OCEAN), Numerai (NMR), Cortex (CTXC), Deep-Brain Chain (DBC), and SingularityDAO (SDAO). Both fundamental and technical analyses are conducted to assess each project's goals, strengths, challenges, market presence, and price behavior over the past five years. The findings illustrate how AI integration with decentralized platforms holds transformative potential but also faces significant hurdles, such as adoption, technical complexity, and market volatility.

Keywords Artificial intelligence (AI) · Blockchain · AI-based cryptocurrencies · AGIX · FET · OCEAN · NMR · CTXC · DBC · SDAO) · ML · Time-series forecasting · Market sentiment analysis · Decentralized finance (DeFi)

1 Introduction

AI and blockchain technologies are revolutionizing a wide array of sectors. AI-based cryptocurrencies leverage the computational power of AI to offer decentralized services, enhance data sharing, and automate financial systems. These coins combine blockchain technology with artificial intelligence to offer innovative solutions in various sectors. Such technologies enhance trade, optimize supply chains, improve data security, and facilitate financial inclusion, among other benefits. This paper focuses on the comparative analysis of key AI-based cryptocurrencies and aims to develop a forecasting model that projects future market trends. The objective is to provide insights into their growth potential and assess the challenges associated with price forecasting in a speculative market.

H. Jafarova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, Murtuza Muxtarov Street, 194, Baku, Azerbaijan
e-mail: hilala-jafarova@unec.edu.az

In recent years, the convergence of artificial intelligence (AI) and blockchain technology has spurred the development of AI-based cryptocurrencies. These digital assets leverage AI algorithms to enhance the capabilities of decentralized networks, creating new applications in fields such as finance, data exchange, and autonomous systems. AI's ability to analyze large datasets and make predictions aligns seamlessly with the decentralized, transparent, and immutable nature of blockchain, making it a powerful combination for solving complex problems across various industries [1, 2].

AI-based cryptocurrencies differ from traditional cryptocurrencies by incorporating AI technologies to drive specific functionalities within their ecosystems. For instance, they enable the automation of financial processes, optimize data-sharing systems, and facilitate autonomous decision-making in decentralized applications [3]. Some well-known AI-based cryptocurrencies like SingularityNET (AGIX) focus on creating decentralized marketplaces for AI services, while others like Fetch.AI (FET) aim to enable decentralized autonomous agents that perform real-world tasks. These projects aim to democratize AI access and reduce dependency on centralized platforms such as Google and Amazon [4]. AI and blockchain are increasingly converging to address the challenges of secure and transparent data management. Chen et al. [5] review the integration of blockchain with AI, emphasizing its role in fostering secure AI applications. Schatsky et al. [6] further highlight how AI applications can be enabled by blockchain, creating new opportunities for innovation across sectors.

Blockchain and AI integration have been widely discussed in scientific literature due to their potential to disrupt several industries. According to Mougayar [7], blockchain provides a secure and distributed infrastructure, while AI contributes data processing, learning, and decision-making capabilities. The combination creates an ecosystem where blockchain ensures transparency, data integrity, and decentralized control, while AI optimizes decision-making processes [8]. Zhang et al. [9] expand on this idea by applying blockchain to smart cities, offering a secure and decentralized approach to managing interconnected devices in urban environments. Blockchain's decentralized nature ensures that IoT systems are more resistant to cyberattacks and data breaches. Research suggests that AI models can be efficiently trained using decentralized computing power, with projects like DeepBrain Chain (DBC) providing blockchain-based infrastructures for lowering the cost of AI training by decentralizing computation [10].

Additionally, AI-based cryptocurrencies address critical issues such as data privacy and security, which are pivotal concerns in the AI industry [9]. For instance, Ocean Protocol (OCEAN) provides a decentralized data exchange platform where users can securely share and monetize their data without relinquishing control over it. This solution is particularly relevant in today's data-driven economy, where companies increasingly require high-quality, secure datasets to train AI models [11].

These projects often face significant hurdles in real-world integration, requiring a large ecosystem of users, developers, and companies to achieve scalability. Moreover, the speculative nature of the cryptocurrency market introduces high volatility, which can affect the long-term growth prospects of these tokens [12]. Blockchain technology has evolved significantly since its inception, offering a range of opportunities across different sectors. As noted by Yli-Huumo et al. [13], blockchain research has expanded from its initial focus on cryptocurrencies to a wide variety of applications. Casino et al. [14] provide a systematic review of blockchain-based applications, categorizing its uses in finance, healthcare, supply chain, and beyond, while highlighting the open issues that need further investigation.

In finance, decentralized finance (DeFi) has emerged as a key area where blockchain has created innovative systems. Gai et al. [15] present an overview of FinTech, which integrates blockchain to enhance security, transparency, and efficiency in financial transactions. Similarly, Catalini and Gans [16] discuss the economics of blockchain, noting that by removing intermediaries, blockchain reduces transaction costs and fosters trust in peer-to-peer systems.

Beyond finance, blockchain is being integrated into various technological infrastructures. Hardjono et al. [11] explore the role of blockchain in data-sharing networks, proposing technical infrastructures that can facilitate more secure and decentralized sharing of data. In the healthcare industry, Jiang et al. [17] introduce BloCHIE, a blockchain-based platform for healthcare information exchange, which enhances patient data security, privacy, and accessibility. This use case illustrates the broader trend of blockchain's application in sensitive data environments, a theme further examined by Liang et al. [18] with ProvChain, which provides data provenance in cloud environments to ensure privacy and data integrity.

Despite its numerous applications, blockchain still faces technical and operational challenges. Vukolić [19] contrasts different consensus protocols, such as proof-of-work and Byzantine fault tolerance, which are critical for ensuring security and scalability in blockchain systems. Zheng et al. [20] provide an overview of blockchain architecture, consensus mechanisms, and future trends, identifying key areas for improvement as blockchain systems continue to evolve.

As blockchain matures, its role as a disruptive technology becomes clearer. Frizzo-Barker et al. [21] conduct a systematic review of blockchain's impact on business, showing how firms are rethinking operations to incorporate distributed ledger technologies. Hughes et al. [22] also discuss how blockchain is reshaping business models, highlighting the implications of distributed ledger technologies beyond Bitcoin.

Moreover, blockchain's promise for fostering sustainable practices has been explored by Saberi et al. [23], who examine its potential in sustainable supply chain management. Blockchain can track and verify supply chains, ensuring transparency, reducing fraud, and promoting environmental sustainability.

Lacity and van Hoek [24] provide a comprehensive summary of what has been learned about blockchain in business thus far, suggesting that while blockchain presents immense potential, the journey is still in its early stages system. Blockchain technology presents both opportunities and critical challenges for modern computing systems [25]. Decentralized finance aims to establish an open and transparent financial ecosystem [26]. Blockchain systems have been analyzed from a data processing perspective in previous studies [27]. Trust remains a major concern in cryptocurrency ecosystems [28]. Various consensus protocols have been reviewed for blockchain applications [29]. Swan provided one of the earliest comprehensive discussions on blockchain-based economic systems [30]. Blockchain applications in healthcare demonstrate significant potential [31]. Numerical system properties have also been investigated in recent mathematical studies [32].

2 Overview of Selected AI-Based Cryptocurrencies

This article explores the landscape of decentralized AI platforms, highlighting the goals, strengths, and challenges associated with seven prominent projects in the field.

SingularityNET (AGIX) is pioneering a decentralized marketplace for AI services, leveraging strategic partnerships, such as with Hanson Robotics, to drive monetization of AI technologies. Despite these advantages, the platform faces hurdles related to its technical complexity and the slow pace of adoption, which are crucial for its success.

Fetch.AI (FET) aims to revolutionize real-world applications through decentralized autonomous agents. Its solutions are applicable across various domains, including smart cities, energy grids, and supply chains. However, the platform's long-term success hinges on its ability to integrate effectively with real-world industries.

Ocean Protocol (OCEAN) is focused on creating a decentralized data exchange protocol, allowing users to securely share and monetize their data for AI development. The platform's effectiveness is dependent on the establishment of a robust ecosystem involving data providers and AI developers.

Numerai (NMR) operates as a decentralized hedge fund, utilizing crowd-sourced predictive models powered by AI to forecast financial markets. While its unique model attracts data scientists and finance professionals, the platform's niche focus limits its broader market appeal.

Cortex (CTXC) enables the integration of AI models into smart contracts and decentralized applications (dApps). Despite its cutting-edge approach, Cortex encounters challenges related to early-stage adoption and competition from established centralized AI applications.

DeepBrain Chain (DBC) offers a decentralized network for AI training, aiming to lower costs by providing computing power to developers and researchers. The platform faces competition from major cloud service providers like AWS and Google Cloud, which dominate the market.

SingularityDAO (SDAO) combines AI with decentralized finance (DeFi) to manage cryptocurrency portfolios using dynamic algorithms known as DynaSets. While the platform benefits from AI-driven optimization, it contends with the inherent volatility and speculative nature of DeFi projects.

3 Fundamental and Technical Analysis

The Each project presents distinct advantages and challenges that influence their potential for adoption and impact within the broader AI and blockchain ecosystems. The fundamental analysis of AI-based cryptocurrencies involves examining their use cases, adoption potential, and project challenges. Below is a summary of the fundamental strengths and weaknesses of each:

Cryptocurrency	Use case diversity	Adoption potential	Innovation level	Challenges
SingularityNET	High	Moderate	High	Slow adoption of decentralized AI
Fetch.AI	High	Moderate	High	Complex integration in real-world systems
Ocean protocol	Moderate	High	Moderate	Requires robust data provider network
Numerai	Low	Low	High	Limited to finance professionals
Cortex	Moderate	Moderate	High	Early adoption of AI smart contracts
DeepBrain chain	Moderate	Moderate	Moderate	Competition from centralized services
SingularityDAO	High	Moderate	High	Highly speculative DeFi space

Technical analysis provides insights into the price behavior of these tokens, their volatility, and market trends over time. The price volatility of each cryptocurrency is indicative of market sentiment and external factors such as AI industry growth or partnerships.

All seven cryptocurrencies exhibit high price volatility due to speculative trading, tokenomics models, and varying degrees of adoption. The tokens are heavily influenced by AI-related developments, partnerships, and news cycles. Based on moving averages and RSI technical indicators, the following can be stated:

- **SingularityNET (AGIX):** Long-term uptrend with regular short-term corrections.
- **Fetch.AI (FET):** Strong bullish momentum followed by periodic corrections.
- **Ocean Protocol (OCEAN):** Steady uptrend supported by growth in data exchange applications.
- **Numerai (NMR):** Short-term volatility driven by AI competitions and financial predictions.
- **Cortex (CTXC):** Low volatility with slow adoption of AI-powered smart contracts.
- **DeepBrain Chain (DBC):** Moderate price fluctuations linked to updates in AI computing capabilities.
- **SingularityDAO (SDAO):** High price sensitivity to DeFi market developments.

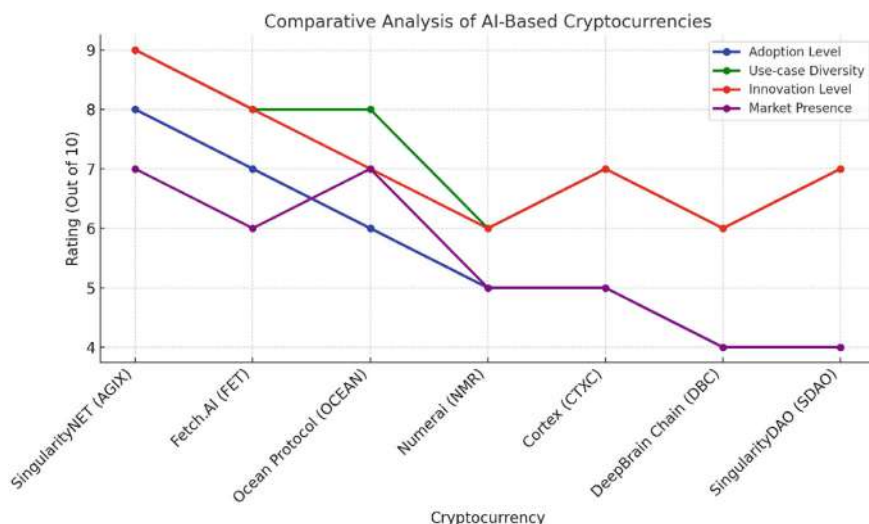
Support and resistance levels for these cryptocurrencies are often determined by project milestones, such as new partnerships, product releases, or successful AI model integration. Investor sentiment also plays a significant role in creating price barriers.

4 Graphical Representation of Price Movements

This section includes a visual analysis of each cryptocurrency's price movements over the past five years on a monthly interval. Graph 1 illustrates price volatility, major price movements, and their corresponding market factors.

5 Discuss and Conclusion

AI-based cryptocurrencies hold the potential to transform sectors like finance, data sharing, autonomous systems, and decentralized applications. However, their success depends on broader adoption, technological advancements, and overcoming challenges such as competition from centralized services. The combination of AI with blockchain technologies presents a promising but still speculative field, with several hurdles to cross in terms of scalability and real-world implementation.



Graph 1 Comparative graph for the AI-based cryptocurrencies: AGIX, FET, OCEAN, NMR, CTXC, DBC, SDAO

Based on Adoption and Diversity of Uses, SingularityNET and Fetch.AI have the widest diversity of use cases, but their adoption is limited by technical complexity and competition.

In terms of market volatility, cryptocurrencies such as Numerai and SingularityDAO are highly sensitive to external factors, making them speculative investments.

Due to Technical and Fundamental Challenges Most projects face significant challenges in adopting centralized platforms and coping with competition. This limits their market presence despite strong levels of innovation.

AI-based cryptocurrencies represent a rapidly evolving intersection of two cutting-edge technologies: artificial intelligence and blockchain. While these tokens offer significant innovation and potential, they remain speculative investments due to the technical challenges and volatility inherent in both fields. Continued development, partnerships, and industry growth will determine the long-term success of these projects. Future research should focus on the integration of AI applications with real-world systems and the scalability of decentralized AI ecosystems.

References

1. S. Nakamoto, *Bitcoin: A Peer-to-Peer Electronic Cash System* (2008). <https://bitcoin.org/bitcoin.pdf>
2. D. Tapscott, A. Tapscott, *Blockchain Revolution: How the Technology Behind Bitcoin is Changing Money, Business, and the World* (Penguin, 2016)

3. Q.K. Nguyen, M. Ding, P.N. Pathirana, A. Seneviratne, Blockchain and AI: a new frontier for secure learning and data management. *IEEE Commun. Surv. Tutor.* **22**(3), 1977–2007 (2020)
4. J. Benet, *IPFS—Content Addressed, Versioned, P2P File System* (2014). [arXiv:1407.3561](https://arxiv.org/abs/1407.3561)
5. W. Chen, W. Xu, Z. Yin, Blockchain for AI: a review. *IEEE Access* **8**, 82538–82555 (2020)
6. D. Schatsky, C. Muraskin, R. Gurumurthy, in *AI Meets Blockchain: AI Applications Enabled by Blockchain* (Deloitte Insights, 2019)
7. W. Mougayar, *The Business Blockchain: Promise, Practice, and the Application of the Next Internet Technology* (Wiley, 2016)
8. M.A. Khan, K. Salah, IoT security: review, blockchain solutions, and open challenges. *Futur. Gener. Comput. Syst.* **82**, 395–411 (2018)
9. L. Zhang, K. Gai, M. Qiu, Y. Zhu, L. Tao, Blockchain-based secure and trustworthy internet of things in SDN-enabled smart cities. *IEEE Internet Things J.* **7**(5), 4169–4181 (2020)
10. J. Xu, S. Wang, M. Zhu, J. Shen, Blockchain-enabled decentralized storage in fog computing: a case study of DeepBrain Chain, in *2018 IEEE Conference on Internet of Things* (2018)
11. T. Hardjono, D. Shrier, A. Pentland, in *Data Sharing Networks: Technical Infrastructure Options*. MIT Connection Science (2019)
12. M. Risius, K. Spohrer, A blockchain research framework: what we (don't) know, where we go from here, and how we will get there. *Bus. Inf. Syst. Eng.* **59**(6), 385–409 (2017)
13. J. Yli-Huoma, D. Ko, S. Choi, S. Park, K. Smolander, Where is current research on blockchain technology?—a systematic review. *PLoS ONE* **11**(10), e0163477 (2016)
14. F. Casino, T.K. Dasaklis, C. Patsakis, A systematic literature review of blockchain-based applications: current status, classification, and open issues. *Telemat. Inform.* **36**, 55–81 (2019)
15. K. Gai, M. Qiu, X. Sun, A survey on FinTech. *J. Netw. Comput. Appl.* **103**, 262–273 (2018)
16. C. Catalini, J.S. Gans, Some simple economics of the blockchain. *Commun. ACM* **63**(7), 80–90 (2020)
17. S. Jiang, J. Cao, H. Wu, Y. Yang, J. He, BloCHIE: a blockchain-based platform for healthcare information exchange. *IEEE Access* **8**, 5919–5930 (2021)
18. X. Liang, S. Shetty, D. Tosh, C.A. Kamhoua, K. Kwiat, L. Njilla, ProvChain: a blockchain-based data provenance architecture in cloud environment with enhanced privacy and availability, in *2017 17th IEEE/ACM International Symposium on Cluster, Cloud, and Grid Computing (CCGRID)* (2017), pp. 468–477
19. M. Vukolić, The quest for scalable blockchain fabric: proof-of-work vs. BFT replication, in *International Workshop on Open Problems in Network Security* (Springer, Cham, 2017), pp. 112–125
20. Z. Zheng, S. Xie, H. Dai, X. Chen, H. Wang, An overview of blockchain technology: architecture, consensus, and future trends. *IEEE Int. Congr. Big Data (BigData Congr.)* **2017**, 557–564 (2017)
21. J. Frizzo-Barker, P.A. Chow-White, P.R. Adams, J. Mentanko, Blockchain as a disruptive technology for business: a systematic review. *Int. J. Inf. Manage.* **51**, 102029 (2020)
22. A. Hughes, A. Park, J. Kietzmann, C. Archer-Brown, Beyond bitcoin: what blockchain and distributed ledger technologies mean for firms. *Bus. Horiz.* **62**(3), 273–281 (2019)
23. S. Saberi, M. Kouhizadeh, J. Sarkis, L. Shen, Blockchain technology and its relationships to sustainable supply chain management. *Int. J. Prod. Res.* **57**(7), 2117–2135 (2019)
24. M.C. Lacity, R. Van Hoek, What we've learned so far about blockchain for business. *MIT Sloan Manage. Rev.* **62**(2), 48–54 (2021)
25. E. Bertino, A. Kundu, Blockchain: a primer on the opportunities and challenges. *Computer* **52**(3), 14–17 (2019)
26. Y. Chen, C. Bellavitis, Decentralized finance: blockchain technology and the quest for an open financial system. *J. Bus. Ventur. Insights* **13**, e00183 (2020)
27. T.T.A. Dinh, R. Liu, M. Zhang, G. Chen, B.C. Ooi, J. Wang, Untangling blockchain: a data processing view of blockchain systems. *IEEE Trans. Knowl. Data Eng.* **30**(7), 1366–1385 (2018)
28. M.H.U. Rehman, K. Salah, E. Damiani, D. Svetinovic, Trust in blockchain cryptocurrency ecosystem. *IEEE Trans. Eng. Manage.* **67**(4), 1196–1212 (2020)

29. Sankar, L.S., Sindhu, M., Sethumadhavan, M., Survey of consensus protocols on blockchain applications, in *2017 4th International Conference on Advanced Computing and Communication Systems (ICACCS)* (2017), pp. 1–5
30. M. Swan, in *Blockchain: Blueprint for a New Economy* (O'Reilly Media, 2015)
31. P. Zhang, D.C. Schmidt, J. White, G. Lenz, Blockchain technology use cases in healthcare. *Adv. Comput.* **111**, 1–41 (2018)
32. Y. Aliyev, Digits of powers of 2 in ternary numerical systems. *Notes Number Theory Discret. Math.* **29**, 474–485 (2023)

Application and Enhancement of New Technologies for Cyber Risk Prediction



Shahnaz N. Shahbazova , Semral Hajiyevev, and Hilala Jafarova 

Abstract The rapid development of cyberspace increases security risks, requiring advanced predictive models and technologies. This study analyzes the application and improvement of AI, machine learning, big data analytics, and behavioral analysis for cyber risk prediction. AI-driven models anticipate cyberattack dynamics, while machine learning detects vulnerabilities early, minimizing risks. Big data analytics enhances real-time security strategies, and behavioral analysis identifies anomalies for early threat detection. Implementing these technologies strengthens cybersecurity and improves future threat prediction. Their integration forms the foundation of effective cyber risk management for organizations and individuals.

Keywords Cyber risks · Prediction · Artificial intelligence · Machine learning · Big data analytics · Behavioral analysis · Security technologies · Cybersecurity

1 Introduction

Cybersecurity has become one of the most important issues in the modern era of rapid development of the digital world. The increasing role of Internet technologies and digital devices, the expansion of the volume of information exchange and the global accessibility of communication tools have led to the emergence of new security problems. Cyber risks are multifaceted threats that threaten the confidentiality, integrity and availability of individual and organizational information. Problems such as data leakage, malware loading into systems, financial data theft and data

S. N. Shahbazova · S. Hajiyevev · H. Jafarova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, UNEC, Baku, Azerbaijan
e-mail: hilala-jafarova@unec.edu.az

S. N. Shahbazova

e-mail: shahnaz_shahbazova@unec.edu.az

S. Hajiyevev

International Center for Master's and Doctoral Studies, Azerbaijan State University of Economics (UNEC), Baku, Azerbaijan

© The Author(s), under exclusive license to Springer Nature Switzerland AG 2026

S. N. Shahbazova et al. (eds.), *Soft Computing Applications*, Lecture Notes in Networks and Systems 2048, https://doi.org/10.1007/978-3-032-04803-5_21

177

manipulation lead to the weakening of the security structure of cyberspace. In such an environment, cyber risk prediction and risk management are considered one of the main tasks in the field of cybersecurity. Cyber risk prediction is an important step in identifying potential risks before attacks occur and minimizing these risks. The process of predicting cyber threats is not only limited to identifying existing vulnerabilities, but also involves developing strategies aimed at reducing the impact of risks by evaluating possible attack scenarios. In this regard, it is necessary to introduce and improve new technologies to make the prediction process more effective. In recent years, artificial intelligence, machine learning, big data analytics, and behavioral analysis technologies have provided important opportunities in the field of cyberthreat prediction. Artificial intelligence technologies analyze big data in real time, identify anomalies, and warn about potential attacks. Machine learning algorithms enable automated detection and classification of risks, which allows for more flexible defense measures against cyberattacks. Big data analytics, on the other hand, analyzes large amounts of information collected from the cyber environment and allows for the prediction of the motives, directions, and potential impact of attacks in advance. In addition, behavioral analysis technologies study the behavioral patterns of users and systems and immediately identify abnormal activities. This technology is used to study repeated patterns of user and system behavior, and any activity that does not match these patterns is evaluated as a threat. This feature ensures the early detection and action of cyber risks.

The main objective of the research is to investigate the application of artificial intelligence, machine learning and other technologies for predicting cyber risks and the effectiveness of these technologies in the field of security. To achieve this goal, the following tasks have been identified:

- Conduct an in-depth analysis of cyberspace security problems and classify various types of cyber risks.
- Study the impact of artificial intelligence and machine learning technologies on the process of predicting cyber risks.
- Investigate innovative approaches to predicting risks through behavioral analysis and big data analytics.
- Develop proposals for the application of existing technologies in the field of predicting and managing cyber risks.

The research will use international scientific research in the field of cybersecurity, practice reports, cutting-edge technology journals and conference proceedings in the field of artificial intelligence and machine learning, as well as existing research on big data technologies as the main information base. These include information on cyber risk analysis provided by organizations.

A study by Abdalla and Arshad [1] highlights the information security issues in the adoption of cybersecurity standards in state-owned companies in Malaysia. The authors emphasize that standards are a key factor in ensuring cybersecurity, considering the need for systematic application of technologies to predict and prevent cyber risks. This is an important study in understanding how new technologies can be applied to increase the resilience of cyberspace.

Abeysekera and Kumarawadu [2] provide an analysis of the application of blockchain technology in the financial sector, showing how this innovative technology creates an advantage in secure information exchange and cyber risk prediction. Blockchain technology provides a secure flow of information and acts as a stronger tool in data protection, revealing the impact of the application of new technologies in preventing cyber risks.

Algarni et al. [3] show that it is possible to improve the prediction process through quantitative assessment of cyber risks to prevent data leaks in the business environment. This study highlights the importance of quantitative approaches in cyber security, highlighting the effectiveness of using big data in analyzing risks related to data breaches.

Alibeigi and Munir [4] discuss the issues of personal data protection and cyber risk prevention in their study on the implementation of the Malaysian Personal Data Protection Act. This study highlights the role of legislative mechanisms in predicting and preventing cyber risks and the need for a legal framework to protect information security.

Aljumah and Ahanger [5] provide extensive information on cybersecurity threats, challenges, and defense mechanisms in cloud computing. Their study reveals the potential of applying artificial intelligence and machine learning to predict and prevent cyber risks in the cloud environment and analyzes the effectiveness of various defense mechanisms.

Ansari et al. [6] highlight the importance of AI-based cybersecurity training in preventing phishing attacks. This study examines the impact of cybersecurity training on the forecasting process, emphasizing the importance of human factors in predicting cyber risks.

Burrell [8] examines the intricacies of managing cybersecurity teams in a remote work environment caused by the COVID-19 pandemic. Burrell's work highlights the importance of collaboration and digital learning technologies in predicting and managing cyber risks in a telework environment.

Campbell et al. (2010) examine the security challenges and lessons learned from autonomous driving systems in urban environments. This study contributes to understanding security aspects in the cyber environment by analyzing the development of technologies applied to predict the security risks of autonomous systems.

Canelon et al. [10] examine the impact of technological strategies to enhance the security of cyberspace by developing a cybersecurity control framework for blockchain ecosystems. This study provides practical recommendations on the potential use of blockchain technology in ensuring security.

Cui et al. [11] propose a framework for safety and security for autonomous vehicles. Their work focuses on studying the impact of collaborative technologies on cyber risk prediction and management.

Balisane et al. [7] study proposes a robust cybersecurity framework to counter evolving cyber threats by integrating artificial intelligence (AI), machine learning (ML), and blockchain for enhanced threat detection, response, and recovery.

Elbahri et al. [17] study developed a comprehensive framework to enhance cybersecurity and improve vulnerability management in digital environments, particularly in cyber-physical power systems (CPPS). By evaluating existing frameworks (NIST, ISO 27001, ISA/IEC 62443) and identifying key vulnerabilities in CPPS, the study highlighted gaps such as the lack of standardized AI integration and real-time threat detection.

Pfleeger and Caputo [18] paper highlights the importance of integrating behavioral science into cybersecurity products and processes. Two examples show how behavioral science can enhance cybersecurity: one demonstrates clear improvements, and the other shows significant potential for increased effectiveness.

Wu et al. [19] study explores emerging technologies developed to combat the COVID-19 pandemic and the challenges related to their design, development, and use.

Yousaf et al. [20] research enhances the FMECA-ATT&CK framework by integrating AI-related risk assessment using the MITRE ATLAS framework, enabling a continuous evaluation of cyber risks in systems like autonomous engine monitoring and control.

2 Material and Method

This study analyzed the application and effectiveness of new technologies, especially artificial intelligence and machine learning technologies, for predicting cyber risks. The methodology combines various technological approaches and tools in accordance with the purpose of the study. The research materials were mainly collected from existing databases and technological resources in the field of cybersecurity and artificial intelligence. This includes international scientific articles, examples of modern cyber attacks, technological platforms and behavioral analysis algorithms. The research methods cover innovations related to the application of the latest technologies used for predicting cyber risks. As one of the main methods, the application of machine learning models in the process of collecting and processing data through artificial intelligence-based analyses was carried out. In addition, classification of attacks and predictive analysis on behavioral models were carried out using big data technologies. The predictive methods were formulated using various analytical programs for analyzing behavioral data recorded during cyber attacks (Table 1).

Cyber risks arise in cyberspace with different types of threats, and their impact levels lead to different consequences for organizations. The main types of cyber risks include phishing, malware, DDoS attacks, insider threats, and ransomware. Phishing attacks are carried out through deceptive tactics aimed at stealing personal and corporate information and have a high impact level. Malware severely degrades system performance and security, therefore it has a high impact level. DDoS attacks aim to overload the network and make it inaccessible and are usually characterized by a medium impact level. Insider threats are threats that violate privacy within

Table 1 Cyberrisk types and impact level.

Cyberrisk type	Description	Potential impact level
Phishing	Deceptive tactics to steal sensitive information	High
Malware	Malware that affects system performance	High
DDoS attacks	Attacks aimed at overloading the network	Medium
Insider threats	Threats originating from within the organization	Medium

Source Burrell [8]

the organization and have a medium impact level. Ransomware, on the other hand, encrypts files and demands a ransom and has a high impact level (Table 2).

Various technologies are employed for cyber risk prediction, with advanced methods such as artificial intelligence (AI), machine learning (ML), big data analytics, and behavioral analysis playing a crucial role. AI enables anomaly detection and risk forecasting in cyberspace by analyzing data from diverse sources to identify potential threats and facilitate proactive measures. ML automates threat detection by recognizing attack patterns, while big data analytics processes large-scale datasets to uncover cyberattack structures and recurring threats. Behavioral analysis enhances security effectiveness by identifying anomalous activities within networks.

Cyber risk prediction relies on various methods, including statistical analysis, machine learning models, pattern recognition, and heuristic analysis. Statistical analysis identifies potential threats based on existing trends in cyberspace, offering moderate effectiveness. Machine learning models, with high efficiency, utilize diverse algorithms to detect cyber threats in advance. Pattern recognition identifies recurring threat patterns, aiding in the optimization of security measures. Heuristic analysis, based on behavioral characteristics, detects threats and identifies emerging risks, demonstrating moderate effectiveness (Table 3).

With the increase in cyber risks and the increasingly complex and damaging nature of cyber attacks, cyberspace security has become one of the main priorities in the modern world. The rapid development of digital technologies necessitates the implementation of new security measures in a wide range, from individual users to government agencies. The use of technologies such as artificial intelligence and machine learning is of great importance in the field of cybersecurity to predict risks and minimize damage. These technologies allow for effective and timely measures

Table 2 Key technologies for cyber risk prediction.

Technology	Description	Application area
Artificial intelligence	Used for prediction and anomaly detection	Risk prediction
Machine learning	Automates threat detection based on various patterns	Threat detection
Big data analytics	Analysis of large volumes of data	Data processing

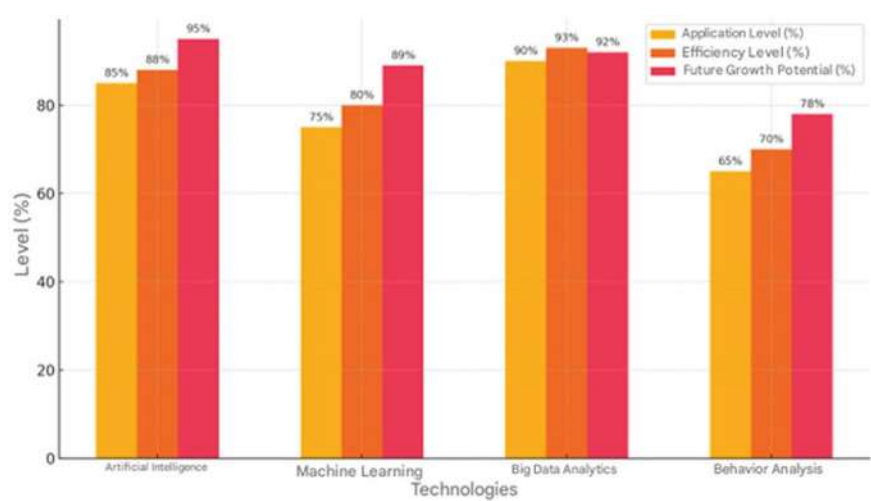
Source Ansari et al. [6]

Table 3 Cyber risk forecasting technologies and application level.

Technology	Implementation level (%)	Effectiveness level (%)	Future growth potential (%)
Artificial intelligence	85	88	95
Machine learning	75	80	89
Big data analytics	90	93	92
Behavioral analytics	65	70	78

Source Alibeigi and Munir [4]

to be taken by analyzing the types of cyber attacks in advance and reducing their potential impact. With the expansion of cyberspace and the changing digital security environment, the need for analytical approaches in the process of detecting new cyber threats is increasing. Big data analytics plays an important role in preventing cyber attacks, as patterns and trends of threats are identified through the analysis of very large volumes of data. Big data analytics helps cybersecurity analysts analyze large data sets to detect system vulnerabilities and apply advanced predictive methods. This technology not only enhances security measures but also establishes a more effective security infrastructure by minimizing the impact of cyber risks (Graph 1).



Graph 1 Application and effectiveness level of technologies in cyberrisk forecasting. Source Algarni et al. [3]

3 Mathematical Formulations

1. **Risk Value Calculation:** The total value of cyber risk is calculated based on probability and impact factors:

$$R = P \times I$$

Example: Suppose the probability of a cyber attack occurring and the impact level $P = 0.3$ is $I = 200,000$ USD. In this case, the risk value will be:

$$R = 0.3 \times 200,000 = 60,000$$

This value estimates the financial losses of cyber risk.

2. **Anomaly Detection Result Formula:** For anomaly detection, the difference between the data in the data set and the average value is taken as the basis. This allows us to detect anomalies due to the normal distribution of the data:

$$Z = \frac{X - \mu}{\sigma}$$

Example: If the given data value is $X = 15$ the mean of the data is $\mu = 10$, and the standard deviation is $\sigma = 2$, the Z-score is calculated as follows:

$$Z = \frac{15 - 10}{2} = 2.5$$

This result indicates that the Z-score is greater than 2 and indicates that the data may be an anomaly.

3. **Cyberattack Prediction Formula (Bayes Theorem):** The probability of cyberattacks is estimated by Bayes Theorem:

$$P(A/B) = \frac{P(A/B) \times P(A)}{P(A)}$$

Example: Suppose that an anomaly was observed in a system (event B) for a cyberattack to occur (event A). It is known that the probability of a cyberattack when an anomaly is observed is $P(A/B) = 0.8$, the total probability of a cyberattack is $P(A) = 0.5$ and the probability of an anomaly being observed in the system is $P(B) = 0.2$. In this case, the probability of an attack will be:

$$P(A/B) = \frac{0.8 \times 0.05}{0.2} = 0.2$$

This calculation indicates that there is a 20% probability of a cyberattack when an anomaly is observed.

Cloud technologies provide a flexible and scalable platform for cyber risk management by centralizing cybersecurity data and offering extensive analytical resources. These platforms enable real-time risk monitoring and forecasting. Cloud-based monitoring and alert systems serve as effective tools for cyber risk prediction.

Quantum computing enhances cybersecurity by enabling the development of more complex and robust threat models. With the ability to process vast and intricate datasets at unprecedented speeds, quantum computers offer greater accuracy in cyber threat forecasting compared to traditional methods. This advancement strengthens and stabilizes security systems.

Blockchain technology ensures data immutability, facilitating cyber risk tracking and auditing. By enabling secure and transparent data exchange, blockchain enhances risk prediction effectiveness. It minimizes data manipulation risks and provides a high level of security, contributing to proactive cyber risk mitigation.

As cyberspace rapidly expands, ensuring data security at both individual and organizational levels has become a top priority. With the increasing frequency and complexity of cyberattacks, effective risk prediction and analysis require the integration of advanced technologies and methodologies.

Cyber risk prediction identifies threat levels before attacks occur, enabling timely security measures. This process enhances the effectiveness of security strategies, helping prevent data loss and financial damage.

Cyber risk management is a structured process involving five key steps: Identification, Assessment, Evaluation, Implementation, and Monitoring. The process begins with identifying potential threats and vulnerabilities in information systems. Next, risk assessment quantifies potential impacts using relevant data. The evaluation phase analyzes existing security measures to identify weaknesses and improve strategies. Implementation involves deploying selected risk management controls to strengthen cybersecurity defenses. Finally, continuous monitoring ensures adaptability to evolving threats, maintaining resilience and optimizing security measures.

4 Conclusion

In general, cyber risk prediction and management play an important role in protecting digital security. With the widespread use of Internet technologies, the complexity and number of cyber attacks are increasing, which emphasizes the need for new security approaches at the individual and organizational levels. The application of advanced technologies such as artificial intelligence, machine learning, big data analytics, and behavioral analysis for analyzing and predicting cyber risks makes a significant contribution to improving security. These technologies allow for continuous monitoring of cyberspace, detection of anomalies, and timely implementation of security measures. In the modern era, cyber risk management is not limited to preventing attacks; at the same time, new approaches are required for predicting threats and optimizing risks. Artificial intelligence and machine learning models allow for more accurate and automated analysis of potential threats. Big data analytics analyzes cyber

attack patterns using extensive databases, while behavioral analysis strengthens the early warning system by monitoring anomalies within the network. These approaches increase the resilience of cybersecurity and ensure the stability of digital infrastructure. Thus, the widespread application of innovative technologies is essential to protect the security of cyberspace and effectively manage cyber risks. These measures not only increase the level of security, but also ensure the sustainable development of the digital environment. These results show that the application of advanced technologies to ensure cybersecurity and continuous innovations in the field of cyber risk prediction make security strategies more effective and reliable.

References

1. M. Abdalla, Y.B. Arshad, Information security: cybersecurity standards adoption among malaysian public listed companies. *Int. J. Eng. Res. Technol.* **9**(8) (2020)
2. M.C. Abeysekera, P. Kumarawadu, Analysis of factors influencing blockchain implementation in finance sector in Sri Lanka. *Ho Chi Minh City Open Univ. J. Sci.—Econ. Bus. Adm.* **12**(2), 3–14 (2022)
3. A. Algarni, V. Thayanathan, Y.K. Malaiya, Quantitative assessment of cybersecurity risks for mitigating data breaches in business systems. *Appl. Sci.* **11**(8), 3678 (2021)
4. A. Alibeigi, A.B.B. Munir, Malaysian personal data protection act, a mysterious application. *Univ. Bologna Law Rev.* **5**(2), 362–374 (2020)
5. A.A. Aljumah, T.A. Ahanger, Cyber security threats, challenges and defence mechanisms in cloud computing. *IET Commun.* **14**(7), 1185–1191 (2020)
6. M.M. Ansari, P.K. Sharma, B. Dash, Prevention of phishing attacks using AI-based cybersecurity awareness training. *Int. J. Smart Sens. Adhoc Netw.* **3**(3), 61–72 (2022)
7. H. Balisane, E.I. Egho-Promise, E. Lyada, F. Aina, Towards improved threat mitigation in digital environments: a comprehensive framework for cybersecurity enhancement. *Int. J. Res.—GRANTHAALAYAH* **12**(5). ISSN 2394-3629 (2024)
8. D.N. Burrell, Understanding the talent management intricacies of remote cybersecurity teams in COVID-19 induced telework organisational ecosystems. *Land Forces Acad. Rev.* **25**(3), 232–244 (2020)
9. M. Campbell, M. Egerstedt, J.P. How, R.M. Murray, Autonomous driving in urban environments: approaches, lessons and challenges. *Philos. Trans. R. Soc. A: Math. Phys. Eng. Sci.* **368**(1928), 4649–4672 (2010)
10. J. Canelon, E. Huerta, J. Incera, T. Ryan, A cybersecurity control framework for blockchain ecosystems. *Int. J. Digit. Account. Res.* **19**, 103–144 (2019)
11. J. Cui, G. Sabaliauskaite, L.S. Liew, F. Zhou, B. Zhang, Collaborative analysis framework of safety and security for autonomous vehicles. *IEEE Access* **7**, 148672–148683 (2019)
12. M.M. Dacorogna, M. Kratz, Managing cyber risk, a science in the making. *ArXiv* **2023**(10), 1000–1021 (2023)
13. B. de Silva, Exploring the relationship between cybersecurity culture and cyber-crime prevention: a systematic review. *Int. J. Inf. Secur. Cybercrime* **12**(1), 23–29 (2023)
14. Y. Ding, K. Li, C. Liu, Z. Tang, K. Li, Short-term load forecasting using hybrid machine learning techniques. *Energy* **223**, 120055 (2021)
15. B. Dow, E. Burke, Cybersecurity and the law: a global perspective. *J. Cybersecur.* **5**(3), 77–95 (2019)
16. E.L. Egho-Promise, E. Lyada, F. Aina, Towards improved vulnerability management in digital environments: a comprehensive framework for cyber security enhancement. *Int. Res. J. Comput. Sci.* **11**(05), 441–449. ISSN 2393-9842 (2024)

17. M. Elbahri, M. Makarov, F. Ahlers, Blockchain integration in smart cities: a review of challenges and opportunities. *Int. J. Urban Sci.* **25**(5), 555–572 (2021)
18. S.L. Pfleeger, D.D. Caputo, Leveraging behavioral science to mitigate cybersecurity risk. *Comput. Secur.* **31**(4), 597–611 (2012). <https://doi.org/10.1016/j.cose.2012.03.003>
19. H. Wu, Z.J. Zhang, W. Li, Information technology solutions, challenges, and suggestions for tackling the COVID-19 pandemic. *Int. J. Inf. Manage.* **57**, 102287 (2021)
20. A. Yousaf, A. Amro, P.T.H. Kwa, M. Li, J. Zhou, Cyber risk assessment of cyber-enabled autonomous cargo vessel. *Int. J. Crit. Infrastruct. Prot.* **46**, 100695 (2024). <https://doi.org/10.1016/j.ijcip.2024.100695>

Soft Computing and Fuzzy Logic

From Question to Answer: Innovative AI-Based System for Exam Preparation



Shahnaz N. Shahbazova  and Agha P. Aghayev

Abstract In today's fast-paced educational environment, preparing for exams, including open-ended exams, is a unique challenge for students. Traditional study and revision methods often require significant time and effort, and the lack of immediate feedback can hinder student progress. To address these issues, we developed an AI-powered system to simplify and streamline the exam preparation process for students. The system automatically extracts questions from provided files and lecture notes and allows students to answer them in an open-ended format. Using natural language processing (NLP) and advanced AI algorithms, the system evaluates answers based on accuracy, relevance, completeness, and originality, and checks them for plagiarism. It also provides personalized suggestions for correction so that students can improve their answers and prepare much better for the exam. Our solution not only optimizes the learning process by providing instant feedback and adaptive study suggestions, but also improves critical thinking and writing skills. This article mentions the features and benefits of the system, as well as its potential to transform traditional test preparation.

Keywords AI-powered system · Open-ended questions · Exam preparation · Natural language processing (NLP) · Automated question extraction · Answer evaluation · Plagiarism detection · Personalized feedback · Adaptive learning · Critical thinking · Educational technology · Instant feedback · Writing skills · Academic improvement

S. N. Shahbazova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, UNEC, Baku, Azerbaijan
e-mail: shahbazova.shahnaz@unec.edu.az

A. P. Aghayev

Azerbaijan State University of Economics, UNEC, Baku, Azerbaijan
e-mail: agaevaga425@gmail.com

1 Introduction

In the modern world, education is getting through significant changes, and one of the key factors influencing these changes is technology development. Traditional methods of learning and preparing for exams are gradually being replaced by innovative solutions, such as online courses, learning management systems, and, of course, artificial intelligence. With an ever-increasing amount of information and an accelerating pace of life, it becomes increasingly difficult to prepare effectively for exams, especially when faced with complex open-ended questions that require in-depth analysis and critical thinking.

Preparing for exams that contain open-ended questions is traditionally one of the most difficult and stressful tasks for students. Unlike multiple-choice exams, where students can intuitively select answers, open-ended questions require students not only to know the facts but also to be able to analyze and interpret information and formulate their thoughts. This can be an overwhelming task for many students, especially those who do not have strong writing and reasoning skills. At the same time, preparing for such exams often takes a lot of time and effort. After all, to pass the exam, students need to study all the material in-depth, learn how to correctly format their answers, and ensure they can answer the questions. In addition, students often face a lack of timely feedback that would help them correct mistakes and improve their results [1].

Our project is dedicated to solving these problems. We have developed an intelligent system designed to simplify and optimize the process of preparing for open-ended test questions. The main goal of the system is to provide students with a tool that can automatically generate questions and intelligently analyze the answers. Artificial intelligence not only checks the correctness, relevance, and completeness of answers but also provides personalized suggestions for improvement. In addition, the system checks answers for plagiarism, promotes the academic integrity of students, and helps avoid the common problem of plagiarism in other people's work.

The program is not just a tool for preparing for exams. It is a sophisticated educational platform that helps students gain a deeper understanding of the material and prepare effectively for exams, thereby acquiring valuable skills for further study. The system uses advanced natural language processing (NLP) technology to automatically extract questions from lecture materials and tutorials. This makes the preparation process more personalized as students can tailor questions to suit their course of study and exam requirements [2].

2 The Main Idea of the Project

The goal of the project is to create an intelligent system that will significantly simplify and automate the process of preparing for exams on open questions. The system is based on the use of artificial intelligence and performs several key tasks, from automatically generating questions to analyzing and evaluating student answers:

Automatic extraction and generation of questions: The system provides students with questions that can be downloaded as pre-prepared files or automatically extracted from the course material. Text analysis algorithms recognize the main topics of the lesson and generate questions based on them. This allows students to immediately begin preparing answers without wasting time on researching and creating questions [3].

Open Answers: Students are given the opportunity to exercise creativity and critical thinking by answering questions in an open-ended format. This format helps not only to evaluate knowledge, but also to develop the ability to correctly organize ideas, discuss and structure information.

Intelligent control using artificial intelligence: AI analyzes each response according to a number of parameters:

- Accuracy—checks whether the answer is factually correct and consistent with theoretical concepts and evidence.
- Relevance—checks how well the answer corresponds to the topic raised and how well it covers it.
- Plagiarism—checks borrowings from third-party sources and ensures that the answer is original and unique [4].
- Integrity—analyzes the answer for the completeness of the argument and the depth of the topic.

Suggestions for improving the answer: After analyzing the answer, the AI not only gives a final score, but also makes suggestions for improving it. These suggestions include tips for correcting deficiencies, improving your writing style, deepening your knowledge of specific topics, and additional self-test questions.

Instant Feedback: One of the key features of the system is the ability to provide instant feedback. Providing students with immediate analysis of their answers allows them to quickly identify gaps in their knowledge and begin to address them. This makes the exam preparation process more flexible and adaptable.

Personalization and adaptation: Based on the results of the analysis, the system can adapt the difficulty of questions and sentences to the level of knowledge of a particular student. This provides an individual approach, allowing each student to develop at their own pace and overcome their weaknesses [5].

3 Advantages of the Project

- *Saving time*: The process of preparing for exams, especially exams with open-ended questions, requires a lot of time searching and structuring information, preparing and checking answers. The system automates much of this process, allowing students to focus on what matters the most: studying material and preparing answers. The system automatically retrieves questions from uploaded files and course materials, eliminating the need to manually prepare questions and significantly reducing preparation time. It also saves time by instantly analyzing answers and eliminating the need to wait for feedback from the teacher [6].
- *Improved Preparation*: Detailed analysis using artificial intelligence allows students to improve their answers on four key criteria:
 - Accuracy—the system points out factual errors and helps students correct them.
 - Relevance—AI helps students avoid unnecessary deviations from the topic, concentrating on the main one.
 - Completeness—the system analyzes the answer to ensure that all necessary aspects have been covered, helping the student provide a more complete and detailed answer.
 - Plagiarism—prevents students from accidentally or intentionally copying information by providing original answers and ideas.
- *Individual suggestions*—AI not only checks the answers, but also gives specific suggestions for improving them. The system can give advice on sentence structure, improve argumentation, point out missing elements in an answer, or even suggest better wording. This helps students not only pass the exam, but also develop skills that will help them in their future studies. Personal tutors allow students to consciously overcome their weaknesses and improve their results through self-study.
- *Instant Feedback*: One of the most important benefits is that students receive instant feedback on every answer. Students do not have to wait for the teacher to check or analyze errors themselves; AI provides an immediate and detailed analysis and solution. Therefore, students can quickly identify weak points in their answers and correct them. This approach makes the learning process more dynamic and interactive [7].
- *Personalization and adaptability*: Based on previous answers, the system can adapt to the student's knowledge level and select questions and suggestions based on their strengths and weaknesses. For example, if a student is having difficulty with a particular subject, the system can suggest additional questions or materials for review. Thus, an individual preparation plan is drawn up for each student, which increases his chances of success. This is especially helpful for students of different skill levels, as they can each progress at their own pace.
- *Objective Score*: AI removes human error from answer-checking and ensures fair and objective scoring. This is important because subjective assessment can lead to errors, underestimation of knowledge and unfairly high grades: AI makes

the assessment process transparent and fair for all students by strictly assessing answers according to certain criteria. This is especially valuable in situations where the teacher cannot pay enough attention to each student.

- *Ease of Use*: The system's interface is designed to make it easy for students to download materials, answer questions, and get results. The intuitive interface and functionality simplify the preparation process and make it accessible to a wide range of users who do not have special technical skills.
- *Boosting student confidence*: Students are often unsure of their knowledge and feel stressed before exams. The system allows students to prepare more consciously, check their results in real time and feel truly prepared for the exam. This reduces anxiety and allows students to approach exams with more confidence in their abilities.
- *Effective plagiarism prevention*: The system not only checks answers for plagiarism, but also helps students better understand how to express their thoughts and ideas correctly. This promotes academic integrity and teaches students original thinking, which is important in an educational environment.
- *Self-assessment and Development Tool*: This project is not just based on answer-checking, it is a complete tool for personal development. It helps students self-assess their knowledge and progress, and gives them the opportunity to improve their answers before contacting teachers. It creates opportunities for independent work and helps develop critical thinking skills by analyzing and structuring information.

4 How the System Works

1. **Loading and searching tasks**: The first step in the system is the process of creating a set of tasks for preparation:
 - a. *Uploading a file with questions for preparation*: a teacher or student can upload a file with a list of questions. These could be ready-made tests, examples from a textbook, or a list of questions provided by the teacher. The system supports various file formats (.txt, .docx, .pdf, etc.) and automatically retrieves questions and configures them for further use.
 - b. *Automatic extraction of questions from lessons*: In the absence of ready-made files with questions, the system can analyze text material from lectures and textbooks. It uses natural language processing (NLP) algorithms to extract key topics and questions based on lecture structure and key concepts. This allows the system to generate relevant questions without teacher input.
2. **Open questions**: the system proposes a question, and the student generates an answer. Answers can be free and provide an opportunity for students to demonstrate their understanding of the topic and develop reasoning skills. Students write their answers in free text without restrictions on structure and volume. This

format encourages creative thinking and promotes better learning. Responses can be uploaded to the system via the web interface or other supported platforms.

3. **Analysis of answers using artificial intelligence:** When a student submits an answer, an analysis process is launched, which includes several stages:

- *Fact checking:* AI analyzes the factual correctness of answers; AI checks the text against the discipline's knowledge base to determine whether the information contained in the answers matches the correct theoretical or practical concepts. For example, if an answer contains incorrect statements, the system flags this and provides feedback.
- *Relevance assessment:* the system checks how relevant the answer is to the question asked. It evaluates how close it is to topic to prevent the student from veering off it or focusing too much on unimportant aspects. This is important as one of the main indicators of a successful answer is its relevance to the question. If a student deviates from the main topic, the system suggests improvements to focus the answer on the essence of the question [8].
- *Plagiarism Analysis:* An important aspect of testing is to identify possible plagiarism. The system uses complex algorithms to check the text for parts borrowed from open sources, scientific papers and previous student work. This prevents academic dishonesty and helps students better understand the need for independent work.
- *Completeness Assessment:* Responses are reviewed for completeness and whether or not they cover all aspects of the question. It is assessed on how deeply the student covered the topic and how fully he defended his point of view. If the answer is incomplete or missing important details, the system suggests ways to improve it by adding the missing elements.

4. **Providing feedback and recommendations:** After the AI has analyzed the answer, the system prepares a detailed report. The report identifies both the strengths of the response and areas for improvement:

- *Error indicator:* Highlights inaccuracies and provides recommendations on the correct data and how to correct them.
- *Relevance hints:* If a student deviates from the topic, AI will take note of this and offer correction to the situation.
- *Tips to improve the completeness of answers:* If the answer is not detailed enough, the system will suggest ways to deepen the answer.
- *Suggestions for eliminating plagiarism:* If the answer contains borrowings, the system will note this and offer to partially paraphrase the answer to make it original.
- *Style and structure recommendations:* Depending on the structure of the text, the system may suggest ways to make answers more logical and consistent, point out flaws in the argument, or suggest stylistic improvements.
- *Adaptive scoring system:* one of the most important features of the system is its adaptability; AI can adapt to the level of knowledge of each student. Based on previous responses, they can determine which points require attention.

For example, if a student is good at solving easy questions, the system can suggest more difficult tasks to deepen his knowledge. This way, each student undergoes an individual preparation process and can move forward at their own pace.

- *Instant Feedback*: One of the key points of this system is instant feedback. When submitting a solution, the student receives an analysis of it almost immediately. This speeds up the preparation process as students can immediately recognize their mistakes and begin to correct them. Instant feedback creates a cyclical learning process where students don't have to wait weeks for results but can immediately work to improve their answers.
- *Results processing and reporting*: the system not only provides feedback on individual answers, but also creates a report on the entire preparation process. In this report, you can see how their results have changed over the course of preparation, which topics require additional study, and how their answers have improved over time. This allows students to track their learning progress and be more conscious about preparing for exams.
- *Post-exam analysis and recommendations*: After the exam, the system analyzes the overall results and makes recommendations for improvement. This could be advice on specific topics that need to be reviewed, or suggestions for how to work on the style and structure of your answers in the future. This analysis helps students learn from their mistakes and prepare more effectively for future exams (Figs. 1 and 2).



Fig. 1 Extracting questions from a file

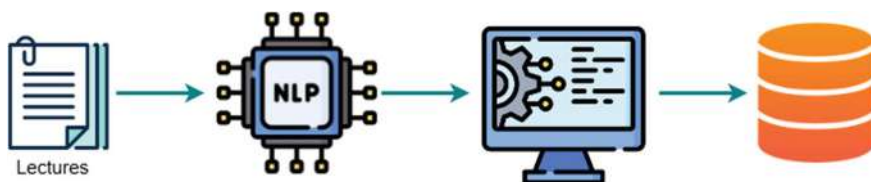


Fig. 2 NLP based question generation process and writing to the database

5 Conclusion

Modern technologies are rapidly changing the approach to education, and our project is a striking example of how artificial intelligence can be effectively integrated into the educational process. Preparing for exams with open-ended questions is always a difficult and responsible task. This is due to the need of deeply analyzing the educational material, preparing reasoned answers, and checking their accuracy and relevance. Our project offers an innovative solution to these problems, combining automation, intelligent analysis and personalized feedback.

The main advantage of this system is that it can be tailored to the needs of each individual student. In a world where everyone has different skill levels, resources and time, a tool that can be tailored to each user's needs is essential. The system becomes a full-fledged learning assistant that not only checks answers, but also helps students better understand their weaknesses and suggests specific ways to improve them. This makes the preparation process more intense and meaningful.

The project also solves one of the most pressing problems of modern education—the fight against plagiarism. Due to open access to a huge amount of information, it is difficult for students to always remain original, and sometimes even accidental borrowing can lead to serious consequences. Our AI tool helps students avoid plagiarism by providing valuable guidance on how to paraphrase and structure their answers to promote academic integrity and original thinking.

Instant feedback is another important advantage of this system. Instead of wasting valuable time with a teacher checking answers, students get immediate results, learn from their mistakes, and adjust their preparation immediately. This approach promotes a more dynamic and effective learning process, allowing students to feel more confident and achieve better exam results.

The project is also of great value to teachers. Automated answer checking significantly reduces teachers' workload, allowing them to spend more time on strategic tasks such as improving the curriculum, individual work with students and improving the quality of the educational process.

In the long term, the project can significantly change the approach to the examination system in educational institutions: the use of artificial intelligence to check open questions will not only speed up and simplify the process, but also improve the quality of education. Students will be more motivated to study the material on their own, and adaptive and personalized recommendations will allow them to study more effectively.

We envision a future where artificial intelligence and advanced technology help students not only pass exams, but also develop the skills they need to succeed in their studies and careers. Our project is a big step in this direction. It is designed to make the preparation process not only simpler, but also more meaningful and effective, which will ultimately improve the quality of knowledge and student performance.

References

1. R. Luckin, W. Holmes, M. Griffiths, L. Forcier, *Artificial Intelligence in Education: Promises and Implications for Teaching and Learning* (2016)
2. M. Potthast, B. Stein, A. Barrón-Cedeño, *Plagiarism Detection Using Machine Learning Algorithms* (2010)
3. L. Pugliese, *Adaptive Learning Technologies in Higher Education* (2016)
4. J. Burstein, M. Chodorow, *Natural Language Processing in Educational Applications* (2012)
5. M.D. Shermis, J. Burstein, *Automated Essay Scoring: Applications to Educational Technology* (2013)
6. R.S.J.D. Baker, G. Siemens, The impact of AI-powered feedback on learning outcomes. *J. Educ. Data Min.* (2014)
7. P. Young, A. Lai, M. Hodosh, J. Hockenmaier, *Natural Language Processing Techniques for Extracting Educational Content* (2014)
8. K. Woo, *AI in Exam Preparation: Improving Efficiency and Outcomes* (2019)

Recursive Approach to Algorithmic Analyses of Language of Mythes and Legends



Mammadova Ana Aga  and Leyla Aliyeva 

Abstract In the article, by means of algorithmic analysis, i.e. by using the method of mathematical modeling, the method of complicated information system organizing is studied. While studying the works of philosophers, basing on myths, legends, fairy tales of different times and peoples, it's shown clearly that in studying complex systems created either by nature or by man, algorithmic analysis of language provides a high degree of understanding of the subject. This is a recursive approach. Nature and art have the same algorithmic principles. This approach is considered only after a certain classification and storage of information.

Keywords Fuzzy logic theory · Artificial intellect · Thought processes · Modeling · Recursion

1 Introduction

In the eternal movement of the world, its own rules, that reflected in language function. The human in the universe who explains the vibrations of the world by sounds, gestures or written signs cognize the laws of all world. While studying the structure of language, the human learns the laws of nature, where universal rules associated with the dynamic structure of change exist. These are the laws of development and interaction. These laws are shown in the world of art, literature, music, painting. The language and human are connected so closely, that scientists had to tackle problem of understanding of the world by human. In different scientific centers of the world, specialists engaged in computer technology began to study mythology, the works of Plato, Hegel, Kant, Kierkegaard, Heidegger, M. Foucault and many other classics of philosophical thought.

M. A. Aga (✉)

Department "Mathematics and Informatics", Baku Slavic University, Baku, Azerbaijan
e-mail: ana.mamedova@outlook.com; ana.mammadova.a@bsu-uni.edu.az

L. Aliyeva

Bolu Abant İzzet Baysal University, Bolu, Turkey

The problem of understanding language by machines has led to the eternal theme of the human phenomenon. Who are we? Where from? Where we are going? "...this is the purpose and meaning of our stay on earth: to think and search, to listen to distant disappeared sounds, since our true homeland is placed behind them" [1].

The searching of exact, accurate formalized means of expressing and studying of development laws began long ago. Pythagoras, who was shocked by the incommensurability of the diagonal and the square sides; Plato, who created theory of dialogical influence for searching the true judgments; young Galois, who in the evening before the mortal duel, proved the impossibility of a solution general polynomial equations in radicals; Einstein, who was looking for the law of interaction between gravity and extractives—all these geniuses thought about this matter. Here it is important to note also that philosophers and mathematicians, having noticed the same laws of development of different objects, proposed a general notion—a complex system. Complex system is a structurally organized object in which conditions, transitions, substructures and interactions of parts are distinguished. The systematic approach makes it possible to analyze them from the position of a unified methodology.

So, we have the mathematical theory of algorithms, invented by logicians in the 30 s of the Twenty century. An algorithm is a complex information system with preset functioning rules. Recursion is one of the most important concepts of this theory and it means a way of system organizing in which certain moments of its development are determined by its rules, it can provoke its own modified copies, can interact with them and include them in its structure. The laws for changing created copies are also included in the system rules and can depend on many parameters. In such case the condition of the system and other subsystems at the moment when the copy is called, the information content of the specified parameters, depending on the rules of the system itself are taken into account. In ancient religions we can find recursive definitions of gods. Recursive methods for description of the world development were expressed in texts for a first time but only many centuries later, they manifested in the theory of algorithms. Recursive procedures have been introduced into modern programming languages.

This concept is well known to programmers who began to study old philosophical tractates. In creative disputes, along with technical terms, the names "Kant", "Hegel", "Heidegger", "Confucius", "Buddha", "Sanskrit" and others are heard. Philosophers, building systems for understanding the world, did not even imagine that their works someday could be useful for practical engineers.

2 Research History

To establish the evolutionary laws of thinking there is no other way than to study the heritage of the distant past given to us in language. All the secrets and clues of existence are hidden in language. Therefore, interest in ancient myths is growing, and systemic mythology and mythical logic of artificial intellect systems becomes possible. The French scientist Levi-Stros [2], having studied the myths of the tribes of

South America, came to conclusion that their plot lines are distinguished by internal stability, i.e. myths reflect the thinking patterns of primitive human. On this basis, scientist managed to discover prehistoric logic in the actions of ancient humans. It can be said, that myths have ontological meaning and in them the basic structures of thought are hidden. It is also interesting that in mythological thinking the same logic works as in scientific thinking is contained, and a person had always thought “well”. Your name alone makes me feel good, Iris is an extraordinary name, I don’t know what it reminds me of.

But you know,” she said, “that it is the name of the blue and yellow cinquefoil”.
[1].

Here the extremely subtle language game draws attention to itself.

Spinoza, back in the Seventeenth century, tried to study the human soul in a logical form: “... there is no doubt, it will seem surprising that I am going to study human vices and stupidities in a geometric way... But my principle is this: in nature there is nothing that one could attribute it to its deficiency, because nature always and everywhere remains the same; ... the laws and rules of nature, according to which everything happens and changes from one form to another, are the same everywhere and always. Therefore, the methods of knowing your nature, what kind of they may be, must be the same, and that is this must be knowing based on the universal laws of the rules of nature” [3].

It is important to have in mind that the main method of myths studying is an inductive technique—i.e. obtaining a general fact based on a comparative analysis of many examples. In the studies of Taylor, the fact of the animation of objects and natural phenomena by primitive people is unquestioning. Taylor called this feature of thinking “animism” (anima—spirit, soul (lat). Taylor affirms that myths came to existence in the period of savagery. All humanity has gone through this stage. Myths contain information that interests us. At the same time, authentic archaic legends are selected, for example, the punishment of a tree if a human fell from it; marriage of the Sun and Moon; devouring of stars by the Sun; turning people into animals, etc.

Human being has always tried to find the beginning of life, “the ground zero”, to understand why “in the beginning was the Word,” why similar stories and legends are repeated all over the world. As older humanity is, the greater abyss of time it can dip into. Moving into the future, we are moving into the past, entering the ancient path. Referring these examples, it can be seen that each purposeful action of a complex system is connected with the concept of an algorithm. Algorithm defines the graduality of actions of an object to achieve a aim. This is how primitive hunters invented algorithms for hunting animals and neighboring tribes, and their wives invented the first culinary recipes, which are also considered to be algorithms. Cognition always searched the ways to describe algorithms. Many ancient logical recipes, math texts, martial arts books have been preserved to the present day.

So, we can fully agree with researchers Novak et al. [4] that the algorithm has many language-dependent phrasings, different manifestastons. An algorithmic system is a system that connects some subsystems in a selected basic language with specified means of interaction and development. When studying complex systems created

by nature or man, the modeling method is used. In this way a higher degree of understanding of the object under study is achieved.

Sometimes the necessary logical strictness of the machine means of algorithms expressing comes into conflict with the principles of functioning of studied object. From other hand, sometimes the exact model of the studied object is replaced by general principles of organization. That's why, researchers often have to use specially a variety of means that imitate uncertainty sets. The famous American scientist (a native of Azerbaijan) Zadeh in his theory of Fuzzy acted exactly in this way. We notice that in literary works the way of closed interacting processes organizing reflects certain algorithmic principles. It is known that programming of algorithmic systems had always existed.

3 Recursion—Analysis and Synthesis of Algorithmic Systems

It should also be noted that recursion, having been enriched in the theory of algorithms and programming, and having become a commonplace method of analysis and synthesis of complex algorithmic systems, returns to the world where it was noticed for a first time and where it has always existed, even remained unrecognized. But now recursion has powerful, developed algorithmic methods in its “store”, and now it is in plain sight. It means that the language itself is recursive phenomenon. From a limited basic set of schemes, myths, tales, and legends, all complex plot structures of modern authors are obtained in a recursive way.

The analysis of the sentences shows that here the recursion manifests itself even more persuasively. It determines the nesting of commentary parts into different parts of the sentence. The sentences that still remain a philological secrecy receive an algorithmic definition, and many problematic facts of syntax become immediately clarified.

So, recursive machines begin to learn to speak and to understand language. The emergence of the concept of “linguistic variable” (the notion was introduced by Zadeh), the meaning of which is a word or sentence both in a natural or artificial language, proves it. More exactly, a “linguistic variable” is described by a set $(X, T(X), B, b, M)$, in which X is the name of the variable $T(X)$ —the term set of X , that is, the set of its linguistic meanings; B —is universal set; b —is syntactic rule that generates the terms of the set M —a semantic rule that associates each linguistic meaning X with its meaning $M(X)$, and $M(X)$ is a fuzzy subset of the set “ u ”. The meaning of linguistic X is characterized by the compatibility function “ b ” and $[0; 1]$, which assigns to each element “ T ” from the set and gives the alignment meaning of this element with X [5].

4 The Essence of Soft Computing

The theory of fuzzy logic is of multifaceted character. This is philosophy, mathematics, cybernetics. Today it has become quite clear that symbol manipulation and predicate logic have serious limitations when working with many world's problems.

In his theory "The role of soft computing and fuzzy logic in the understanding, design and development of information intellectual systems" Zadeh reveals to us that the essence of soft computing (Soft Computing) is that, in differ from traditional, hard computing, they are aimed at adapting to the comprehensive inaccuracy of real world. We can see it also in the recursive understanding of the language of myths and legends. The guiding principle of soft computing is: "tolerance to inexactness, uncertainty and partial truth in order to achieve ease of manipulation, robustness, low cost of solutions and better agreement with reality" [6].

Human thinking is the starting model for soft computing. The author notes that soft computing is not a separate methodology. It is, rather, an unification, a partnership of different directions and trends. Fuzzy logic, neurocomputing, genetic computing and probabilistic computing, with the later inclusion of chaotic systems, trust networks and branches of learning theory are the main partners in this unification.

So, summarizing all points, mentioned above, we can give a definition: soft computing is a term introduced by Zadeh in 1994. This term denotes a set of inexact, approximate methods for solving problems that often do not have a solution in polynomial time.

According to Zadeh, the purpose of a semantic rule is to link the compatibility of primary terms in a composite linguistic meaning. With compound meaning compatibility, by means of linguistic variables one can describe approximately phenomena that are so complex or badly defined that they cannot be described in conventional quantitative terms. Thus, we rely on Zadeh's fuzzy logic, and it can serve as a more realistic diagram of human reasoning than traditional two-valued logic. Computing with linguistic probabilities comes down to solving non-linear programming problems.

5 Describing Phenomena in Myths and Legends Using Linguistic Variables

The main applications of this linguistic approach are contained in the sphere of humanistic systems (i.e. systems on which human behavior, human judgments, perceptions or emotions have strong influence). The human, the person himself, his intellectual processes are considered as a humanistic system. Thus, the law of recursive development of the world was reflected in ancient texts for a first time, then it manifested in literature, and only after it that law got mathematical embodiment in the theory of algorithms and programming. After it the reverse reflection occurred.

In the book “Secret Dede—Gorgud” by the famous Azerbaijani scientist—linguist and writer Abdullah, the author notes: “One cannot help but notice how in dastan the collision of “Nature” and “Culture” is conveyed in artistic form, the presence and opposition of its components in the process of transition from chaotic imaginations about the infinite and undivided world to an orderly, knowable world, to a world that can be explained” [7].

In his book the author points out that the mythological line of thinking, which means a collective image, the thinking of the entire society is gradually left behind and the individual is being separated from the collective. This process is very painful and full of drama. As the writer emphasizes, it is enough to follow the fate of one of dastan heroes- Beyrek, to be convinced of this. Researches of the famous scientist, writer, professor Abdullah, by establishing semantic connections between individual plot situations and characters, writer’s comparisons and commensuration with samples of ancient mythology, raises and analyzes many issues regarding the content of “Kitabi—Dede Gorgud”. This point once again convinces us in the verity of the algorithmic analysis of the history of language emergence.

6 Conclusion

All these peculiarities testify that algorithmic analysis becomes a surprisingly powerful means of cognition and it confirms the unity of the expression of the world by means of technical and humanitarian sciences. The same algorithmic principles function both in nature and creativity. As it is known, the question of the nature and status of human intellect has not been resolved in philosophy. Although a number of hypotheses have been proposed in the early days of artificial intelligence, e.g. the Turing test or the Newell-Simon hypothesis, there is no exact definition whether computers have achieved “rationality,”. That’s why, despite the presence of many approaches, both to understanding problems and to creating intelligent information systems, two approaches to the development of AI stand out.

Top-down, semiotic approach means the creation of expert systems, knowledge bases and logical inference systems that imitate high-level mental processes, such as thinking, reasoning, speech, emotions, etc. Another type approach is ascending, biological one, it means the study of neural nets and evolutionary calculations that model intellectual behavior based on biological elements, as well as the creation of corresponding computing systems, such as a neurocomputer or a biocomputer.

English ethnographers Taylor [8] (1832–1917) and Frazier (1854–1941) analyzed and classified thousands of myths, rituals and customs of the peoples of different countries. Using the formula of primitive thinking, one can determine the unity of the structures of ancient myths of different peoples. This unity surprises researchers still. Then the myths were included in biblical legends, the plots of paintings and folklore. But in the 60 s of the Twenty century, a new scientific study of cultural systems, in language particularly, was formed, that study got the name “structuralism”.

From what moment can we talk about the appearance of “human being”? We don’t know how many centuries it took to consolidate this scheme, we still don’t know how the genetic transmission of thinking schemes occurs, we don’t know in which brain structures they are formed, but we know for sure that they exist, and now we are even able to catch their general appearance. What do we risk seeing when entering the country of myths and legends? We can see vague, forgotten shadows, but they are still new. This country is open to us, but it is also a labyrinth where you wander along lost paths. This path tempt us like a distant light in the darkness of the night.

I would like to end the article with the following words of Atkins: “We are children of chaos, and at the heart of every change decay is hidden very deeply. Initially only a process of dissipation, degradation exists; everyone is overflowed by waves of chaos that has no reason or explanation. There is no initial aim in this process, there is only ongoing movement. However, as we ensured, in this movement different directions are possible, the choice of which is dictated by chance” [9].

References

1. Hesse Hermann. Iris. <http://samlib.ru/k/koncheew/iris.shtml>
2. K. Levi-Stros, *Primitive thinking* M.: TERRA-Book Club; Republic (Library of philosophical thought), p. 392 (1999)
3. B. Spinoza, *World of Book* (2010), p. 352
4. V. Novak, I. Perfilyeva, I. Mochkrozh, *Mathematical principles of fuzzy logic*. M.: Fizmatlit. (2006)
5. L.A. Zadeh, Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets Syst.* **1**(1), 3–28 (1978)
6. L. Zadeh, The Concept of a Linguistic Variable and its application to approximate decision making. M. Mir. 167 (1976)
7. K. Abdullah, *The Secret Dede Gorgud* (Baku, Mutarjim, 2006), p. 346
8. E.B.Taylor, in *Primitive culture* (M. Political literature publisher, 1989), p. 573
9. P. Atkins, Order and disorder in nature. M. Mir. (70), 224 (1987)
10. J. Frezier, *Rituals and Customs of Different Countries*, p. 1039. <https://www.rulit.me/books/zolotaya-vetv-read-301132-1.html>
11. V.V. Kruglov, M.I. Dli, R.Y. Galunov, *Fuzzy logic and artificial neural networks*. M.: Fizmatlit. 36–38 (2001)

A CNN-Based Handwritten Digit Classifier: Architecture, Training, and Performance Evaluation



Shahnaz N. Shahbazova 

Abstract This research proposes a convolutional neural network (CNN)-based architecture for digit classification, trained on the MNIST dataset. The model incorporates a hierarchical structure encompassing two convolutional layers, each followed by max-pooling operations, dropout regularization for mitigating overfitting, and a series of fully connected layers. Implemented using TensorFlow, the architecture achieves an average classification accuracy of 97%. Additional evaluation using a confusion matrix confirms the robustness of the model.

Keywords Digit recognition · Convolutional neural network · Digit classification · MNIST dataset · Deep learning · Image classification · Dropout regularization · Max-pooling · Fully connected layers · Supervised learning

1 Introduction

With the rapid advancement of deep learning technologies, the number of domains where such techniques are applied continues to grow. One of the most widely used deep learning models for visual data processing is the Convolutional Neural Network (CNN) [1, 2]. CNNs have become essential in fields such as object detection, facial recognition, robotics, video analysis, image segmentation, pattern recognition, natural language processing, spam filtering, topic classification, regression, speech recognition, and especially image classification. Among these applications, CNNs have achieved near-human or even superior performance in tasks such as handwritten digit recognition.

The foundational architecture of CNNs is inspired by mammalian visual processing systems, specifically the concept of localized receptive fields in the visual cortex, as first discovered by Hubel and Wiesel [3]. Their work laid the groundwork for Fukushima's Neocognitron [4–6], the first hierarchical model for shift-invariant

S. N. Shahbazova (✉)

Department of Digital Technologies and Applied Informatics, Azerbaijan State University of Economics, UNEC, Baku, Azerbaijan

e-mail: shahbazova.shahnaz@unec.edu.az

pattern recognition. In 1998, LeCun et al. proposed the first functional CNN for digit recognition, leveraging convolutional layers, pooling, and backpropagation to classify digits from pixel-level image input [7, 8].

The CNN used in this study aligns with this legacy. It consists of two convolutional layers with ReLU activation functions and max-pooling operations for dimensionality reduction. To mitigate overfitting, dropout regularization is applied after the second convolutional layer. The resulting feature maps are flattened and passed through a fully connected layer with 50 hidden units, followed by another dropout and a final layer that outputs class probabilities for ten digits (0–9). Unlike traditional artificial neural networks (ANNs), this CNN exploits local connectivity and weight sharing, significantly reducing parameter count while retaining representational power.

Training is carried out by minimizing a cross-entropy loss function, with gradients propagated backward through the network to iteratively adjust weights and biases. This process increases model performance and enables accurate classification of handwritten digits.

2 Methodology

This section provides a comprehensive description of the methodology adopted in this study. It encompasses several key subsections: (2.1) the dataset utilized for training and evaluation, (2.2) the architectural design of the proposed convolutional neural network (CNN), (2.3) the training protocol and (2.4) the development and integration of a graphical user interface (GUI). Each of these components is discussed in detail in the subsequent subsections.

2.1 Dataset

The **MNIST dataset (Modified National Institute of Standards and Technology dataset)** is one of the most widely used datasets for training and evaluating the performance of the algorithms designed for handwritten digit recognition. It consists of grayscale images of handwritten digits (0 through 9), each with a corresponding label indicating the digit it represents. The total number of images is 70,000, which is divided into two parts: 60,000 images for training and remaining 10,000 for evaluating. Each image is 28 pixels in height and 28 pixels in width, resulting in a total of 784 pixels per image. Moreover, pixel values are often normalized to a range of [0, 1] by dividing by 255, thus helping to improve the convergence of learning models during training (Fig. 1).

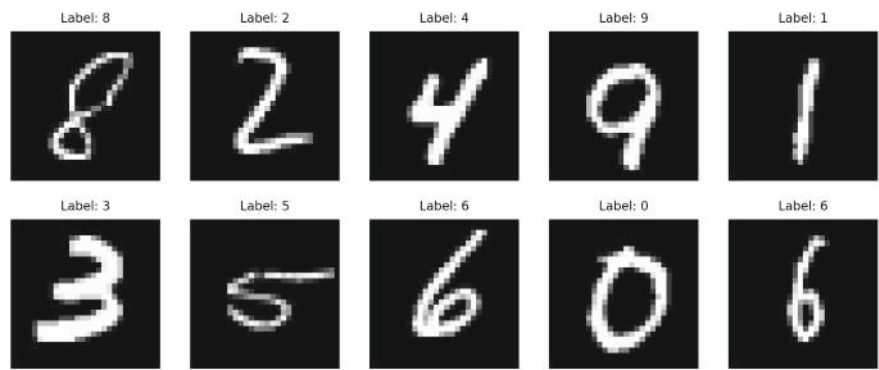


Fig. 1 10 random images from MNIST dataset

2.2 Model Architecture

Table 1 illustrates the layers we defined for training the model and how they affected the shape of the image.

The first convolutional layer applies 10 convolutional filters, each of size 5×5 , to the input image. Given that the input consists of a single grayscale image of dimensions 28×28 , each of the 10 filters generates a separate feature map. Since no padding is applied, the spatial dimensions of the output feature maps are reduced as follows:

$28 \times 28 \rightarrow 24 \times 24$. At this stage, the convolutional layer primarily captures low-level features such as edges and corners.

A max pooling operation is performed with a 2×2 filter, selecting the maximum value from each 2×2 region. This results in a downsampling effect, reducing the spatial dimensions by half: $24 \times 24 \rightarrow 12 \times 12$. The purpose of max pooling is

Table 1 Layer-wise architecture and shape transformations of the proposed CNN model

Layer	Operation	Shape transformation
Conv2D (1 $\rightarrow$ 10, 5×5 kernel)	Convolution	(1, 28, 28) $\rightarrow$ (10, 24, 24)
Max pooling (2 $\times$ 2)	Downsampling	(10, 24, 24) $\rightarrow$ (10, 12, 12)
Conv2D (10 $\rightarrow$ 20, 5×5 kernel)	Convolution	(10, 12, 12) $\rightarrow$ (20, 8, 8)
Dropout	Regularization	No shape change
Max pooling (2 $\times$ 2)	Downsampling	(20, 8, 8) $\rightarrow$ (20, 4, 4)
Flatten	Reshape	(20, 4, 4) $\rightarrow$ (320,)
Fully connected (320 $\rightarrow$ 50)	Dense layer	(320) $\rightarrow$ (50)
Dropout	Regularization	No shape change
Fully connected (50 $\rightarrow$ 10)	Output layer	(50) $\rightarrow$ (10)

twofold: Firstly, to reduce computational complexity, secondly, extract the dominant features, ensuring the most relevant spatial information is retained.

The second convolutional layer applies 20 convolutional filters of size 5×5 to the previous feature maps. Since there were 10 feature maps from the first layer, this layer now operates on these learned representations, capturing more complex hierarchical feature such as specific shapes. As in the first convolutional layer, the lack of padding results in a reduction of spatial dimensions: $12 \times 12 \rightarrow 8 \times 8$.

A dropout layer is applied to randomly disable a subset of neurons during training, thus preventing overfitting. While dropout has no effect on the shape of the feature maps, it improves generalization by reducing dependency on specific neurons and encouraging the network to learn more robust features.

Similar to the first max pooling layer, the second 2×2 max pooling operation is applied to the feature maps, further reducing the spatial dimensions: $8 \times 8 \rightarrow 4 \times 4$. Before passing the learned representations to fully connected layers, the feature maps are flattened into a one-dimensional vector. Given that there are 20 feature maps of size 4×4 , the resulting vector has the following shape: $(20, 4, 4) \rightarrow (320)$.

A fully connected layer applies a linear transformation to the flattened vector, mapping it to a feature representation of dimension 50: $(320) \rightarrow (50)$.

A second dropout layer is introduced to further enhance generalization and prevent overfitting. Again, this does not affect the shape of the data but improves robustness.

The final layer of the network is a fully connected layer that outputs 10 values, each corresponding to a digit class (0–9): $(50) \rightarrow (10)$. These 10 values represent the raw logits, which are subsequently passed through a softmax activation function to produce class probabilities.

The softmax function is applied to the output layer to convert the raw scores (z_i) into probabilities (p_i), ensuring that they sum to 1. The predicted class is determined by selecting the index with the highest probability. The softmax function is defined as:

$$p_i = \frac{e^{z_i}}{\sum_{j=1}^{10} e^{z_j}}$$

where:

- z_i represents the raw output (logit) for class i ,
- e^{z_j} exponentiates each logit, ensuring non-negative values,
- The denominator normalizes the values so they sum to 1.

2.3 Training the Model

The model training and evaluation process was implemented using the PyTorch deep learning framework. To maximize computational efficiency, the model was configured to utilize a Graphics Processing Unit (GPU) when available, with a fallback to CPU otherwise. This conditional device allocation ensures the parallel processing

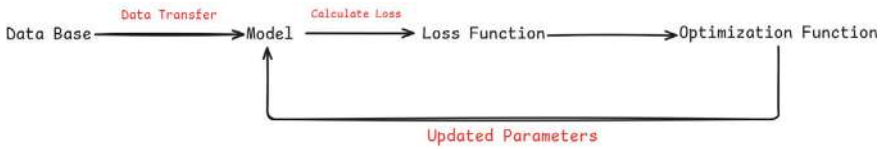


Fig. 2 Model training procedure

capabilities of GPUs are fully exploited, significantly reducing training time for deep neural networks.

The core architecture, as mentioned earlier, was a Convolutional Neural Network (CNN), optimized specifically for the handwritten digit classification task. For model optimization, the Adam optimizer was employed—a popular choice due to its adaptive learning rate and momentum-based updates. The initial learning rate was set to 0.001, a standard baseline that balances convergence speed with gradient stability.

To handle the multi-class classification objective, the Cross-Entropy Loss function was used as the training criterion. The model underwent training for several epochs, each representing a complete iteration over the training dataset. During each epoch, the model was explicitly set to training mode, enabling dynamic components such as dropout layers to operate properly (Fig. 2).

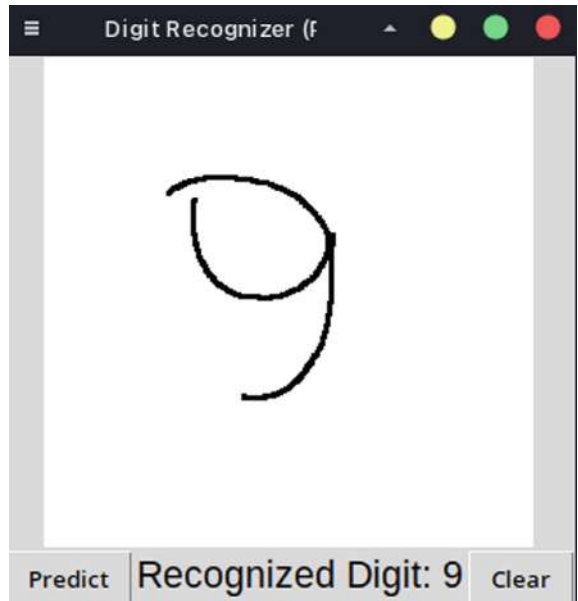
2.4 GUI Integration

The application is developed using the Tkinter library for the graphical user interface (GUI) and PyTorch for loading and utilizing a pre-trained model. Users can draw a digit on a canvas, and the application will preprocess the drawing, pass it through the model, and display the predicted digit. The canvas is a 28×28 pixel area, matching the input dimensions of the MNIST dataset. Users can draw on the canvas using their mouse, and the drawing is stored as a grayscale image. To provide visual feedback, the preprocessed image is displayed using Matplotlib.

Before performing recognition, the drawn image undergoes preprocessing. The image is converted to grayscale and resized to 28×28 pixels. Although resizing is not strictly necessary, this step ensures alignment with the standard transformations applied to MNIST images. Subsequently, pixel values are normalized to the range $[0, 1]$ and inverted, as MNIST images typically have white backgrounds and black digits.

Finally, the image is converted into a PyTorch tensor with dimensions $(1, 1, 28, 28)$, representing batch size, channels, height, and width, respectively. The preprocessed image is then passed through the model to generate predictions. The model outputs a probability distribution across the 10-digit classes (0–9). The predicted digit is determined by selecting the class with the highest probability (Fig. 3).

Fig. 3 GUI app illustration



3 Evaluation and Results

Below we present the model’s performance using metrics such as accuracy and confusion matrix (Fig. 4).

The provided **line graph** represents the model’s accuracy over 10 training epochs. The x-axis denotes the number of epochs, while the y-axis shows the accuracy values.

Initially, the model starts with an accuracy slightly above 93%, showing rapid improvement in the first few epochs. Between epochs 1 and 3, there is a steep increase in accuracy, reaching around 96%. Following this, from epoch 4 onward, the rate of accuracy improvement slows down, indicating that the model is converging. By epoch 6, the accuracy reaches approximately 97%, with minor improvements in the later epochs. This implies that the model has almost saturated its learning capacity, and further training may cause overfitting or undesirable results (Fig. 5).

The given **confusion matrix** compares model predictions with actual values. In essence, the confusion matrix divides predictions into 4 categories: true positives, true negatives, false positives and finally false negatives.

First of all, the majority of predictions lie along the diagonal, indicating that most samples were classified correctly. For instance, the model accurately classified 968 instances of digit 0.

However, the model exhibits certain misclassification patterns, specifically for digits with similar curves and shapes. Notably, digit 5 was misclassified as 3 in 17 instances. Similarly, digit 2 was mistaken for 7 in 16 cases, indicating potential ambiguity in distinguishing these numbers. Additionally, some confusion exists between

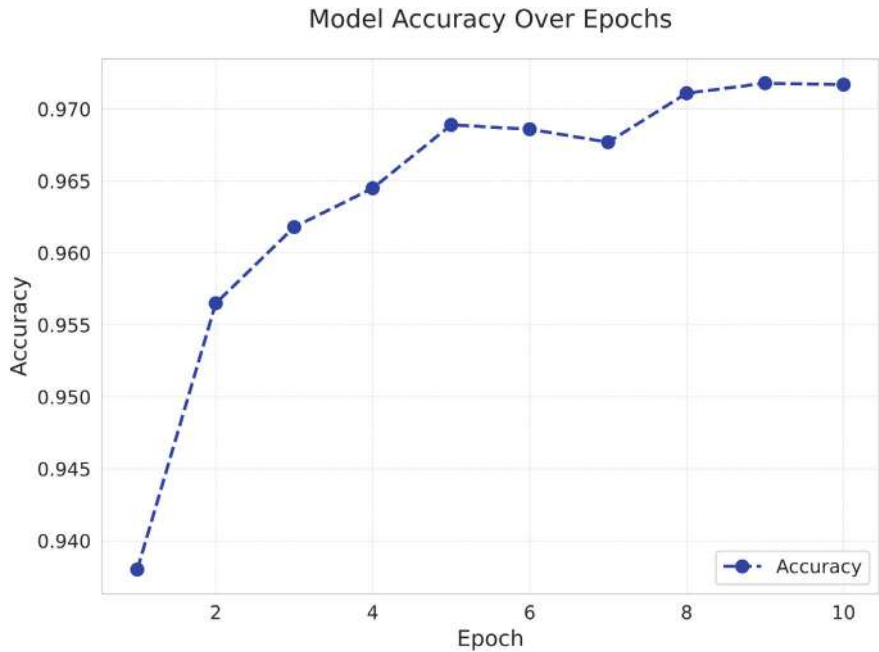


Fig. 4 Model accuracy over epochs

8, 3, and 9, further reinforcing the idea that visually similar digits pose a challenge for the model.

4 Conclusion

The analysis of the model’s performance using both accuracy metrics and the confusion matrix provides a view of the effectiveness and limitations of the implemented Convolutional Neural Network (CNN) for handwritten digit recognition.

The model demonstrated a strong learning curve during training. As shown in the accuracy plot over 10 epochs, the model quickly attained high accuracy in the early stages of training, with significant performance gains within the first three epochs. This rapid increase reflects the model’s ability to capture fundamental patterns in the MNIST dataset, such as simple edges and curves. The plateauing of the accuracy beyond the sixth epoch suggests convergence, indicating that the model had extracted most of the relevant features and was no longer making significant improvements.

The confusion matrix offers deeper insight into the model’s classification behavior than accuracy. The strong presence of values along the matrix diagonal confirms high overall accuracy and indicates that the model correctly identified most digit classes.

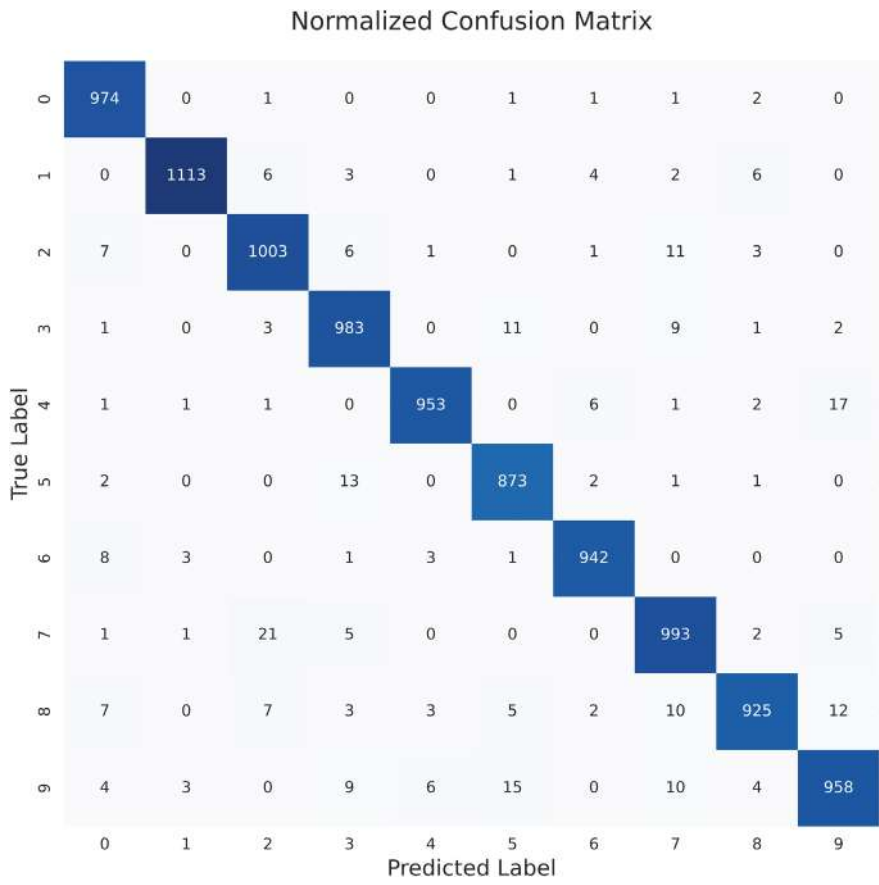


Fig. 5 Confusion matrix

This reflects the successful learning of both low-level and higher-order features necessary for digit recognition.

Nevertheless, the matrix also reveals key areas for improvement. Misclassifications, though relatively rare, tend to occur among digits with visually similar edges and corners. For example, digits such as 5 and 3, or 2 and 7, share structural similarities that can confuse the convolutional layers—especially when written with stylistic variation or noise. This suggests that while the model has effectively generalized across most digit classes, its feature representations for certain ambiguous cases may not be sufficiently distinct.

In conclusion, the implemented CNN model achieves commendable performance on the MNIST digit recognition task, with high accuracy and a confusion matrix dominated by correct predictions. The findings validate the effectiveness of CNNs for visual classification tasks but also highlight the intrinsic challenges posed by ambiguous handwritten input.

When it comes to practical sides of our model, one of the most prominent areas of application is in automated form processing, such as the digitization of handwritten data from bank forms, tax documents, surveys, and institutional records. In such scenarios, the CNN model can replace manual data entry, thus reducing labor costs and human error. Another critical application is in postal services, where handwritten postal codes or numeric addresses on packages must be interpreted rapidly. The integration of digit recognition systems into sorting machines helps to automate this process.

References

1. Y LeCun 1989 Backpropagation applied to handwritten zip code recognition *Neural Comput.* 1 4 541–551
2. A. Krizhevsky, I. Sutskever, G.E. Hinton, Imagenet classification with deep convolutional neural networks. NIPS (2012)
3. D Hubel T Wiesel 1971 Aberrant visual projections in the Siamese cat *J. Physiol.* 218 1 33–62
4. Y LeCun Y Bengio G Hinton 2015 Deep learning *Nature* 521 7553 436
5. D. Cireşan, U. Meier, J. Schmidhuber, *Multi-column Deep Neural Networks for Image Classification*. arXiv preprint [arXiv:1202.2745](https://arxiv.org/abs/1202.2745) (2012)
6. K. Fukushima, S. Miyake, Neocognitron: a self-organizing neural network model for a mechanism of visual pattern recognition, in *Competition and Cooperation in Neural Nets* (Springer, 1982), pp. 267–285
7. Y. LeCun et al., Handwritten digit recognition with a back-propagation network. NIPS (1990)
8. Y LeCun L Bottou Y Bengio P Haffner 1998 Gradient-based learning applied to document recognition *Proc. IEEE* 86 11 2278–2324

Author Index

A

Abdullaev, Fatali, [151](#)
Aga, Mammadova Ana, [199](#)
Aghayev, Agha P., [189](#)
Aliev, Telman, [105](#)
Aliyeva, Leyla, [199](#)
Alvarez, Robert, [161](#)

B

Baquier, Alondra, [19](#)
Baral, Chitta, [79](#)
Bayramov, Azad, [143](#), [151](#)
Beltran, Bradley, [19](#)

C

Ceberio, Martine, [161](#)
Chekina, A. V., [85](#)
Choudhury, Anuradha, [71](#)

E

Erentürk, Köksal, [113](#)
Erentürk, Saliha, [113](#)
Escamilla, Jeffrey, [65](#)

F

Filippov, Aleksey A., [55](#)

G

Grinberg, Galyna, [131](#)
Gunel, Aliyeva, [121](#)
Guskov, Gleb Yu., [55](#)

H

Hajiyev, Semral, [177](#)
Haque, Md Ahsanul, [71](#)

J

Jafarova, Hilala, [167](#), [177](#)

K

Kosheleva, Olga, [19](#), [95](#)
Kreinovich, Vladik, [3](#), [11](#), [19](#), [65](#), [71](#), [79](#),
[95](#), [161](#)

L

Lyubchyk, Leonid, [131](#)

M

Miki-Silva, Gabriel, [19](#)
Muñoz, Fernando, [11](#)
Musaeva, Naila, [105](#)

N

Nowmi, Saeefa Rubaiyet, [71](#)

R

Rena, Huseynova, [121](#)
Revayet, Huseynova Zinyet, [43](#)
Romanov, Anton A., [55](#)
Ruiz, Salvador, [161](#)
Ryen, Ahmed Ann Noor, [71](#)

S

Saika, Sabrina, [71](#)

Servin, Christian, [11](#)

Shahbazova, Shahnaz N., [3](#), [27](#), [33](#), [95](#), [177](#),
[189](#), [207](#)

Suleymanov, Samir, [143](#), [151](#)

Y

Yamkovyi, Klym, [131](#)

Yarushkina, Nadezhda G., [55](#), [85](#)

Yu, Guskov G., [85](#)