

**МАТЕРІАЛИ НАУКОВО-ПРАКТИЧНИХ
КОНФЕРЕНЦІЙ ЧОРНОМОРСЬКОГО
НАЦІОНАЛЬНОГО УНІВЕРСИТЕТУ
ІМЕНІ ПЕТРА МОГИЛИ**

№ 1(1), 2025

**Матеріали ХХІІ Міжнародної наукової конференції
«ОЛЬВІЙСЬКИЙ ФОРУМ – 2025: стратегії країн
Причорноморського регіону в геополітичному просторі»**

16-21 червня 2025 р., м. Миколаїв, Україна

СЕРІЯ: ТЕХНІЧНІ НАУКИ

НАУКОВИЙ ЗБІРНИК

Заснований у 2025 році

Виходить 2 рази на рік



Миколаїв-2025

Засновано у 2025 році. Виходить 2 рази на рік.

Засновник і видавець:

Чорноморський національний університет імені Петра Могили

РЕДАКЦІЙНА РАДА

Голова редакційної ради:

Леонід Клименко – доктор технічних наук, професор, заслужений діяч науки і техніки України, професор, в.о. ректора університету, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна).

Заступники голови:

Юрій Котляр – доктор історичних наук, професор, перший проректор, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Роман Дінжос – доктор технічних наук, професор, проректор з наукової роботи, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Володимир Ємельянов – доктор наук з державного управління, заслужений діяч науки і техніки України, професор, директор Навчально-наукового інституту публічного управління та адміністрування, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна).

Члени редакційної ради:

Світлана Белінська – доктор економічних наук, професор, декан факультету економічних наук, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Анжела Бойко – кандидат технічних наук, доцент кафедри комп'ютерної інженерії, декан факультету комп'ютерних наук, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Віталій Вербицький – кандидат історичних наук, доцент, в.о. декана факультету фізичного виховання і спорту, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Олена Кузнецова – кандидат педагогічних наук, доцент кафедри соціальної роботи, управління і педагогіки, в.о. директора Навчально-наукового медичного інституту, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Ганна Норд – доктор економічних наук, професор, директор Навчально-наукового інституту післядипломної освіти, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Дмитро Січко – кандидат юридичних наук, доцент, декан юридичного факультету, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Анастасія Хмель – кандидат історичних наук, доцент, декан факультету політичних наук, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна);

Людмила Шерстюк – кандидат педагогічних наук, доцент, в.о. декана факультету філології, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна).

**Свідоцтво про реєстрацію
друкованого засобу
масової інформації:**
Серія, номер R30-05946
від 27.03.2025 року

*Рекомендовано до друку та
поширення мережею
рішенням Вченої ради ЧНУ
№ 5 від 29.05.2025 р.*

Мови видання:
українська, англійська,
мови ЄС

Електронна версія:
[//mspc.mk.ua](https://mspc.mk.ua)

Адреса редакції:
м. Миколаїв,
вул. 68 Десантників, 10,
54003
Тел.: 8 (0512) 76-55-99,
8 (0512) 76-55-81.
факс 50-00-69, 50-03-33
E-mail: rvv@chmnu.edu.ua

М 32 Матеріали науково-практичних конференцій Чорноморського національного університету імені Петра Могили. Серія: Технічні науки : Ольвійський форум – 2025 : стратегії країн Причорноморського регіону в геополітичному просторі : XXII міжнар. наук. конф. 16-22 черв. 2025 р., м. Миколаїв : матеріали конф. / М-во освіти і науки України ; ЧНУ імені Петра Могили. – Миколаїв : Вид-во ЧНУ ім. Петра Могили, 2025. – 136 с.

УДК 004+3+62+7/9]:001.891](477)(06)

**MATERIALS OF THE SCIENTIFIC
AND PRACTICAL CONFERENCE
OF PETRO MOGYLA BLACK SEA
NATIONAL UNIVERSITY**

№ 1(1), 2025

**Materials of the XXII International Scientific Conference
«Olbia Forum – 2025: strategies of the countries of the Black
Sea region in the geopolitical space»**

June 16-22, 2025, Mykolaiv, Ukraine

SERIES: TECHNICAL SCIENCES

SCIENTIFIC COLLECTION

Founded in 2025

Published 2 times a year



Mykolaiv-2025

Founded in 2025. It is published 2 times a year.
Founder and publisher:
Petro Mohyla Black Sea National University

EDITORIAL COUNCIL

Chairman of the Editorial Board:

Leonid Klymenko - Doctor of Technical Sciences, Professor, Honored Worker of Science and Technology of Ukraine, Professor, Acting Rector of the University, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine).

Deputy Chairmen:

Yurii Kotliar - Doctor of History, Professor, First Vice-Rector, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Roman Dinzhos - Doctor of Technical Sciences, Professor, Vice-Rector for Research, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Volodymyr Yemelianov - Doctor of Science in Public Administration, Honored Worker of Science and Technology of Ukraine, Professor, Director of the Educational and Research Institute of Public Administration and Management, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine).

Members of the Editorial Board:

Svitlana Belinska - Doctor of Economics, Professor, Dean of the Faculty of Economics, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Anzhela Boyko - PhD in Engineering, Associate Professor of the Department of Computer Engineering, Dean of the Faculty of Computer Science, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Vitalii Verbytskyi - PhD in History, Associate Professor, Acting Dean of the Faculty of Physical Education and Sports, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Olena Kuznetsova - PhD in Pedagogy, Associate Professor of the Department of Social Work, Management and Pedagogy, Acting Director of the Educational and Research Medical Institute, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Hanna Nord - Doctor of Economics, Professor, Director of the Educational and Research Institute of Postgraduate Education, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Dmytro Sichko - PhD in Law, Associate Professor, Dean of the Faculty of Law, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Anastasiia Khmel - PhD in History, Associate Professor, Dean of the Faculty of Political Science, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine);

Liudmyla Sherstiuk - PhD in Pedagogy, Associate Professor, Acting Dean of the Faculty of Philology, Petro Mohyla Black Sea National University (Mykolaiv, Ukraine).

The conference is registered in the catalog of international conferences in Academic Resource Index / Research Bib.

The conference is registered in UkrNTI №358 dated April 07, 2025.

Identifier in the Register of Subjects in the Media Sector:
R30-05946 (Decision of the National Council of Ukraine on Television and Radio Broadcasting No. 656 dated March 27, 2025).

Recommended for publication by the decision of the Academic Council of Petro Mohyla Black Sea National University (protocol № 5, 29.05.2025).

Languages of the publication:
Ukrainian, English,
EU languages

Electronic version:
[//mspc.mk.ua](https://mspc.mk.ua)

Editorial address:
Mykolaiv,
10 Desantnykov St., 68, 54003
tel.: 8 (0512) 76-55-99,
8 (0512) 76-55-81.
fax: 50-00-69, 50-03-33
e-mail: rvv@chmnu.edu.ua

M 32 Materials of scientific and practical conferences of the Petro Mohyla Black Sea National University. Series: Technical sciences : Olbia Forum – 2025: strategies of the countries of the Black Sea region in the geopolitical space: XXII int. scientific conf. June 16-22, 2025, Mykolaiv: materials of the conf. / Ministry of Education and Science of Ukraine; Petro Mohyla Black Sea National University. – Mykolaiv: Publishing house of Petro Mohyla Black Sea National University, 2025. – 136 p.

UDC 004+3+62+7/9]:001.891](477)(06)

**РЕДАКЦІЙНА КОЛЕГІЯ
СЕРІЯ: ТЕХНІЧНІ НАУКИ**

Головний редактор:

Анжела Бойко кандидат технічних наук, доцент, доцент кафедри комп'ютерної інженерії, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Заступники головного редактора:

Едуард Лисенков доктор фізико-математичних наук, професор, завідувач кафедри фізики та математики, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Свген Давиденко кандидат технічних наук, доцент, завідувач кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Юрій Кондратенко доктор технічних наук, професор, завідувач кафедри інтелектуальних інформаційних систем, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Члени редакційної колегії:

Свген Сіденко кандидат технічних наук, доцент, доцент кафедри інтелектуальних інформаційних систем, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Ірина Манькусь кандидат педагогічних наук, доцент, доцент кафедри фізики та математики, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Альона Швед докторка технічних наук, професорка, професорка кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Інесса Кулаковська кандидат фізико-математичних наук, доцент, доцент кафедри інтелектуальних інформаційних систем, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Микола Фісун доктор технічних наук, професор, професор кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Сергій Гузій кандидат фізико-математичних наук, доцент, доцент кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

Відповідальний секретар:

Ігор Кандиба PhD, старший викладач кафедри інтелектуальних інформаційних систем, Чорноморський національний університет імені Петра Могили (м. Миколаїв, Україна)

**РЕДАКЦІЙНА КОЛЕГІЯ
СЕРІЯ: ТЕХНІЧНІ НАУКИ**

Головний редактор:

Anzhela Boiko Candidate of technical sciences, Associate Professor, Associate Professor of the Department of Computer Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Заступники головного редактора:

Eduard Lysenkov Doctor of Physical and Mathematical Sciences, Professor, Head of the Department of Physics and Mathematics, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Yevhen Davydenko Candidate of technical sciences, Associate Professor, Head of the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Yurii Kondratenko Doctor of Technical Sciences, Professor, Head of the Department of Intelligent Information Systems, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Члени редакційної колегії:

Yevhen Sidenko Candidate of technical sciences, Associate Professor, Associate Professor of the Department of Intelligent Information Systems, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Iryna Mankus Candidate of Pedagogical Sciences, Associate Professor, Associate Professor of the Department of Physics and Mathematics, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Alona Shved Doctor of Technical Sciences, Professor, Professor of the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Inessa Kulakovska Candidate of Physical and Mathematical Sciences, Associate Professor, Associate Professor of the Department of Intelligent Information Systems, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Mykola Fisun Doctor of Technical Sciences, Professor, Professor of the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Serhii Huzii Candidate of Physical and Mathematical Sciences, Associate Professor, Associate Professor of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

Відповідальний секретар:

Ihor Kandyba PhD, Senior Lecturer at the Department of Intelligent Information Systems, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine

ЗМІСТ

ТЕХНІЧНІ НАУКИ	10
Катерина Антіпова АВТОМАТИЗОВАНЕ ВИЯВЛЕННЯ ТЕКСТУ, ЗГЕНЕРОВАНОГО ВЕЛИКИМИ МОВНИМИ МОДЕЛЯМИ	13
Андрій Блиндарук, Олена Шаповалова END-TO-END ДИФЕРЕНЦІЙОВАНИЙ ЛАНЦЮГ NURBS – GEOM-GNN – TGN ДЛЯ ПРОСТОРОВО-ЧАСОВОЇ ІДЕНТИФІКАЦІЇ РУХОМИХ ОБ'ЄКТІВ.....	16
Надія Болюбаш, Володимир Блинцов ІНФОРМАЦІЙНА СИСТЕМА АНАЛІЗУ ТА МОНІТОРИНГУ ПРОДУКТИВНОСТІ ВЕБРЕСУРСІВ	20
Стефанія Бондаренко, Микита Ковальчук КІБЕРФІЗИЧНІ СИСТЕМИ В МОДЕЛЮВАННІ ТА ОРГАНІЗАЦІЇ КОМП'ЮТЕРНИХ СИСТЕМ ЛОГІСТИКИ	26
Оксана Брагінець ВІЗУАЛІЗАЦІЯ ЧИСЕЛЬНИХ РОЗВ'ЯЗКІВ РІВНЯНЬ МАТЕМАТИЧНОЇ ФІЗИКИ З ВИКОРИСТАННЯМ MAPLE.....	30
Віктор Гожий ПІДХІД ДО МОДЕЛЮВАННЯ ВЗАЄМОДІЇ КОМПОНЕНТІВ ІНТЕЛЕКТУАЛЬНИХ WEB-СЕРВІСІВ	35
Олександр Гожий, Ірина Калініна МЕТОДОЛОГІЯ АНАЛІЗУ ТА ПОПЕРЕДНЬОЇ ОБРОБКИ ДАНИХ ДЛЯ ВИРІШЕННЯ ЗАДАЧ МАШИННОГО НАВЧАННЯ	39
Олександр Єлісєєв, Гліб Горбань АВТОМАТИЧНА ГЕНЕРАЦІЯ ПРОЄКТІВ ЗА ДОПОМОГОЮ LLM-АГЕНТІВ	45
Ірина Калініна, Олександр Гожий ПОБУДОВА БАГАТОРІВНЕВИХ АНСАМБЛЕЙ В ЗАДАЧАХ МАШИННОГО НАВЧАННЯ	49

Назарій Кіріяк, Гліб Горбань ІНФОРМАЦІЙНА СИСТЕМА ЕЛЕКТРОННОЇ КОМЕРЦІЇ НА ОСНОВІ СУЧАСНИХ ВЕБТЕХНОЛОГІЙ	54
Олександр Ковалів, Євген Сіденко, Галина Кондратенко ПОРІВНЯЛЬНИЙ АНАЛІЗ РЕЗУЛЬТАТІВ ЗАСТОСУВАННЯ РІЗНИХ АРХІТЕКТУР НЕЙРОННИХ МЕРЕЖ ДЛЯ ПРОГНОЗУВАННЯ ІНФЕКЦІЙНИХ ЗАХВОРЮВАНЬ	58
Роман Кошовий, Катерина Кірей ГЕЙМИФІКАЦІЯ НАВЧАЛЬНОГО ПРОЦЕСУ ПІДГОТОВКИ МАЙБУТНІХ ІТ- ФАХІВЦІВ	63
Олександр Обушко, Інесса Кулаковська СИСТЕМА РОЗПІЗНАВАННЯ І КЛАСИФІКАЦІЇ ФІНАНСОВИХ ДОКУМЕНТІВ	67
Віктор Раленко, Євгеній Стоєв ЗАСТОСУВАННЯ FIREBASE STUDIO В РОЗРОБЦІ REACT ЗАСТОСУНКІВ	73
Максим Салютін, Інесса Кулаковська АСПЕКТИ ЗАСТОСУВАННЯ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ НА ПРИКЛАДІ СКРИПТІВ JAVASCRIPT НА ПЛАТФОРМІ MOODLE	77
Максим Салютін, Євген Сіденко, Юрій Кондратенко СУЧАСНІ ПІДХОДИ ДО РОЗПІЗНАВАННЯ ВІЙСЬКОВИХ ОБ'ЄКТІВ ІЗ ВИКОРИСТАННЯМ МОДЕЛІ YOLOV8	84
Олександр Скакодуб РОЗРОБКА ЛЮДИНО-МАШИННОГО ІНТЕРФЕЙСУ ДЛЯ ПОШУКУ ЗНИКЛИХ ТУРИСТІВ ГРУПОЮ БПЛА	91
Микита Смоленський, Євген Сіденко ФОРМАЛІЗОВАНІ ПІДХОДИ ДО ІНТЕГРАЦІЇ RFID-ТЕХНОЛОГІЙ У ЦИФРОВУ ІНФРАСТРУКТУРУ МЕДИЧНИХ ЗАКЛАДІВ	96
Іван Сова, Олексій Козлов АДАПТИВНІ МЕТОДИ МАШИННОГО НАВЧАННЯ ДЛЯ РОЗПІЗНАВАННЯ БПЛА В УМОВАХ ЗМІННОГО РАДІОЧАСТОТНОГО СЕРЕДОВИЩА	100
Богдан Сомряков, Світлана Боровльова КЛАСИФІКАЦІЯ ВІЙСЬКОВОЇ ТЕХНІКИ ЗА ДОПОМОГОЮ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ (CNN)	104
Євгеній Стоєв, Віктор Раленко АРХІТЕКТУРНІ ПРИНЦИПИ ПОБУДОВИ ВІДМОВСТІЙКИХ І ВИСОКОПРОДУКТИВНИХ ОБЧИСЛЮВАЛЬНИХ СИСТЕМ РЕАЛЬНОГО ЧАСУ: ВПЛИВ ПАТЕРНІВ ПРОЄКТУВАННЯ НА ФУНКЦІОНАЛЬНІ ХАРАКТЕРИСТИКИ ...	108

Микола Фісун, Ігор Кандиба СТВОРЕННЯ СЛОВНИКА ПРЕДМЕТНОЇ ГАЛУЗІ ЗАСОБАМИ PYTHON.....	111
Альона Швед, Євген Давиденко ЗАСТОСУВАННЯ СВР-ПІДХОДУ В СИСТЕМАХ СИТУАЦІЙНОГО УПРАВЛІННЯ	115
Олександр Яновський, Марина Фаленкова, Юрій Решетнік ЕЛЕКТРОННА КОМЕРЦІЯ: СУЧАСНІ ТЕНДЕНЦІЇ, ВИКЛИКИ ТА ПРОГНОЗИ.....	119
Roman Dinzhos, Nataliia Fialko NON-ISOTHERMAL CRYSTALLIZATION OF POLYMER NANOCOMPOSITES BASED ON POLYCARBONATE FILLED WITH CARBON NANOTUBES	123
Bohdan Somryakov, Viktor Ralenko SQL DATA ANALYSIS FOR BIG DATA OPTIMIZATION: BEST PRACTICES AND PITFALLS	128
Bohdan Somriakov, Ievgen Sidenko, Yuriy Kondratenko COMPARATIVE ANALYSIS OF THE APPLICATION RESULTS OF MAX-MIN CONVOLUTION BASED ON VARIOUS T-NORM AND S-NORM OPERATORS	132

CONTENTS

TECHNICAL SCIENCES	12
Kateryna Antipova AUTOMATED DETECTION OF TEXT GENERATED BY LARGE LANGUAGE MODELS ..	13
Andrii Blyndaruk, Olena Shapovalova END-TO-END DIFFERENTIABLE NURBS – GEOM-GNN – TGN PIPELINE FOR SPATIOTEMPORAL IDENTIFICATION OF MOVING OBJECTS	16
Nadiia Boliubash, Volodymyr Blyntsov INFORMATION SYSTEM FOR ANALYSIS AND MONITORING OF WEB RESOURCES PERFORMANCE	20
Bondarenko Stefaniia, Kovalchuk Mykyta CYBER-PHYSICAL SYSTEMS IN MODELING AND ORGANIZATION OF COMPUTER SYSTEMS IN LOGISTICS	26
Oksana Brahinets VISUALIZATION OF NUMERICAL SOLUTIONS OF MATHEMATICAL PHYSICS EQUATIONS USING MAPLE	30
Viktor Gozhyj APPROACH TO MODELING INTERACTION OF COMPONENTS OF INTELLECTUAL WEB SERVICES	35
Oleksandr Gozhyj, Iryna Kalinina METHODOLOGY OF DATA ANALYSIS AND PRE-PROCESSING FOR SOLVING MACHINE LEARNING PROBLEMS	39
Yelisieiev Oleksandr, Horban Hlib AUTOMATIC PROJECT GENERATION WITH THE HELP OF LLM AGENTS	45
Iryna Kalinina, Oleksandr Gozhyj CONSTRUCTION OF MULTILEVEL ENSEMBLES IN MACHINE LEARNING PROBLEMS.....	49

Kiriiak Nazarii, Horban Hlib INFORMATION SYSTEM OF E-COMMERCE BASED ON MODERN WEB TECHNOLOGIES	54
Oleksandr Kovaliv, Ievgen Sidenko, Galyna Kondratenko COMPARATIVE ANALYSIS OF THE RESULTS OF THE APPLICATION OF DIFFERENT NEURAL NETWORK ARCHITECTURES FOR THE PREDICTION OF INFECTIOUS DISEASES	58
Roman Koshovyi, Kateryna Kirei, GAMIFICATION OF THE EDUCATIONAL PROCESS OF FUTURE IT SPECIALISTS	63
Oleksandr Obushko, Inessa Kulakovska SYSTEM FOR RECOGNIZING AND CLASSIFYING FINANCIAL DOCUMENTS	67
Viktor Ralenko, Yevhenii Stoiev USING FIREBASE STUDIO IN DEVELOPING REACT APPLICATIONS	73
Maksym Salyutin, Inessa Kulakovska ASPECTS OF THE APPLICATION OF INFORMATION AND COMMUNICATION TECHNOLOGIES USING THE EXAMPLE OF JAVASCRIPT SCRIPTS ON THE MOODLE PLATFORM	77
Maksym Salyutin, Ievgen Sidenko, Yuriy Kondratenko, MODERN APPROACHES TO MILITARY OBJECT RECOGNITION USING THE YOLOV8 MODEL	84
Oleksandr Skakodub DEVELOPMENT OF A HUMAN-MACHINE INTERFACE FOR SEARCHING MISSING TOURISTS BY A UAV GROUP	91
Mykya Smolenskyi, Ievgen Sidenko, FORMALIZED APPROACHES TO INTEGRATING RFID TECHNOLOGIES INTO THE DIGITAL INFRASTRUCTURE OF HEALTHCARE INSTITUTIONS	96
Ivan Sova, Oleksiy Kozlov, ADAPTIVE MACHINE LEARNING METHODS FOR UAV RECOGNITION IN A VARIABLE RADIO FREQUENCY ENVIRONMENT	100
Bohdan Somriakov, Svitlana Borovlova CLASSIFICATION OF MILITARY VEHICLES USING CONVOLUTIONAL NEURAL NETWORKS (CNNs)	104
Yevhenii Stoiev, Viktor Ralenko, ARCHITECTURAL PRINCIPLES FOR BUILDING FAULT-TOLERANT AND HIGH- PERFORMANCE REAL-TIME COMPUTING SYSTEMS: THE IMPACT OF DESIGN PATTERNS ON FUNCTIONAL CHARACTERISTICS	108

Mykola Fisun, Ihor Kandyba PYTHON-BASED CREATION OF A DOMAIN DICTIONARY	111
Alona Shved, Yevhen Davydenko APPLICATION OF THE CBR APPROACH IN SITUATION MANAGEMENT SYSTEMS....	115
Oleksandr Yanovskyi, Maryna Falenkova, Jurij Reshetnik ELECTRONIC COMMERCE: CURRENT TRENDS, CHALLENGES AND FORECASTS	119
Roman Dinzhos, Nataliia Fialko NON-ISOTHERMAL CRYSTALLIZATION OF POLYMER NANOCOMPOSITES BASED ON POLYCARBONATE FILLED WITH CARBON NANOTUBES	123
Bohdan Somryakov, Viktor Ralenko SQL DATA ANALYSIS FOR BIG DATA OPTIMIZATION: BEST PRACTICES AND PITFALLS	128
Bohdan Somriakov, Ievgen Sidenko, Yuriy Kondratenko COMPARATIVE ANALYSIS OF THE APPLICATION RESULTS OF MAX-MIN CONVOLUTION BASED ON VARIOUS T-NORM AND S-NORM OPERATORS	132

Розділ
ТЕХНІЧНІ НАУКИ

Section
TECHNICAL SCIENCES

DOI 10.34132/mspc2025.01.14.01

Катерина Антіпова

АВТОМАТИЗОВАНЕ ВИЯВЛЕННЯ ТЕКСТУ, ЗГЕНЕРОВАНОГО ВЕЛИКИМИ МОВНИМИ МОДЕЛЯМИ

Тези представляють аналіз можливостей сучасних детекторів щодо виявлення тексту згенерованого великими мовними моделями. Зокрема розглянуто використання штучного інтелекту у сфері освіти для академічного шахрайства. Якість згенерованого тексту зростає, що ускладнює виявлення плагіату. Ідентифікація згенерованих робіт та звітів у сфері освіти має деякі відмінності порівняно з іншими сферами в основному через специфіку та контекстуальність навчальних завдань. На теперішній час розроблено багато детекторів для виявлення згенерованого тексту. Сучасний детектор має бути надійним у розпізнаванні таких текстів, щоб запобігти неправомірному використанню штучного інтелекту.

Ключові слова: академічне шахрайство, великі мовні моделі, генерація тексту, детектори плагіату, штучний інтелект.

Поява генеративних інструментів штучного інтелекту (ШІ), які базуються на використанні можливостей великих мовних моделей (ВММ), суттєво вплинула на різні галузі людської діяльності. Здатність цих інструментів генерувати текст, схожий на людську мову, становить значну загрозу для академічної доброчесності. Представлення текстів, створених великими мовними моделями, як власної роботи, називається ШІ-плагіатом. Залежність здобувачів вищої освіти від інструментів генерування тексту призводить до втрати креативності та здатності до навчання. Якщо здобувачі отримують оцінки за роботу інших, тоді знижується важлива зовнішня мотивація для отримання знань і компетенцій. Так само спотворюється й оцінка компетентності здобувачів, що може призвести до отримання ними неправомірної вигоди.

Можливості передових ВММ вплинули на сферу освіти різними способами. Наприклад, здобувачі вищої освіти використовують ChatGPT для виконання практичних робіт і складання іспитів. Це викликало занепокоєння щодо поточних систем оцінювання, які використовуються у вищих навчальних закладах. Викладачі намагаються виявляти шахрайські дії здобувачів, і плагіат є однією з головних проблем. У минулому плагіат здебільшого полягав у представленні робіт і есе, які містили абзаци з різних джерел без їх цитування, але з появою ВММ здобувачі тепер можуть використовувати ШІ для генерації текстів для виконання різнопланових завдань.

Типологія плагіату варіюється залежно від типу даних або рівня затемнення. Автори роботи [1] представили різні типології, визначені в кількох дослідницьких роботах, і запропонували нову типологію плагіату за рівнем затемнення: плагіат зі збереженням символів, плагіат зі збереженням синтаксису, плагіат зі збереженням семантики, плагіат зі збереженням ідеї та «гострайтинг» (англ. *ghostwriting*). Тип плагіату може також відрізнитися залежно від типу даних, наприклад, плагіат коду або плагіат тексту. Виявлення плагіату також можна класифікувати на основі кількості мов, використаних в одному тексті. Крім того, коли плагіат виявляється лише за допомогою самого тексту, це внутрішнє виявлення плагіату. Якщо ж плагіат виявляють у порівнянні з іншим текстом, то йдеться про зовнішнє виявлення плагіату.

З передбачуваним швидким розвитком високопродуктивних ВММ якість вихідних текстів зростає, що ускладнює їх виявлення. Навіть людині важко розпізнати згенерований текст через високу здатність ВММ робити його все більше й більше схожим на написаний людиною. Ідентифікація згенерованих робіт та звітів у сфері освіти має деякі відмінності порівняно з іншими сферами в основному через специфіку та контекстуальність навчальних завдань. Хоча загальний контент, створений ШІ, може демонструвати помітні невідповідності в широкому контексті, більш вузький і більш структурований обсяг академічних завдань може маскувати такі аномалії, роблячи процес виявлення більш складним. Отже, виявлення практичних завдань, згенерованих ШІ, вимагає більш досконалих і контекстно-залежних алгоритмів.

Основні відмінності між згенерованим текстом та написаним людиною:

– ШІ дає організовані відповіді з чіткими, структурованими абзацами тексту або пунктами переліку. Текст може виглядати надто відшліфованим або навіть механічним. З іншого боку, текст, написаний людиною, зазвичай

більш варіативний по структурі. Може спостерігатися певна дезорганізація або зміна тону, які виникають внаслідок природного ходу думок, через що текст здається менш «досконалим», ніж згенерований.

– III, як правило, використовує нейтральний, майже офіційний тон, якщо немає інших вказівок. Він уникає емоційних крайнощів, на відміну від людини, яка надає емоційні відповіді, наприклад, з доданням певного гумору. Такий текст часто більш динамічний, з різними настроями залежно від ситуації. Загалом, люди, як правило, вносять більше індивідуальності у свої тексти.

– III добре комбінує наявні знання, але він не є по-справжньому творчим у людському розумінні. Його відповіді є похідними, тобто базуються на великій базі знань, і тому приклади можуть не справляти враження свіжих чи оригінальних. Люди набагато кращі у творчому мисленні, часто представляють нові перспективи, несподівані ідеї та унікальні підходи до вирішення проблем.

– ВММ навчені розуміти контекст, але мають певні обмеження при роботі зі складними контекстами певними нюансів. Модель може давати занадто буквальну відповідь, оскільки покладається на узагальнені знання, а не на конкретний особистий досвід.

– У великих згенерованих текстах III, як правило, зберігає зв'язність протягом деякого часу, але може почати втрачати фокус або повторюватися через декілька абзаців. Люди краще зберігають довгострокову зв'язність, але вони також можуть повторювати ідеї або змінювати перспективу.

– Хоча III створює граматично правильний і безпомилковий контент, він все одно може припускатися помилок у міркуваннях або фактах, які людина може легко помітити. Він також може іноді генерувати незграбні фрази, через що текст може здаватися неприродним. Люди частіше роблять граматичні або друкарські помилки, пропускають слова тощо.

Ці узагальнені характеристики вказують на те, що III значно покращився в завданнях обробки природної мови для широкого спектру областей. Порівняно з людиною, йому може бракувати індивідуальності, але він має більш комплексний і нейтральний погляд на питання. Інакше кажучи, згенерований текст має тенденцію бути граматично відшліфованим, структурно чітким і нейтральним за тоном, з тенденцією до точності та узагальнення. З іншого боку, написаний людиною текст має тенденцію до більшої варіативності в тоні, структурі та глибині, з унікальними особистими штрихами та випадковими недосконаlostями, які роблять його більш автентичним.

На теперішній час розроблено багато детекторів для виявлення згенерованого тексту, наприклад, DetectGPT [2], RADAR [3], Ghostbuster [4], GPT-Sentinel [5]. Детектор має бути надійним у розпізнаванні текстів, згенерованих III, щоб запобігти неправомірному використанню ВММ. Сучасний детектор контенту, створеного ВММ, повинен мати такі ключові характеристики:

– *Точність* означає, що модель має бути здатною розрізняти згенерований і написаний людиною текст, досягаючи відповідного компромісу між точністю та швидкістю відкликання.

– *Ефективність* означає, що детектор повинен оперувати якомога меншою кількістю прикладів з мовної моделі;

– *Узагальнення* означає, що детектор повинен мати можливість працювати узгоджено, незалежно від будь-яких змін в архітектурі моделі, довжини запиту або набору навчальних даних.

– *Пояснюваність* означає, що детектор повинен надавати чіткі пояснення причин своїх рішень. Вона спрямована на те, щоб забезпечити надання зрозумілих для людини міркувань для рішень щодо виявлення, зазвичай представлені у вигляді пояснень природною мовою або візуального представлення основних ознак.

Завдання можна розділити на три рівня:

– пряме пояснення (безпосереднє визначення ознак підробки за допомогою невеликих прикладів у контексті);

– пояснення на основі міркувань (обґрунтування з кількома кроками та оцінка логічної узгодженості);

– детальний аналіз у вільній формі (детальний аналіз аспектів підробки, узгоджений із задалегідь визначеною таксономією ознак підробки).

Незважаючи на те, що автоматичне виявлення може виявити певний плагіат, існуючі дослідження [6] показали, що програмне забезпечення зіставлення тексту не тільки не знаходить увесь III-плагіат, але й позначає написаний текст як плагіат, таким чином надаючи хибнопозитивні результати. Це неприйнятний сценарій в академічному середовищі, оскільки чесного здобувача можуть звинуватити в порушенні академічної доброчесності.

Сучасні дослідження показують, що найсучасніші детектори III-текстів демонструють значне погіршення продуктивності, коли стикаються з текстами, написаними не носіями англійської мови [7]. Інструменти виявлення також виявляються нездатними впоратися з текстами, перекладеними з інших мов. Крім того, продуктивність детекторів ще більше погіршується з використанням методів обфускації, таких як ручне редагування або машинне перефразування. Згідно зі звітом, опублікованим OpenAI, їхній III-детектор текстів не є повністю надійним у цьому плані. У звіті наводиться оцінка деяких складних випадків для англійських текстів, їхній класифікатор правильно ідентифікує лише 26% згенерованого тексту, тоді як неправильно класифікує 9% тексту, написаного людиною.

Точність і надійність інструментів виявлення текстів, створених III, може варіюватися залежно від кількох факторів, таких як конкретний інструмент, що використовується, тип моделі III, що генерує текст, і контент,

який аналізується. Більшість інструментів виявлення досягають 70-80% точності при виявленні тексту, згенерованого такими моделями, як GPT-3. Детектори також не можуть впоратися з короткими абзацами тексту, а також з результатами роботи більш розвинутих моделей пізніших поколінь, таких як GPT-4.

Отже, проблему академічного шахрайства неможливо вирішити лише за допомогою суто технічних заходів, оскільки така конкуренція може тривати вічно. Таким чином, освітні заклади повинні розглянути можливість прийняття альтернативних освітніх рішень. Одним із рішень для запобігання академічним порушенням може бути вдосконалення поточних стратегій оцінювання в університетах. Оцінювання має більше зосереджуватися на критичному мисленні здобувача та навичках розв'язання проблем, а не на навичках запам'ятовування фактів. Адже прості питання можна легко вирішити за допомогою ВММ. Таким чином, викладачі повинні призначати більш творчі проекти, створювати екзаменаційні питання з відкритою відповіддю, які вимагають від здобувачів критичного мислення та застосування їхніх навичок вирішення проблем.

References

1. Foltyniek, T., Meuschke, N., & Gipp, B. (2019). Academic plagiarism detection: a systematic literature review. *ACM Computing Surveys (CSUR)*, 52(6), pp. 1-42. URL: <https://doi.org/10.1145/3345317>
2. Mitchell, E., Lee, Y., Khazatsky, A., Manning, C.D., & Finn, C. (2023). DetectGPT: Zero-Shot Machine-Generated Text Detection using Probability Curvature. *International Conference on Machine Learning*. URL: <https://doi.org/10.48550/arXiv.2301.11305>
3. Hu, X., Chen, P., & Ho, T. (2023). RADAR: Robust AI-Text Detection via Adversarial Learning. URL: <https://doi.org/10.48550/arXiv.2307.03838>
4. Verma, V. K., Fleisig, E., Tomlin, N., & Klein, D. (2023). Ghostbuster: Detecting Text Ghostwritten by Large Language Models. *North American Chapter of the Association for Computational Linguistics*. URL: <https://doi.org/10.48550/arXiv.2305.15047>
5. Chen, Y., Kang, H., Zhai, V., Li, L., Singh, R., & Ramakrishnan, B. (2023). GPT-Sentinel: Distinguishing Human and ChatGPT Generated Content. URL: <https://doi.org/10.48550/arXiv.2305.07969>
6. Tang, R., Chuang, Y., & Hu, X. (2023). The Science of Detecting LLM-Generated Text. *Communications of the ACM*, 67, 50-59. URL: <https://doi.org/10.48550/arXiv.2303.07205>
7. Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J.Y. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4. DOI: <https://doi.org/10.48550/arXiv.2304.02819>

DOI 10.34132/mspc2025.01.14.01

Kateryna Antipova

AUTOMATED DETECTION OF TEXT GENERATED BY LARGE LANGUAGE MODELS

The theses present an analysis of the capabilities of modern detectors to detect text generated by large language models. In particular, the use of artificial intelligence for academic fraud is considered. The quality of the generated text is increasing, which makes plagiarism detection more difficult. Identification of generated papers and reports in the field of education has some differences compared to other areas, mainly due to the specificity and contextual nature of educational tasks. Currently there are many instruments developed for detecting generated text. A modern detector should be reliable in recognizing such texts to prevent the misuse of artificial intelligence.

Keywords: academic fraud, artificial intelligence, large language models, plagiarism detectors, text generation.

Відомості про автора / Information about the Author

Катерина Антіпова, PhD, доцент (б. в. з.) кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: antipova.katerina@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-9012-5290>.

Kateryna Antipova, PhD, Associate Professor (w/o a. d.) of the Department of Software Engineering, Faculty of Computer Sciences, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: antipova.katerina@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-9012-5290>.

DOI 10.34132/mspc2025.01.14.02

*Андрій Блиндарук
Олена Шаповалова*

END-TO-END ДИФЕРЕНЦІЙОВАНИЙ ЛАНЦЮГ NURBS – GEOM-GNN – TGN ДЛЯ ПРОСТОРОВО-ЧАСОВОЇ ІДЕНТИФІКАЦІЇ РУХОМИХ ОБ'ЄКТІВ

Запропоновано комплексний підхід до ідентифікації рухомих об'єктів, що поєднує параметрично-гладку репрезентацію контурів раціональними NURBS-кривими (Non-Uniform Rational B-Splines), геометрично-зважену обробку графів у Geom-GNN та часово-подієве моделювання пам'яті у TGN.

Неперервні параметри кривої перетворюються в компактний просторовий граф, де вершини відповідають контрольним точкам, а ребра – сегментам та k -найближчим з'єднанням. Geom-GNN агрегує повідомлення з урахуванням відстаней, а Temporal Graph Network (TGN) підтримує рекурентну пам'ять, забезпечуючи стійку ідентифікацію під час руху об'єкта траєкторією.

Ключові слова: *штучний інтелект, нейронні мережі, NURBS, GNN, TGN, ідентифікація об'єктів, просторово-часові дані.*

Точна ідентифікація об'єктів, що рухаються, є критичною для комп'ютерного зору, робототехніки та розширеної реальності. Більшість сучасних трекерів використовує піксельні дескриптори, point-cloud тощо, які мають високу пам'ятну та обчислювальну складність і погано відтворюють гладку геометрію. Раціональні B-сплайни (NURBS) забезпечують компактний опис контуру з гарантією безперервності похідних, а графові нейронні мережі (GNN) природно працюють із топологічними структурами [1].

Окремо Geom-GNN демонструють високі показники на статичних 3-D об'єктах, тоді як TGN ефективна на потокових даних подій [2]. Однак, дослідження для об'єднання цих методів у єдиний градієнтно-спроможний ланцюг, що враховує і геометрію, і довгу часову залежність не є розвинутих досі.

Наявні роботи та дослідження, зосереджуються переважно на статичній формі об'єктів, без розв'язання проблеми оновлення геометрії в реальному часі. Роботи з динамічними графами, зокрема Temporal Graph Network, навпаки оперують подієвими потоками, але базуються на «плоских» дескрипторах, що не містять детальної просторової інформації про криву. Тобто на звичайних векторних ознаках, які описують об'єкт лише як набір чисел без внутрішньої просторової структури. Наприклад, кольорові гістограми, середні значення пікселів чи 2-D координати рамки. Такі дескриптори називаємо «плоскі» за наступних причин:

- 1) Не містять геометрії контуру чи форми — немає інформації про кривини, взаємне розташування точок або топологію.
- 2) Не зберігають зв'язки між елементами — кожна ознака розглядається ізольовано, без контексту сусідніх точок чи сегментів.
- 3) Обмежуються фіксованим вектором — зміна роздільної здатності або деталізації не відбивається у структурі ознак.

Відсутність єдиної системи для ідентифікації об'єктів при просторово-часовому спостереженні пояснюється щонайменше трьома наступними викликами:

- 1) Диференційоване перетворення геометричної інформації про об'єкт спостереження в граф: адаптивне додавання або видалення контрольних точок для кожного кадру створює змінну топологію, що руйнує сталість розмірності. Тобто має місце обмеження фіксованою сіткою.
- 2) Узгодження геометричних і часових ознак: геометричні атрибути (кривина, орієнтація) живуть у різних масштабах та одиницях виміру порівняно з часовими інтервалами. Без спеціального нормування та механізму уваги ці домени «конкурують», що знижує якість повідомлень.
- 3) Керування пам'яттю при високій частоті подій: інтеграція потоків даних вже при швидкості по 30–60 FPS веде до стрімкого росту таблиці пам'яті в мережі TGN. Питання ефективного «забування» старих станів або їх-агрегації залишається відкритим.

Таким чином, запропонований у роботі диференційований ланцюг NURBS – Geom-GNN – TGN закриває прогалину між двома окремими ділянками досліджень: високоточною геометричною обробкою та динамічним графовим моделюванням [3]. Підхід демонструє, що шляхом об'єднання диференційованої вихідної обробки кривих, геометрично-зважених згорток і рекурентної пам'яті можна досягти універсальної, масштабованої та стійкої системи для просторово-часової ідентифікації рухомих об'єктів.

На етапі попередньої обробки при відео спостереженні кожен силует об'єкта у відеокадрі вилучається за допомогою стандартних методів сегментації (наприклад, U2-Net або Canny).

Отриманий контур апроксимується раціональною B-сплайн кривою фіксованого ступеня. Кількість контрольних точок нормується до верхньої межі: якщо точок замало, застосовується вставка вузлів, якщо забагато — зайві точки усуваються шляхом злиття збереженням геометричної похибки в межах заздалегідь заданого порогу. У результаті формується компактний вектор параметрів, інваріантний до роздільної здатності вихідного зображення та масштабу дискретизації.

Контрольні точки кривої розглядаються як вершини графа, а їх послідовні з'єднання — як ребра. Для збереження локальної геометрії додаються додаткові ребра між декількома найближчими за відстанню вершинами.

Кожна вершина збагачується атрибутами: тривимірними координатами, ваговим коефіцієнтом, оцінкою локальної кривини, кутом дотичної тощо. Таким чином граф зберігає усю форму кривої у дискретному, але топологічно зв'язному вигляді й займає на порядок менше пам'яті, ніж point-cloud або воксельне представлення.

На побудованому графі виконується багатопарова згортка Geom-GNN [4]. На відміну від класичних графових згорток, кожне повідомлення між вершинами модулюється наборами факторів, наприклад, евклідовою відстанню між їхніми координатами та кутковою різницею між напрямками нормалей.

Такий механізм дозволяє мережі диференціювати близьких за формою сусідів від далеких та ефективно працювати на графах із низькою гомофілією, де суміжні вершини можуть мати різні семантичні мітки. Тобто, в контексті графів, що розглядаються, ступінь, до якого взаємопов'язані вершини, мають схожі характеристики або однакові мітки (класи).

Щоб перетворити послідовність незалежних графів на динамічну структуру, здійснюється soft-matching контрольних точок між сусідніми кадрами.

Для кожної пари потенційно відповідних вершин створюється подія, що містить індекси вершин попереднього та поточного кадру, часову мітку та опис змін геометричних атрибутів (зсув центру, різницю кривини, зміну ваги).

Набір таких подій задає часовий граф, де ребра існують лише тоді, коли між вершинами відбувається взаємодія у часі.

Таким чином soft-matching не лише визначає попарні відповідності, а й формує багату подієву семантику, що кодує як геометричні, так і часові зміни, перетворюючи набір статичних NURBS-графів на єдиний потоковий просторово-часовий граф (рис. 1), придатний для обробки TGN без втрати градієнтної цілісності - властивості обчислювального ланцюга, при якому всі перетворення даних залишаються диференційованими:

1. Параметри NURBS, геометрично-зважені операції Geom-GNN і модуль пам'яті TGN утворюють єдину обчислювальну графову схему, де всі кроки (у тому числі soft-matching, нормування ознак, обчислення ваг повідомлень тощо) реалізовані диференційованими функціями.
2. Зворотне поширення помилки (back-propagation) може безперешкодно проходити крізь усі вузли цього графа, корегуючи як ваги нейронних шарів, так і параметри, що впливають на зіставлення контрольних точок або формування подій.
3. Не потрібно окремих кроків (наприклад, кластеризації), які переривали б ланцюг градієнта й змушували б використовувати два роздільні етапи — спершу препроцесинг, потім навчання нейромережі.

TGN підтримує окремий вектор пам'яті для кожної вершини просторово-часового графа [5]. Коли надходить нова подія, генерується повідомлення, що оновлює пам'ять як відправника, так і отримувача, із врахуванням давності події. Таким чином TGN зберігає довгострокову історію руху об'єкта, роблячи ідентифікацію стійкою до короточасних оклюзій, злиття силуетів або пропусків кадрів.

Для сцен із високою деталізацією передбачено багаторівневу схему: на грубому рівні крива має мінімальну кількість контрольних точок і використовується для швидкого пошуку відповідності; на тонкому — підвищена щільність точок додається лише для частини кадрів, де потрібна висока точність.

Інтегрована модель NURBS – Geom-GNN – TGN демонструє, що поєднання гладкої геометричної репрезентації, геометрично-зважених графових операцій та часової пам'яті є життєздатним шляхом до високоточних систем відстеження та ідентифікації об'єктів [6].

Розвиток дослідження можливий у напрямі автоматичного регулювання розмірності графа залежно від складності сцени, а також у залученні наборів даних для покращення здатності на унікальні траєкторії та деформовані рухомі об'єкти [7].

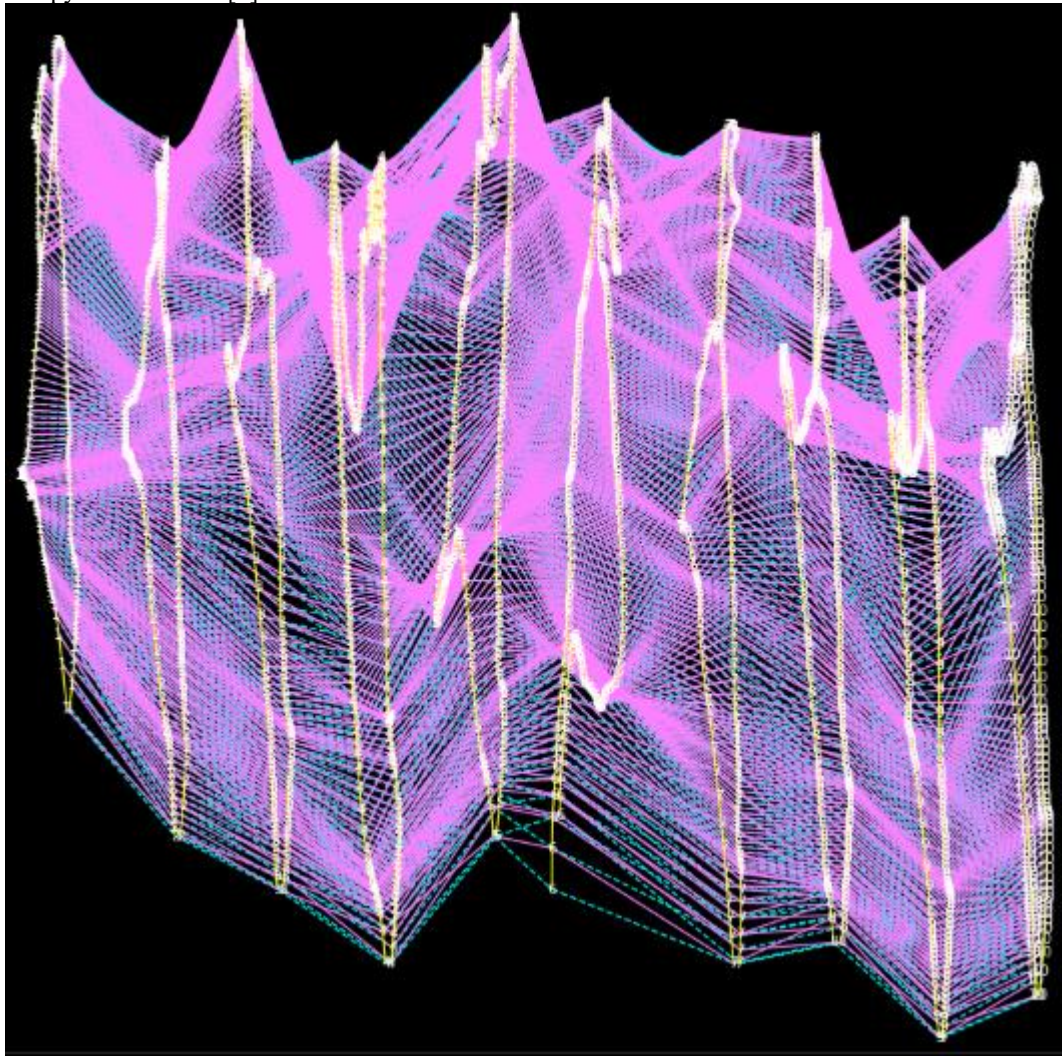


Рис. 1. Приклад просторово-часового графа, що побудовано на послідовних кадрах рухомого об'єкту.

Запропонований end-to-end диференційований ланцюг демонструє, що поєднання гладкої параметризації NURBS, геометрично-зваженої агрегації Geom-GNN та довго-часової пам'яті TGN забезпечує стійку ідентифікацію рухомих об'єктів при значному скороченні обчислювальних витрат.

Подальші напрями роботи включають адаптивне керування розмірністю графа в реальному часі та інтеграцію симулятивних даних для кращого узагальнення на унікальні траєкторії.

References

1. Fan J., Gholami B., Bäck T., Wang H. (2024) *NeuroNURBS: Learning Efficient Surface Representations for 3D Solids*. <https://arxiv.org/html/2411.10848v1> [in English].
2. Rossi E., Chamberlain B., Frasca F. (2020) *Temporal Graph Networks for Deep Learning on Dynamic Graphs*. <https://arxiv.org/abs/2006.10637> [in English].
3. A.O. Blyndaruk, O.O. Shapovalova. (2024), Poiednannia metodiv GNN ta NURBS dlia identyfikatsii rukhomykh ob'ektiv. *Systemy obrobky informatsii*, № 1 (176). С. 7-11. DOI: 10.30748/soi.2024.176.01 [in Ukrainian].
4. Garcia Satorras V., Hoogeboom E., Welling M. (2021) *E(n) Equivariant Graph Neural Networks*. <https://arxiv.org/abs/2102.09844> [in English].

5. Zhang Y., Liang Y., Leng J., Wang Z. (2024) *SCGTracker: Spatio-Temporal Correlation and Graph Neural Networks for Multiple Object Tracking*. *Pattern Recognition*, 110249.
<https://www.sciencedirect.com/science/article/abs/pii/S0031320323009469?via%3Dihub> [in English].
6. Lee J., Yeo C., Cheon S.-U., Park J. H., Mun D. H. (2023) *BRepGAT: Graph Neural Network to Segment Machining Feature Faces in a B-rep Model*. *Journal of Computational Design and Engineering*, 10 (6), 2384–2400.
<https://academic.oup.com/jcde/article/10/6/2384/7453688?login=false> [in English].
7. Blyndaruk A.O, Shapovalova O.O. (2024), Dosiahnennia ta vyklyky u sferi vidstezhennia rukhomykh ob'ektiv ta analizu yikhoi povedinky. *Molodizhnyi ekonomichnyi visnyk KhNEU im. S. Kuznetsia*, № 1, S.17-21 [in Ukrainian].

DOI 10.34132/mspc2025.01.14.02

Andrii Blyndaruk
Olena Shapovalova

END-TO-END DIFFERENTIABLE NURBS – GEOM-GNN – TGN PIPELINE FOR SPATIOTEMPORAL IDENTIFICATION OF MOVING OBJECTS

A comprehensive approach to moving-object identification is proposed, combining parametric smooth contour representation with rational NURBS curves (Non-Uniform Rational B-Splines), geometry-aware graph processing in Geom-GNN, and event-based temporal memory modeling in TGN.

Continuous curve parameters are transformed into a compact spatial graph whose vertices correspond to control points, while edges capture segment connectivity and k -nearest relations. Geom-GNN aggregates messages with distance weighting, and a Temporal Graph Network (TGN) maintains recurrent memory, enabling robust identification as an object traverses its trajectory.

Keywords: artificial intelligence; neural networks; NURBS; GNN; TGN; object identification; spatiotemporal data.

Відомості про автора / Information about the Author

Блиндарук Андрій Олександрович, аспірант з комп'ютерних наук кафедри кібербезпеки та інформаційних технологій Харківського національного економічного університету, ім. Семена Кузнеця, м. Харків, Україна, <https://orcid.org/0009-0009-8596-4020>

Andrii Blyndaruk, Ph.D. student in Computer Science of the Cybersecurity and Information Technologies Department, Kharkiv National University of Economics named after Simon Kuznets, Kharkiv, Ukraine, <https://orcid.org/0009-0009-8596-4020>

Шаповалова Олена Олександрівна, кандидат технічних наук, доцент Харківського національного економічного університету, ім. Семена Кузнеця, м. Харків, Україна, <https://orcid.org/0000-0003-4566-6634>

Olena Shapovalova, Assis. professor, PhD, Kharkiv National University of Economics named after Simon Kuznets, Kharkiv, Ukraine, <https://orcid.org/0000-0003-4566-6634>

© А. Блиндарук, О. Шаповалова

Стаття отримана редакцією 19.05.2025.

DOI 10.34132/mspc2025.01.14.03

*Надія Болюбаиш,
Володимир Блинцов*

ІНФОРМАЦІЙНА СИСТЕМА АНАЛІЗУ ТА МОНІТОРИНГУ ПРОДУКТИВНОСТІ ВЕБРЕСУРСІВ

Тези представляють опис розробленої інформаційної системи, яка на основі CLI-застосунку здійснює моніторинг продуктивності вебресурсів, відслідковуючи широкий спектр ключових показників Performance, SEO і Best Practices аудиту та технічного аналізу, надаючи рекомендації з їх покращення..

Ключові слова: оцінка продуктивності, управління продуктивністю застосунків, SEO аудит, технічний аналіз.

В умовах високих темпів цифровізації безперебійна робота вебзастосунків стала критично важливим фактором для успішного функціонування в усіх сферах суспільства. Сучасні цифрові платформи виступають провідними каналами комунікації і взаємодії, через які здійснюється отримання інформації, знайомство з брендами, здійснення покупок, просування товарів на ринку, спілкування зі співробітниками, партнерами та клієнтами тощо. Тому ключем для успіху вебресурсів, які розробляють ІТ-компанії, є їх висока продуктивність на різних платформах у будь-якій точці земної кулі.

На сьогоднішній день існує багато сервісів, які перевіряють продуктивність вебзастосунків. До найбільш поширених відносять SE Ranking, WebPageTest, PageSpeed Insights, GTmetrix, Pingdom Tools. Проте більшість існуючих інструментів не враховують широкий спектр критично важливих метрик, що робить їх недостатньо ефективними для комплексного аналізу та оптимізації продуктивності вебзастосунків. Проведений аналіз показав, що вебінструменти для аналізу продуктивності мають низку обмежень: обмежений доступ до метрик, відсутність глибоких рекомендацій, низьку гнучкість налаштувань, використання сторонніх API та обмежену кількість безкоштовних аудитів, що обмежує їх ефективне використання. Це обумовлює необхідність розробки проєкту, який усуває недоліки вже існуючих вебінструментів, об'єднуючи широкий спектр критично важливих метрик, пропонуючи конкретні рекомендації для їх покращення, підтримуючи запуск аналізу продуктивності вебсторінки з різних географічних локацій, а також можливість швидкого та зручного аналізу в універсальному безкоштовному форматі.

Розроблена інформаційна система аналізу та моніторингу продуктивності вебресурсів базується на основі розробленого консольного CLI-застосунку (англ. Command Line Interface, CLI), як складової частини АРМ-систем. АРМ-системи є спеціалізованими системами управління продуктивністю вебзастосунків (англ. Application Performance Management, АРМ). Їх впровадження забезпечує моніторинг показників продуктивності вебресурсів, проводить діагностику проблем й надає рекомендації з їх розв'язання. Сучасні АРМ-рішення охоплюють численні аспекти роботи програмного забезпечення, що дозволяє розглядати їх як універсальні інструменти для підтримки високої продуктивності. Проте, незважаючи на наявність великої кількості інструментів для аналізу продуктивності вебсторінок, більшість з них мають суттєві обмеження. Зокрема, вони покладаються на синтетичні тести, які не завжди відображають реальний досвід користувачів. В умовах, коли вебресурси стають дедалі складнішими, із великою кількістю зовнішніх залежностей та інтеграцій, традиційні інструменти часто не здатні адекватно проаналізувати вебзастосунок [1]. Крім того, вебінструменти можуть вимагати значного часу на обробку даних, оскільки результати часто генеруються з урахуванням великої кількості факторів, що ускладнює швидке прийняття рішень.

Розробка саме CLI-застосунку пропонує цілу низку важливих переваг: високу швидкість роботи, незалежність від зовнішніх серверів, можливість повної автоматизації процесу аналізу та максимальну адаптацію під специфіку конкретного середовища. На відміну від традиційних веборієнтованих АРМ-рішень, CLI-інструмент не потребує мережевого підключення до сторонніх сервісів для обробки й візуалізації даних. Це дозволяє

уникнути затримок, пов'язаних із передачею інформації через мережу, що особливо актуально в умовах високого навантаження або обмеженої пропускної здатності каналу.

Такий підхід дає змогу інтегрувати інструмент безпосередньо в CI/CD (англ. Continuous Integration / Continuous Delivery), забезпечуючи автоматичний запуск аналізу та моніторинг стану системи в реальному часі. Технологія неперервної інтеграції та доставки CI/CD, як технологія автоматизації тестування та доставки нових модулів проєкту, що розробляється, зацікавленим сторонам: розробникам, аналітикам, інженерам якості та кінцевим користувачам, дозволяє виявляти критично низькі показники продуктивності ще до релізу продукту в продакшн, що знижує ризики та вартість усунення помилок на пізніших етапах життєвого циклу програмного забезпечення (англ. Software Development Life Cycle, SDLC) [2].

Для виявлення методів оцінки продуктивності вебзастосунків та основних факторів, які впливають на продуктивність вебресурсів і ключових показників, що відображають ефективність конкретних оптимізаційних заходів, було здійснено аналіз досліджень науковців та практиків у цій сфері. Проведений аналіз дозволив виявити широкий спектр критично важливих метрик, таких як Performance, SEO, Best Practices та технічні аналіз, які було реалізовано у розробленій інформаційній системі.

До ключових метрик Performance, які вимірюють реальний досвід користувача стосовно швидкості завантаження контенту вебресурсу, інтерактивності та візуальної стабільності сторінки, відносять [3]:

- TTFB (Time to First Byte) – час, необхідний для отримання першого байта відповіді сервера;
- FCP (First Contentful Paint) – момент, коли браузер вперше відображає будь-який вміст (текст, зображення тощо);
- LCP (Largest Contentful Paint) – час відображення найбільшого за розміром елемента сторінки;
- CLS (Cumulative Layout Shift) – стабільність макета (сукупний зсув), що показує, наскільки змінюється положення елементів при завантаженні;
- INP (Interaction to Next Paint) – час відгуку на взаємодію користувача до моменту візуального оновлення.

При визначенні ключових метрик враховують порогові значення, які визначають зони:

- *зелена зона*: показники в межах оптимальних значень, що свідчить про те, що вебресурс функціонує на високому рівні;
- *жовта зона*: показники, які виходять за межі норми, але ще не критично, що є сигналом для команди розробників звернути увагу на показники системи з метою оптимізації, адже вони можуть перетворитися на серйозні проблеми з ефективністю вебресурсу;
- *червона зона*: показники перебувають на небезпечному рівні, який негативно впливає на працездатність вебресурсу та користувацький досвід, що може призводити до високого відсотка відмов від користування ресурсом, зниження задоволеності користувачів і, як наслідок, погіршення позицій у пошукових системах.

Для кожної метрики в рамках порогових значень визначаються відповідні штрафи, які залежать від ступеня перевищення встановлених порогів. Оцінка продуктивності вебресурсу базується на загальній сумі штрафів, яка нараховується відповідно до певних метрик. Початкова оцінка складає 100 балів, із якої відраховуються штрафи за кожну метрику, що перевищує встановлені порогові значення. Загальна формула для оцінки продуктивності (англ. Performance Score) вебресурсу S_{perf} є наступною:

$$S_{perf} = \max(100 - D), \quad (1)$$

де D – загальний штраф усіх метрик.

У контексті підвищення ефективності роботи вебресурсу не менш важливе значення має SEO (англ. Search Engine Optimization) оптимізація, яка охоплює як органічний, так і платний пошуковий трафік. Результати дослідження впливу SEO факторів на залучення користувачів, свідчать про те, що ефективна оптимізація вебсайту для пошукових систем значно збільшує його видимість, що, в свою чергу, веде до зростання органічного трафіку [4]. Оцінка ефективної SEO-стратегії базується на елементах, які допомагають пошуковим системам краще зрозуміти вміст сторінки і правильно її ранжувати для пошукових ботів із метою органічного просування [5]:

- *title*: головний заголовок сторінки, який має бути унікальним, точним і включати основні ключові слова вебзастосунку;
- *description*: короткий опис сторінки, який теж відображається в пошуку і його добре продумана реалізація може збільшити кількість переходів на сайт;
- *H1 tag*: основний заголовок на самій сторінці, який має відображати її головну тему та допомагає пошуковим системам визначити, про що саме йде мова на вебсторінці;
- *alt image attribute*: альтернативний текст для зображень є важливим як для SEO (описує зміст зображення), так і для доступності сайту для людей із вадами зору (англ. accessibility);
- *canonical tag*: тег, що показує пошуковим системам, яка версія сторінки є основною, щоб уникнути дублювання контенту;

- *crawable*: означає, що сторінка відкрита для індексації і доступна для сканування пошуковими ботами (якщо доступ обмежений, сторінка може не потрапити в результати пошуку).

Для кожного SEO-компонента вказують вагу, яка відображає його значущість, та валідність, яка може приймати значення 0, 0,5 та 1 і відображає відповідність компонента оптимальним налаштуванням. Тоді розрахунок комплексного показника SEO-оцінки здійснюється за формулою:

$$S_{SEO} = \left(\frac{\sum_{i=1}^k w_i \cdot m_i}{\sum_{i=1}^k w_i} \right) \cdot 100, \quad (2)$$

де i – індекс SEO-компонента із множини вище перерахованих; w_i – значення ваги i -го компонента; m_i – значення показника валідності i -го компонента; k – кількість SEO компонент.

Установлено, що ефективність бізнес-процесів у вебзастосунках визначається також дотриманням цілого комплексу стандартів і рекомендацій, відомих як Best Practices. Вони охоплюють критично важливі аспекти – доступність контенту, продуктивність роботи (на рівні як frontend, так і backend), безпеку даних та архітектурну стійкість системи. Їх підтримка та контроль є надійним фундаментом для забезпечення окупності бізнесу. До основних показників Best Practices вебресурсів відносять оцінку доступності, геолокації та безпеки, використання cookies та застарілих API, коректне співвідношення розмірів зображень, оцінку налаштувань метатегів та кодування символів, декларацію типу контенту [6]. Агрегована оцінка Best Practices вебресурсу здійснюється за формулою:

$$S_{BP} = \max(100 - \sum_{i=1}^n w_i (1 - m_i), \quad (3)$$

де w_i – вага, m_i – коефіцієнт валідності аудиторського елемента.

Вага аудиторського елемента відображає його важливість, а коефіцієнт валідності може приймати значення 0, 0,5 та 1 і оцінює коректність елемента.

Виявлено, що продуктивність не обмежується лише швидкістю завантаження, а включає й якість реалізації frontend архітектури, безпеки та доступності [7]. Це обумовлює необхідність здійснення технічного аналізу, який передбачає моніторинг таких показників: перевірка ефективності кешування та обробки даних: дозволяє оптимізувати витрати ресурсів і підвищити швидкодію застосунку; валідація HTML DOM структури у випадку інтеграції з веб-компонентами або генерацією звітів: забезпечує коректність візуалізації та сумісність із різними платформами; оптимізація роботи із зображеннями та медіа-ресурсами, що критично важливо для швидкодії та адаптивності вебчастини системи; перевірка адаптивності дизайну: дозволяє гарантувати коректне відображення інтерфейсів на різноманітних пристроях та розширеннях екранів; аналіз логування, обробки помилок та безпеки: допомагає мінімізувати ризики вразливостей і втрати даних.

Аналіз цих параметрів дає змогу виявляти технічні недоліки як на ранніх етапах розробки, так і під час впровадження та експлуатації, забезпечуючи стабільну, швидку та безпечну роботу вебресурсу з урахуванням потреб усіх категорій користувачів.

Для розробки інформаційної системи аналізу та моніторингу продуктивності вебресурсів було використано мову програмування JavaScript та середовище розробки VS Code. Для реалізації серверної частини було обрано Node.js. Бібліотека Puppeteer застосовувалася для керування браузером через Node.js, забезпечуючи програмний доступ до функціональності браузера та автоматизацію доступу до вебресурсів, збору даних про час завантаження, структуру DOM, аналіз SEO-метрик і відповідність Best Practices. Це дозволило отримувати комплексну аналітику сторінок без необхідності ручного втручання. Для створення ядра Node.js застосунку використано бібліотеку Yargs, яка дозволила організовувати обробку аргументів командного рядка. Бібліотека Axios використовувалася для здійснення асинхронних HTTP-запитів. Застосування бібліотеки cli-spinner забезпечило візуалізацію результатів роботи системи.

Для покращення візуального відображення результатів аудиту було реалізовано інтеграцію розробленого CLI-застосунку до вебзастосунку інформаційної системи. На рисунку 1 показано, як виводиться загальна оцінка аудиту вебресурс.

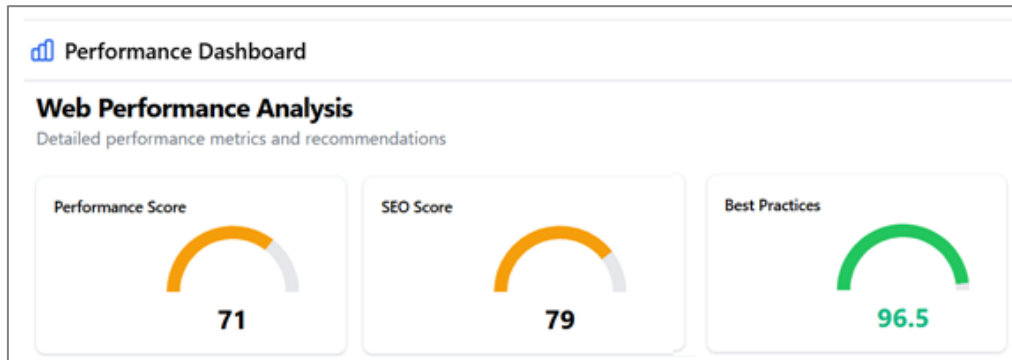


Рис. 1. Загальна оцінка аудиту вебресурсу.

На рисунку 2 відображено результати оцінки метрик завантаження сторінки Performance на різних платформах.

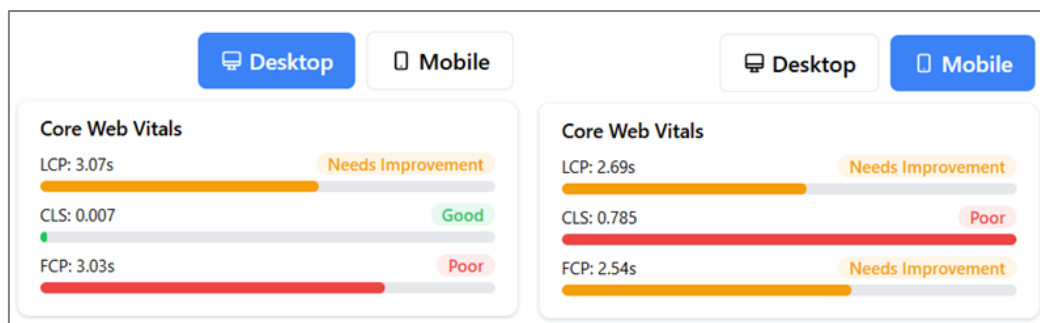


Рис. 2. Оцінка показників Performance вебресурсу.

Окрім графічних елементів, у системі передбачено виведення звіту по технічному аналізу, SEO і Best Practices аудиту та рекомендацій по забезпеченню оптимальної продуктивності й ефективності роботи вебресурсу. На рисунку 3 наведено приклад результатів оцінки метрик SEO та рекомендацій по їх покращенню.

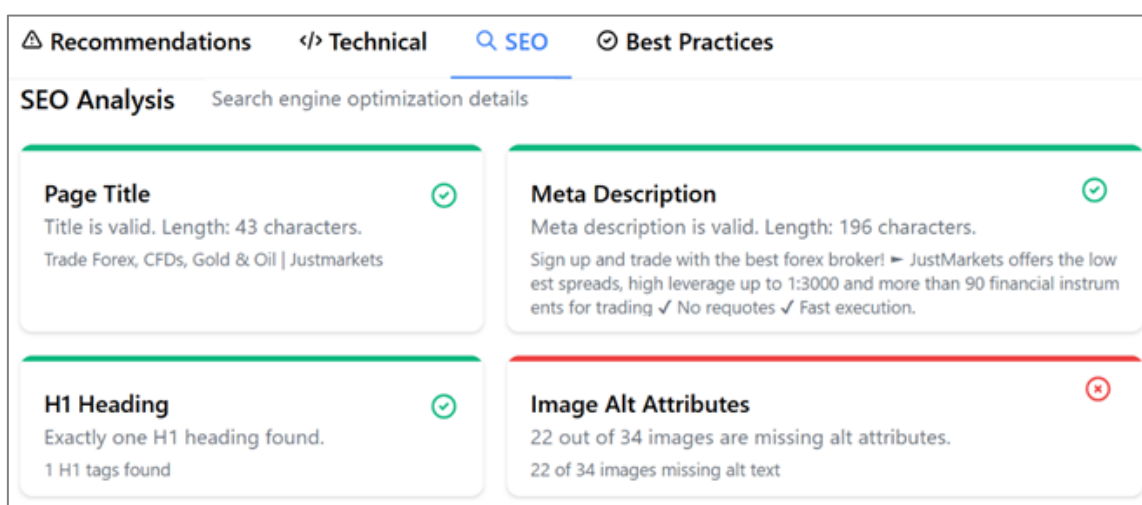


Рис. 3. Оцінка та рекомендації по оптимізації показників SEO.

У таблиці 1 наведено показники якості розробленого CLI-застосунку у порівнянні з іншими програмними засобами, які оцінюють продуктивність вебресурсів. Порівняльний аналіз демонструє, що розроблений

інструмент не лише відповідає, але й перевершує функціональність багатьох популярних рішень, які представлені на ринку.

Таблиця 1.

Порівняльний аналіз розробленого CLI-застосунку з популярними сервісами аналізу продуктивності вебресурсів

Показники аналізу	Поширені ПЗ					CLI (розроблений)
	SE Ranking	Web Page Test	PageSpeed Insights	GTmetrix	Pingdom Tools	
Performance	+	+	+	+	+	+
SEO	+	–	+	–	–	+
Best Practise	–	+	+	–	–	+
Suggestions	–	–	+	–	–	+
GEO runs	–	+ (обмеж.)	–	+ (обмеж.)	+ (обмеж.)	+ (без обмеж.)
3 rd API for run	+	+	own API	+	own API	own API
Free or Paid	Free plan, but paid	Free plan, but paid	Free	Free plan, but paid	Free plan, but paid	Free
Audit report	–	+	+	+	+	+
Час аналізу, с	~20	~35-40	~20	~25	~10-15	~15

Таким чином, розроблена інформаційна система на основі CLI-застосунку для локального аналізу продуктивності вебресурсів є перспективним і стратегічно обґрунтованим рішенням. Її застосування дозволяє проводити комплексний аналіз та моніторинг продуктивності, бізнес параметрів, SEO-оптимізації і технічного стану вебресурсів, орієнтований на швидке, автономне тестування у реальному середовищі без потреби у сторонніх онлайн-сервісах. Що дає змогу оперативно виявляти проблемні місця та недоліки на усіх етапах життєвого циклу вебресурсу із врахуванням геолокації та потреб усіх користувачів.

References

1. Xilogianni, C., Doukas, F., & Drivas, I. (2022). Speed Matters: What to prioritize in Optimization for Faster Websites. *Analytics*, 1(2), (pp/ 175-192). Retrieved from: <https://doi.org/10.3390/analytics1020012>
2. Ale, N.K. (2020). Integration Performance Testing into CI/CD Pipelines for Test Automation. *Journal of Scientific and Engineering Research*, 7(6), (pp. 273-278). Retrieved from: <https://jsaer.com/download/vol-7-iss-6-2020/JSAER2020-7-6-272-278.pdf>
3. Google Documentation. Understanding Core Web Vitals and Google search result. Retrieved from: <https://developers.google.com/search/docs/appearance/core-web-vitals>.
4. Bansal, D. (2024). How SEO Makes Website Loads Faster and Helps in User Engagement. *International Journal for Multidisciplinary Research*, 6(2). Retrieved from: <https://www.ijfmr.com/papers/2024/2/15291.pdf>
5. Coursaris, C., Osch, W., Lores-Nicolas, C., Molina-Castillo F., & Rapp, N. (2013). Driving Website performance using Web Analytics: A Case Study. In *Proceedings of the 19th Americas Conference on Information Systems (AMCIS)*. Retrieved from: <https://aisel.aisnet.org/amcis2013/HumanComputerInteraction/RoundTablePresentations/9/>
6. Vepsalainen, J., Hellas, A., & Vuorimaa, P. (2024). Overview of Web Application Performance Optimization Techniques. Retrieved from: <https://arxiv.org/html/2412.07892v1>
7. Marang, A.Z. (2018). Analysis of web performance optimization and its impact on user experience. Retrieved from: <http://www.diva-portal.org/smash/get/diva2:1228341/FULLTEXT01.pdf>

INFORMATION SYSTEM FOR ANALYSIS AND MONITORING OF WEB RESOURCES PERFORMANCE

The abstracts present a description of the developed information system, which, based on a CLI application monitors the performance of web resources, tracking a wide range of key indicators of Performance, SEO and Best Practices of audit and technical analysis, providing recommendations for their improvement.

Key words: *performance score, application performance management (APM), SEO audit, technical analysis.*

Відомості про авторів / Information about the Authors

Надія Болюбаш, канд. пед. наук, доцент кафедри інтелектуальних інформаційних систем Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: Nadiya.Bolubash@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-2274-2422>.

Nadiia Boliubash, PhD in Education, Associate Professor at the Department of Intelligent Information Systems, Petro Mohyla Black Sea National University, Ukrainian. E-mail: Nadiya.Bolubash@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-2274-2422>.

Володимир Блинцов, студент факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: volodymyr.blyntsov@gmail.com.

Volodymyr Blyntsov, student at the Faculty of Computer Science, Petro Mohyla Black Sea National University, Ukrainian. E-mail: volodymyr.blyntsov@gmail.com.

DOI 10.34132/mspc2025.01.14.04

Стефанія Бондаренко
Микита Ковальчук

КІБЕРФІЗИЧНІ СИСТЕМИ В МОДЕЛЮВАННІ ТА ОРГАНІЗАЦІЇ КОМП'ЮТЕРНИХ СИСТЕМ ЛОГІСТИКИ

Тези аналізують цифровізацію як ключовий фактор розвитку цифрового суспільства та демократичного врядування в Україні. Цифровізація розглядається як інструмент модернізації державного управління, що базується на співпраці з громадянами, захисті суспільних інтересів та розвитку державно-приватних партнерств з бізнес-структурами. Особливу увагу приділено тому, як цифрові технології сприяють об'єднанню зусиль для досягнення цілей сталого розвитку країни. У контексті сталого розвитку підкреслюється важливість добровільної ініціативи як індивідуумів, так і держав, а також необхідність підвищення прозорості державного управління, що дозволяє зменшити можливості для фальсифікацій і забезпечити доступ до необхідної інформації для всіх зацікавлених сторін.

Ключові слова: цифровізація, цифрові технології, електронізація, інтернет речей, кіберфізичні системи, логістичні процеси, оптимізація, сталий розвиток.

Сучасна логістика вимагає застосування інноваційних підходів до управління складними системами, оскільки обсяги даних, що обробляються, постійно зростають, а попит і транспортні умови змінюються в режимі реального часу. У такому середовищі важливими інструментами стають кіберфізичні системи (КФС) (рис. 1), які пропонують інтеграцію фізичних об'єктів – транспортних засобів, складів і товарів – з кібернетичними компонентами, такими як алгоритми, датчики та програмне забезпечення. Це дозволяє створювати гнучкі та адаптивні комп'ютерні системи, здатні автоматично налаштовуватися під зміни умов, оптимізуючи процеси перевезень і управління ресурсами в реальному часі [1].

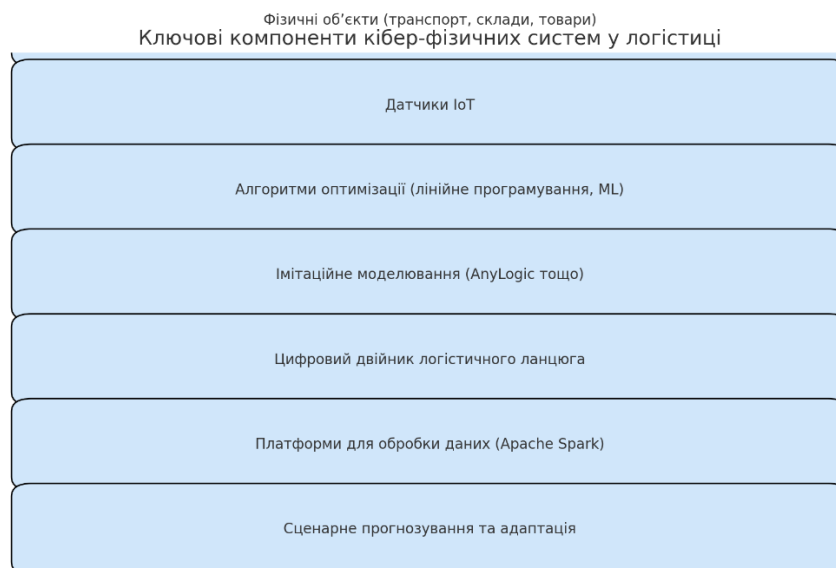


Рисунок 1 – Ключові компоненти кібер-фізичних систем у логістиці
Джерело: розроблено авторами

Однією з основних характеристик КФС у логістиці є синхронізація фізичного та віртуального світів. Цей процес реалізується через використання датчиків Інтернету речей (IoT), що встановлюються на транспортних засобах і в логістичних вузлах. Вони збирають дані про стан транспортної мережі, залишки товарів, запити та попит на різних етапах ланцюга постачання. Ці дані обробляються за допомогою потужних інструментів моделювання, таких як імітаційні платформи, наприклад, AnyLogic, і алгоритмів оптимізації, зокрема лінійного програмування та методів машинного навчання [2].

Цифровий двійник, що створюється за допомогою таких технологій, є віртуальним відображенням реального логістичного ланцюга. Цей двійник дозволяє здійснювати детальне прогнозування розвитку процесів, враховуючи випадкові або непередбачувані фактори, такі як затримки, зміни попиту чи коливання в доступності ресурсів. Завдяки можливості моделювання різних сценаріїв і варіантів розвитку ситуацій, КФС сприяють підвищенню ефективності управління ресурсами і скороченню витрат на доставку товарів.

Впровадження КФС у логістику також дає змогу реалізувати більш інтелектуальні системи управління, здатні не лише реагувати на зміни в реальному часі, а й самостійно оптимізувати маршрутні плани, контролювати запаси і передбачати проблеми до їх виникнення. Оскільки попит на швидку і ефективну доставку товарів тільки зростає, застосування таких інноваційних технологій дозволяє підтримувати високий рівень обслуговування клієнтів і забезпечувати безперебійну роботу логістичних процесів.

Математична основа транспортної задачі, яка включає мінімізацію цільової функції та обмеження на запаси і потреби, має тісний зв'язок із концепцією кіберфізичних систем (КФС) у контексті моделювання та організації комп'ютерних систем логістики. У такій системі цільова функція, яка мінімізує витрати на транспортування вантажів між постачальниками і споживачами, є важливим елементом для оптимізації логістичних процесів. КФС дозволяють створювати автоматизовані системи, які враховують не лише фізичні обмеження, такі як обсяги транспортування або географічні умови, але й інформацію, отриману через датчики та інші кібернетичні компоненти, що стежать за станом транспорту і процесами доставки.

Обмеження на запаси і потреби, що використовуються в транспортній задачі, є важливими для моделювання логістичних процесів у КФС. Дані, що постійно збираються за допомогою датчиків, дозволяють реальним системам, таким як автоматизовані склади або транспортні мережі, оперативно коригувати маршрути і планування, оптимізуючи потоки товарів. КФС інтегрують ці дані в автоматизовані алгоритми управління, що забезпечують доставку з урахуванням фізичних обмежень, наприклад, ємності складів або вантажопідйомності транспорту.

Таким чином, математичні моделі транспортної задачі стають основою для побудови ефективних автоматизованих логістичних систем, які є частиною кіберфізичних систем. Важливою є інтеграція фізичних та інформаційних компонентів, що дозволяє враховувати як реальні умови, так і математичні обмеження для досягнення мінімальних витрат і ефективного управління доставкою товарів.

Математична основа транспортної задачі виражається у вигляді мінімізації цільової функції:

$$\min Z = \sum_{i=1}^m \sum_{j=1}^n C_{ij} x_{ij} \quad (1)$$

– за умов обмежень на запаси:

$$\sum_{j=1}^n x_{ij} = a_i, \quad i = 1, 2, \dots, m \quad (2)$$

– за умов обмежень на потреби:

$$\sum_{i=1}^m x_{ij} = b_j, \quad j = 1, 2, \dots, n, \quad (3)$$

де

C_{ij} – вартість перевезення одиниці товару,

x_{ij} – кількість товару, що перевозиться з пункту i до пункту j ,

a_i – запаси у пункті відправлення,

b_j – потреби у пункті призначення.

Після визначення математичної основи транспортної задачі, яка включає мінімізацію цільової функції та обмеження на запаси і потреби, можна стверджувати, що ці формули є важливими інструментами для досягнення оптимізації в реальних кіберфізичних системах. Завдяки використанню таких математичних моделей, які відображають як економічні, так і фізичні обмеження, можна розробляти ефективні алгоритми для автоматизованого управління логістичними процесами.

Зокрема, цільова функція, яка мінімізує витрати на транспортування, є основою для оптимізації маршрутів та

розподілу товарів між різними пунктами. У реальних умовах, коли обсяги даних, що надходять від сенсорів і систем моніторингу, постійно зростають, ці математичні моделі допомагають зберегти високий рівень ефективності управління. Інтеграція таких моделей із кіберфізичними системами дозволяє швидко реагувати на зміни умов доставки, такі як затримки, зміни в попиту або доступності ресурсів.

Обмеження на запаси і потреби (формули (2) і (3)) визначають допустимі межі для кількості товару, який можна транспортувати з певних пунктів відправлення до пунктів призначення, і дозволяють адаптувати систему до змін у реальному часі. У контексті КФС ці обмеження можуть бути динамічно адаптовані в залежності від актуальних даних, що надходять із датчиків IoT, створюючи гнучкі системи, які враховують зміни у запасах або потребах на різних етапах ланцюга постачання.

У вирішенні задач застосування принципів кіберфізичних систем (КФС) до комп'ютерних систем логістики важливими є три основні етапи, які дозволяють досягти оптимізації та ефективності в управлінні логістичними процесами. Першим етапом є розробка структурно-алгоритмічних моделей, що поєднують фізичні об'єкти, такі як транспортні засоби та склади, з алгоритмами управління. Це дозволяє створювати моделі для оптимізації маршрутів доставки та розподілу ресурсів, враховуючи динамічні зміни в реальному часі, наприклад, зміни попиту, затримки чи зміни в транспортних умовах. За допомогою таких моделей можна розробити адаптивні системи управління, які забезпечують максимальну ефективність логістичних операцій навіть в умовах постійних змін.

Другим етапом є створення засобів опису та стандартизації компонентів, що взаємодіють у системі. Використання формальних специфікацій дозволяє стандартизувати взаємодію між різними елементами системи, такими як датчики IoT, транспортні засоби, склади і програмне забезпечення. Це є необхідним для забезпечення безперебійного збору даних та їх обміну між різними компонентами, що створює стабільну основу для подальшої обробки і аналізу інформації. Завдяки цьому, система стає здатною ефективно взаємодіяти з різними технологіями, забезпечуючи надійність і точність у зборі та обробці даних.

Третім етапом є моделювання логістичних процесів і аналіз різних сценаріїв. За допомогою імітаційного моделювання можна прогнозувати різні варіанти розвитку подій, враховуючи непередбачувані фактори [3], такі як зміни в попиту, затримки або коливання ресурсів. Моделювання дозволяє протестувати різні сценарії і вибрати найбільш оптимальний, що допомагає в прийнятті обґрунтованих рішень щодо управління запасами та транспортними потоками. Це дає можливість зменшити витрати, підвищити ефективність і адаптивність логістичних операцій до змін у реальному часі.

Практична реалізація цього підходу включає інтеграцію даних з Інтернету речей (IoT) з платформами для обробки великих даних, такими як Apache Spark, що дозволяє обробляти великі обсяги інформації, зібраної від сенсорів і пристроїв в реальному часі. Ці дані постійно оновлюються, що дає змогу адаптувати алгоритми управління до змін у логістичних умовах. Імітаційне моделювання дає змогу оцінити оптимальність маршрутів і розподілу товарів, а також прогнозувати потреби і змінювати стратегії управління залежно від ситуації. Це дозволяє зменшити витрати на перевезення та підвищити гнучкість логістичних операцій, забезпечуючи більш ефективне використання ресурсів.

Таким чином, інтеграція КФС, IoT-технологій та алгоритмів оптимізації дозволяє створити адаптивні системи, здатні ефективно управляти величезними обсягами даних і оперативно реагувати на зміни в умовах доставки. Це є важливим кроком до досягнення більш високої ефективності, гнучкості та сталості в управлінні логістичними процесами.

References

1. Lee J., Bagheri B., Kao H.-A. A Cyber-Physical Systems architecture for Industry 4.0-based manufacturing systems. *Manufacturing Letters*, 2015, Vol. 3, pp. 18–23. DOI: 10.1016/j.mfglet.2014.12.001.
2. Wang S., Wan J., Zhang D., Li D., Zhang C. Towards smart factory for Industry 4.0: A self-organized multi-agent system with big data based feedback and coordination. *Computer Networks*, 2016, Vol. 101, pp. 158–168. DOI: 10.1016/j.comnet.2015.12.017.
3. Monostori L. Cyber-physical production systems: Roots, expectations and R&D challenges. *Procedia CIRP*, 2014, Vol. 17, pp. 9–13. DOI: 10.1016/j.procir.2014.03.115.

CYBER-PHYSICAL SYSTEMS IN MODELING AND ORGANIZATION OF COMPUTER SYSTEMS IN LOGISTICS

The theses present a theoretical and methodological analysis of digitalization as a crucial factor in the development of a digital society and democratic governance in Ukraine. Digitalization is characterized as a tool for modernizing public administration, based on cooperation with citizens, protection of public interests, and fostering public-private partnerships with commercial business structures. The paper emphasizes how digital technologies have contributed to uniting efforts toward achieving the global goal of sustainable development. In the context of sustainable development, the importance of voluntary initiatives from individuals, states, and the global community is highlighted, as well as the need to increase the transparency of public administration, which reduces the possibilities for falsification and ensures access to necessary information for all stakeholders.

Keywords: digitalization, digital technologies, e-governance, Internet of Things, cyber-physical systems, logistics processes, optimization, sustainable development.

Відомості про автора / Information about the Author

Бондаренко Стефанія, викладачка кафедри інженерії програмного забезпечення, Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: bondarenko.stefania@chmnu.edu.ua, ORCID: <https://orcid.org/0009-0007-8767-1032>

Bondarenko Stefaniia, Lecturer, Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: bondarenko.stefania@chmnu.edu.ua, ORCID: <https://orcid.org/0009-0007-8767-1032>

Ковальчук Микита, здобувач третього рівня вищої освіти (PhD) зі спеціальності 123 «Комп'ютерна інженерія», Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: sakovalchuk99@chmnu.edu.ua, ORCID: <https://orcid.org/0009-0002-9214-8039>

Kovalchuk Mykyta PhD student, Specialty 123 «Computer Engineering», Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: sakovalchuk99@chmnu.edu.ua, ORCID: <https://orcid.org/0009-0002-9214-8039>

ВІЗУАЛІЗАЦІЯ ЧИСЕЛЬНИХ РОЗВ'ЯЗКІВ РІВНЯНЬ МАТЕМАТИЧНОЇ ФІЗИКИ З ВИКОРИСТАННЯМ MAPLE

У тезах розглянуто візуалізацію розв'язків деяких задач математичної фізики, отриманих чисельними методами із застосування математичного пакету Maple, таких як: змішана крайова задача для рівняння коливання струни, рівняння теплопровідності та задачі Діріхле для рівняння Лапласа. Хоча змішані крайові задачі мають певні складнощі, існує потужний арсенал чисельних методів та програмних інструментів, які дозволяють їх розв'язувати. Використання Maple для чисельного розв'язання задач математичної фізики дає змогу студентам-комп'ютерникам та інженерам сконцентруватися на суті проблеми, їм не доводиться витрачати час на складну реалізацію чисельних алгоритмів. Вивчення чисельних методів передбачає їхню програмну реалізацію, що сприяє розвитку навичок програмування, особливо в області роботи з масивами даних, розробки ефективних алгоритмів та оптимізації коду. Фахівці з глибоким розумінням чисельних методів є затребуваними у багатьох галузях, включаючи IT-компанії, саме тому при вивченні таких методів ми поєднуємо застосування математичного пакету Maple та програмної реалізації на різних мовах програмування.

Ключові слова: диференціальні рівняння з частинними похідними, чисельні методи, метод сіток, рівняння математичної фізики, Maple.

Рівняння математичної фізики – це клас диференціальних рівнянь з частинними похідними (ДРЧП), які описують різноманітні фізичні явища. Невідома функція $U = U(t, x, y, z)$ у таких рівняннях залежить від багатьох змінних (зазвичай це час t та просторові координати x, y, z). Найчастіше використовуються рівняння другого порядку, оскільки згідно із законами механіки перша похідна – швидкість, а друга – прискорення.

Окрім ДРЧП другого порядку математична модель задач математичної фізики містить ще додаткові умови: початкові та/або крайові. Ці умови потрібні для того, щоб виділити єдиний фізично змістовний розв'язок. У фізиці, механіці та інженерії часто зустрічаються змішані крайові задачі, в яких комбінують різні типи умов (початкові та крайові). Використання чисельних методів до розв'язування таких задач зумовлене багатьма вагомими причинами: вони подають наближений розв'язок у вигляді набору чисел у вузлах сітки, що легко піддається аналізу та візуалізації, на відміну від аналітичного розв'язку, який може мати вигляд нескінчених рядів чи інтегралів, крім того, іноді вони дозволяють отримувати наближені розв'язки для задач, які є абсолютно недоступними для аналітичних методів.

Існує багато потужних програмних пакетів, які реалізують чисельні методи та надають інструменти для візуалізації результатів. Одним із таких пакетів є Maple. Maple є лідером у знаходженні аналітичних розв'язків ДРЧП, він володіє великою бібліотекою відомих розв'язків ДРЧП та різними методами, такими як метод характеристик, використання симетрій Лі та інтегральних перетворень. У випадку неможливості відшукування аналітичного розв'язку у Maple можна знайти чисельні розв'язки, підключивши спеціальні пакети. Також Maple дає потужні інструменти для візуалізації отриманих розв'язків ДРЧП у вигляді графіків, поверхонь, що допомагає зрозуміти поведінку розв'язку та вплив крайових умов. Тому саме Maple з його потужними символічними та чисельними можливостями, а також з можливістю візуалізації, роблять його ефективним для використання студентами, інженерами та дослідниками при вивченні рівнянь математичної фізики [1, 2].

Найбільш поширеним методом розв'язання задач математичної фізики є метод сіток [3, 4]. Він є простим у розумінні та реалізації, ефективним для задач з простими геометричними областями. Сутність цього методу полягає в тому, що неперервну область, в якій шукають розв'язок задачі заміняють дискретною сіткою вузлів. Вузлами цієї сітки є точки $(x_i, t_j) = (x_0 + ih, t_0 + jk)$, де h, k – відповідно крок по осях Ox та Ot , $i, j = 0, 1, \dots$. Невідому функцію $U(x, t)$ шукають наближено у вузлах сітки, використовуючи позначення $U_{ij} = U(x_i, t_j)$. Частинні похідні у вузлах сітки апроксимуються скінчено-різницевиими співвідношеннями:

$$\begin{aligned} U_t &= \frac{U_{i,j+1} - U_{i,j}}{k}, & U_x &= \frac{U_{i+1,j} - U_{i,j}}{h}, \\ U_{tt} &= \frac{U_{i,j+1} - 2U_{i,j} + U_{i,j-1}}{k^2}, & U_{xx} &= \frac{U_{i+1,j} - 2U_{i,j} + U_{i-1,j}}{h^2}. \end{aligned} \quad (1)$$

Рівняння коливання струни. Розглянемо змішану крайову задачу для рівняння коливання струни:

$$U_{tt} = U_{xx}, \quad (2)$$

з граничними умовами:

$$U(0, t) = p(t), \quad U(L, t) = q(t), \quad (3)$$

та початковими умовами:

$$U(x, 0) = f(x), \quad U_t(x, 0) = g(x). \quad (4)$$

Розв'язок будемо шукати в області $D: \{0 \leq x \leq L, 0 \leq t \leq +\infty\}$. Початкові та крайові умови повинні бути узгодженими, тобто:

$$p(0) = f(0), \quad q(0) = f(L), \quad p'(0) = g(0), \quad q'(0) = g(L).$$

Використовуючи (1) замінимо задачу (2)-(4) її скінчено-різницевою аналогом:

$$\frac{U_{i,j+1} - 2U_{i,j} + U_{i,j-1}}{k^2} = \frac{U_{i+1,j} - 2U_{i,j} + U_{i-1,j}}{h^2}, \quad (5)$$

$$U_{0,j} = p_j, \quad U_{n,j} = q_j, \quad j = 0, 1, 2, \dots, \quad (6)$$

$$U_{i,0} = f_i, \quad U_{i,1} = U_{i,0} + kg_i, \quad i = 0, 1, 2, \dots \quad (7)$$

Умови (6)-(7) задають значення функції $U(x, t)$ на границі області D . Якщо позначити через $\lambda = \frac{k}{h}$, то рівняння (5) зведеться до вигляду

$$U_{i,j+1} = 2(1 - \lambda^2)U_{i,j} + \lambda^2(U_{i+1,j} + U_{i-1,j}) - U_{i,j-1}, \quad i = \overline{1, n-1}; j = 1, 2, \dots \quad (8)$$

Розв'язок різницевої задачі (5)-(7) рівномірно збігається з розв'язком крайової задачі (2)-(4) при $h \rightarrow 0, k \rightarrow 0$, якщо $\lambda < 1$, тобто $k < h$.

Розв'язок у Maple подається у вигляді таблиці значень шуканої функції $U(x, t)$ у вузлах сітки. Більш цікавим є отримання візуалізації розв'язків задачі таких як, форма струни в певний момент часу t або зміна $U(t)$ у фіксованих точках x :

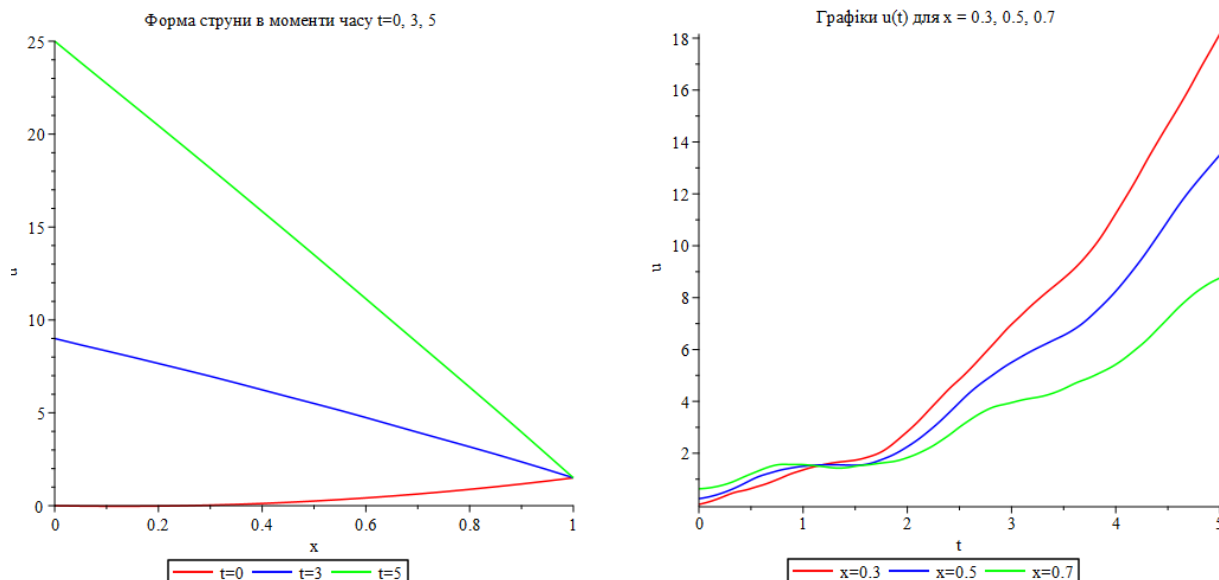


Рис. 1. Візуалізація розв'язків, отриманих для змішаної крайової задачі рівняння коливання струни в Maple.

Рівняння теплопровідності. Розглянемо ДРЧП параболічного типу на прикладі рівняння теплопровідності. Змішана крайова задача для такого рівняння має вигляд:

$$U_t = U_{xx}, \quad (9)$$

з граничними умовами:

$$U(0, t) = p(t), \quad U(L, t) = q(t), \quad (10)$$

та початковою умовою:

$$U(x, 0) = f(x). \quad (11)$$

Розв'язок будемо шукати в області $D: \{0 \leq x \leq L, 0 \leq t \leq +\infty\}$. Початкова та крайові умови повинні бути узгодженими, тобто:

$$p(0) = f(0), \quad q(0) = f(L).$$

Використовуючи (1) замінімо задачу (9)-(11) її скінчено-різницевою аналогом:

$$\frac{U_{i,j+1} - U_{i,j}}{k} = \frac{U_{i+1,j} - 2U_{i,j} + U_{i-1,j}}{h^2}, \quad (12)$$

$$U_{0,j} = p_j, \quad U_{n,j} = q_j, \quad j = 0, 1, 2, \dots, \quad U_{i,0} = f_i, \quad i = \overline{0, n}. \quad (13)$$

Умови (13) задають значення функції $U(x, t)$ на границі області D . Якщо позначити через $\lambda = \frac{k}{h^2}$, то рівняння (12) зведеться до вигляду

$$U_{i,j+1} = \lambda U_{i-1,j} + (1 - 2\lambda)U_{i,j} + \lambda U_{i+1,j}, \quad i = \overline{1, n-1}; j = 1, 2, \dots \quad (14)$$

Процес розв'язання збігається і є стійким, якщо $\lambda \leq 0,5$, тобто $k < \frac{h^2}{2}$.

У рівнянні (9) функція $U(x, t)$ задає температуру в будь-якій точці стержня довжиною L на момент часу t . Візуалізація розподілу температур в різні моменти часу, отримана в Maple та Phytон продемонстрована на рис. 2.

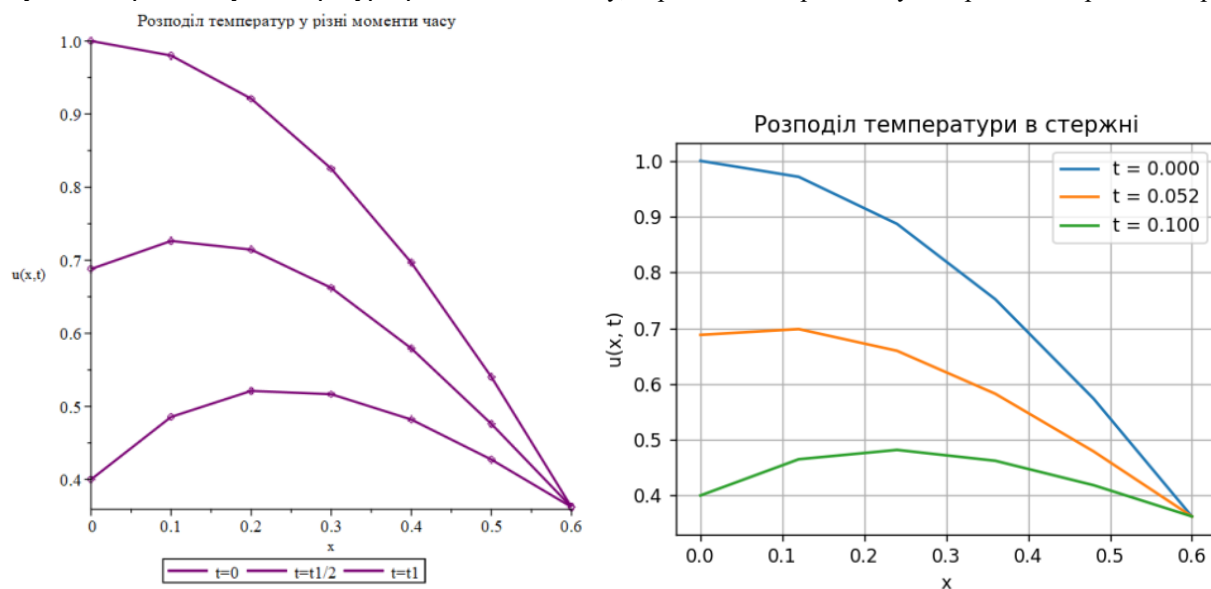


Рис. 2. Візуалізація розподілу температур в різні моменти часу, отримана в Maple та Phytон.

Задача Діріхле для рівняння Лапласа. Розглянемо рівняння еліптичного типу, до яких приводять дослідження стаціонарних процесів, коли шукана функція не залежить від часу. Прикладом такого рівняння є рівняння Лапласа. Розглянувши його з граничними умовами по координатах x, y , отримаємо першу крайову задачу або задачу Діріхле:

$$\Delta U = 0, \quad U(x, 0) = f(x), \quad U(x, b) = g(x), \quad U(0, y) = p(y), \quad U(a, y) = q(y). \quad (15)$$

Розв'язок будемо шукати в області $D: \{0 \leq x \leq a, 0 \leq y \leq b\}$.

Використовуючи (1) замінімо задачу (15) її скінчено-різницевою аналогом:

$$\frac{U_{i,j+1} - 2U_{i,j} + U_{i,j-1}}{k^2} + \frac{U_{i+1,j} - 2U_{i,j} + U_{i-1,j}}{h^2} = 0, \quad (16)$$

або

$$2(h^2 + k^2)U_{i,j} = k^2(U_{i+1,j} + U_{i-1,j}) + h^2(U_{i,j+1} + U_{i,j-1}), \quad i = \overline{1, n-1}; j = 1, 2, \dots \quad (17)$$

Знаходимо значення $U_{i,j}$, які відповідають граничним вузлам завдяки граничним умовам (15). Для всіх внутрішніх вузлів сітки записуємо рівняння (17) та отримаємо СЛАР з досить великою кількістю рівнянь та невідомих. Отриману систему алгебраїчних рівнянь можна розв'язати, використовуючи прямі або ітераційні методи [5].

Візуалізація отриманого розв'язку задачі (15) в Maple:

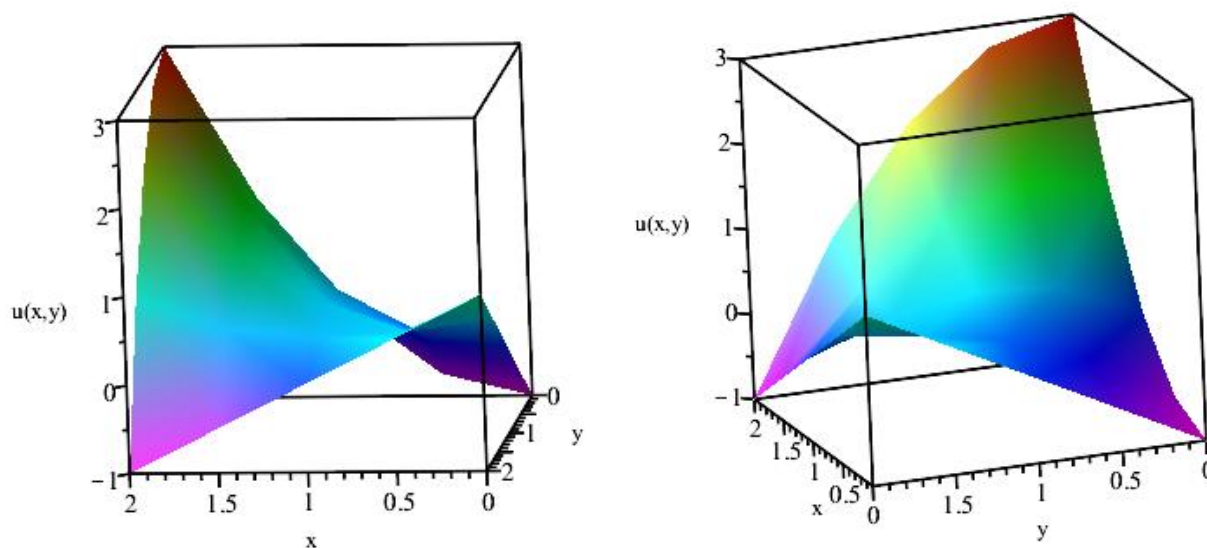


Рис. 3. Візуалізація розв'язку задачі Діріхле для рівняння Лапласа в Maple.

Використовуючи чіткі алгоритми чисельних методів для розв'язання задач математичної фізики студентам пропонується написати програми з візуалізацією отриманих розв'язків будь-якою мовою програмування, наприклад на C++, Python або Java, та провести аналіз програмної реалізації з розв'язанням в середовищі Maple.

Висновки. Використання програмних пакетів та програмної реалізації до вивчених методів значно скорочує час, який необхідно витратити для отримання розв'язку складної задачі. Є можливість отримати результат за досить короткий час замість багатьох годин ручної роботи, що дозволяє більше зосередитися на аналізі результатів та прийнятті рішень. Крім того, програмні пакети добре протестовані, тому забезпечують більш точні та надійні результати. За допомогою програмної реалізації є можливість отримати не тільки розв'язок поставленої задачі, а ще й візуалізацію результатів, що значно полегшує розуміння складних процесів.

References

1. Nikitenko O.M. (2014). Maple: Solving engineering and scientific problems: a training manual. Kharkiv: NURE, 289 p. [in Ukrainian].
2. Vorobyova A.I., Kuriksha O.V. (2009) Symmetric analysis of partial differential equations using the Maple mathematical package. Scientific works of the Petro Mohyla Black Sea State University. Ser.: Computer Technologies, 106, vol. 93, pp. 75-79. [in Ukrainian].
3. Andrunyk V.A., Vysotska V.A., Pasichnyk V.V., Chyrun L.B., Chyrun L.V. (2020). Numerical Methods in Computer Science: Textbook - Lviv: Vydavnytstvo «Novyy svit – 2000», 470 p. [in Ukrainian].
4. Goy T.P., Kopach M.I., Fedak I.V. (2010) Approximate methods for solving differential equations. A textbook for students of the “mathematics” and “applied mathematics” training areas. – Ivano-Frankivsk: Publishing and design department of the Center for Information Technologies of Vasyl Stefanyk Precarpathian National University, 148 p. [in Ukrainian].
5. Brahinets O., & Vorobyova A. (2022). Using the mathematical program Maple to solve systems of linear algebraic equations by direct numerical methods. Computer-integrated technologies: education, science, production, (47), 55-63. <https://doi.org/10.36910/6775-2524-0560-2022-47-08> [in Ukrainian].

**VISUALIZATION OF NUMERICAL SOLUTIONS OF MATHEMATICAL PHYSICS EQUATIONS USING
MAPLE**

The theses consider the visualization of solutions to some mathematical physics problems obtained by numerical methods using the Maple mathematical package, such as: a mixed boundary value problem for the string oscillation equation, the heat conduction equation, and the Dirichlet problem for the Laplace equation. Although mixed boundary value problems have certain difficulties, there is a powerful arsenal of numerical methods and software tools that allow them to be solved. Using Maple for numerical solution of mathematical physics problems allows computer science students and engineers to concentrate on the essence of the problem, they do not have to waste time on complex implementation of numerical algorithms. The study of numerical methods involves their software implementation, which contributes to the development of programming skills, especially in the field of working with data sets, developing effective algorithms and optimizing code. Specialists with a deep understanding of numerical methods are in demand in many industries, including IT companies, which is why when studying such methods we combine the use of the mathematical package Maple and software implementation in various programming languages.

Keywords: partial differential equations, numerical methods, grid method, equations of mathematical physics, Maple.

Відомості про автора / Information about the Author

Оксана Брагінець, канд. фіз.-мат. наук, доцент кафедри інтелектуальних інформаційних систем Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: oksana.brahinets@gmail.com, <https://orcid.org/0000-0003-2635-7203>.

Oksana Brahinets, Candidate of Physical and Mathematical Sciences, Docent of the Department of Intellectual Information Systems of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: oksana.brahinets@gmail.com, <https://orcid.org/0000-0003-2635-7203>.

ПІДХІД ДО МОДЕЛЮВАННЯ ВЗАЄМОДІЇ КОМПОНЕНТІВ ІНТЕЛЕКТУАЛЬНИХ WEB-SERVISІВ

Анотація. Розглянуто підхід до моделювання взаємодії компонентів інтелектуальних web-сервісів на основі алгебри для моделювання взаємодії компонентів інтелектуальних web-сервісів. В якості основних базових конструкцій для алгебри використовуються: послідовність, альтернативний вибір, пріоритетний вибір, цикл та цикл з пріоритетом. Крім того, визначено комбіновані конструкції, якими є паралельність, довільна послідовність, дискримінатор, динамічний вибір та пріоритетний динамічний вибір. Конструкції були обрані для забезпечення загального та розширеного набору операцій для взаємодії комбінації web-сервісів. Після визначення кожної конструкції дається формальна семантика оператора в термінах мереж Петрі.

Ключові слова: Алгебра, web-сервіс, взаємодія web-сервісів, мережі Петрі.

Вступ

Використання методів візуального представлення та моделювання, таких як мережі Петрі на етапі розробки складних технічних систем ефективні за декількома причинами. По-перше візуальні представлення забезпечують високорівневу, але точну мову опису та представлення складної системи, яка дозволяє формально описувати та моделювати на різних рівнях абстракції. По-друге дозволяють моделювати складні системи в динаміці. Одним із прикладів складних систем, які працюють динамічно, це web-сервіси. Web-сервіс представляє собою компонент, написаний у будь-якій мові, який може бути розгорнутим на будь-якій платформі, яка має стандартний інтерфейс на основі XML.

При розробці інтелектуальних web-сервісів особливе значення має поведінка web-сервісів та їх компонент. Поведінка web-сервісів в основному це частково впорядкований набір операцій. Це робить їх зручними для відображення за допомогою мереж Петрі. Операції моделюються переходами, а стани сервісу моделюються позиціями. Стрілки між позиціями та переходами використовуються для позначення причинно-наслідкових зв'язків [4,5].

Компоненти web-сервісів можуть взаємодіяти з іншими додатками, які самі відповідають стандартам web-сервісів. В інтелектуальних web-сервісах встановлюється окрема послідовність використання інших додатків і їх компонентів. Як правило, один сервіс не задовольняє потреб користувачів, і сервіси стають все більш складними. Фактично сучасний web-сервіс створюється шляхом поєднання різних web-сервісів та їх компонентів для створення компонентного сервісу, який пропонує набір нових функціональних послуг. Цей процес називається композицією web-сервісів [7].

При композиції web-сервісів самим критичним є взаємодія web-сервісів та їх компонентів між собою, що вимагає детального формального опису процесів функціонування та дослідження і моделювання їх поведінки. В цьому контексті пропонується стохастична алгебра для моделювання взаємодії web-сервісів та їх компонент на основі мережі Петрі, за допомогою якої провадиться формальний опис та моделювання складних композицій web-сервісів. Базові та розширені конструкції, які підтримуються запропонованою алгеброю, синтаксично і семантично визначені.

Постановка проблеми. Метою даної роботи є розробка сервісної алгебри для побудови моделей взаємодії компонентів інтелектуальних web-сервісів та їх взаємодія з іншими сервісами на основі мереж Петрі, та їх модифікацій.

Аналіз останніх досліджень.

Використання мереж Петрі, як інструмента графічного і математичного моделювання складних систем та процесів останнім часом отримало широке розповсюдження [4,5]. Мережа Петрі (PN) це орієнтований двухдольний граф з двома типами вузлів: позиціями (представлені колами), та переходами (представлені прямокутниками). Дуги графа з'єднують позиції і переходи таким чином, що позиції можуть бути пов'язані

тільки з переходами і навпаки. Позиції в мережі Петрі можуть містити дискретне число токенів. Розподіл токенів над місцями називається маркуванням. Коли К.А. Петрі вперше представив мережі Петрі, вони служили для опису паралельних систем з точки зору причинно-наслідкових зв'язків без обліку часу [11]. Впровадження часових концепцій в мережі Петрі було запропоновано декількома роками пізніше іншими дослідниками [6]. Для моделювання складних систем та процесів різної природи використовуються різні модифікації мереж Петрі: стохастичні мережі Петрі (SPN), часові мережі Петрі (TPN), кольорові мережі Петрі (CPN), нечіткі мережі Петрі (FuzzyPN) та інші. [2,6,8,9,10]. Вперше мережі Петрі були використані, як інструмент для побудови та композиції web-сервісів в роботі [7]. В роботах [7,12] мережі Петрі різних типів були використані для побудови різних web-сервісів. Але задача моделювання взаємодії складних web-сервісів та їх компонентів залишається актуальною. Для вирішення даної задачі розроблено алгебру сервісів із розширеним набором операцій, що дозволяє детально описати взаємодію між сервісами та їх компонентами при побудові і моделюванні складних web-сервісів.

Формальний опис процесу функціонування сервісів

Взаємодія web-сервісів може бути *статичною* (коли компоненти сервісів взаємодіють по попередньо визначеному алгоритму), або *динамічною* (коли компоненти взаємодіють в процесі виконання операції, або процесу). В термінах мереж Петрі пропонується наступний опис.

Сервісна мережа являє собою позначену мережу, тобто це множина

$$SN = (P, T, W, i, o, l, k),$$

де:

- P – кінцева множина позицій,
- $T = IT \cup TT$ – кінцева множина переходів, що представляють операції сервісів. IT – множина безпосередніх переходів, які схематично позначені чорними прямокутниками. TT – це множина часових переходів, що позначені порожніми прямокутниками.
- $W \subseteq (P \times T) \cup (T \times P)$ являє собою набір спрямованих дуг (напрям відносин),
- i – вхідна позиція, $z \bullet i = \{x \in P \cup T \mid (x, i) \in W\} = \emptyset$,
- o – вихідна позиція $z \bullet o = \{x \in P \cup T \mid (o, x) \in W\} = \emptyset$,
- $l : T \rightarrow AU\{\tau\}$ – функція маркування, де A набір імен операцій. Припустимо, що $\tau \notin A$ і позначає приховану операцію,
- k – індекс пріоритету над іншими компонентами.

Набір спрямованих дуг W також можна інтерпретувати як функцію. $W : (P \times T) \cup (T \times P) \rightarrow \{0, 1\}$. **Приховані операції** – це переходи, які не можуть бути спостережуваними. Вони використовуються для розділення зовнішньої і внутрішньої поведінки сервісу.

Web-сервіс – це множина $WS = (NameWs, Desc, Loc, URL, CS, SN, Pws)$, де:

- $NameWs$ – це назва сервісу, яка використовується як її унікальний ідентифікатор,
- $Desc$ – опис сервісу, що надається,
- Loc – це сервер, на якому знаходиться сервіс,
- URL – це виклик web-сервісу,
- CS є набір сервісів-компонентів web-сервісу. Якщо $CS = \{NameWs\}$, тоді WS – базовий сервіс. В іншому випадку WS є складовою послугою і
- SN – це модель сервісної мережі динамічної поведінки сервісів.
- Pws – індекс пріоритета сервісу.

Позиція i є початковим маркуванням сервісу WS_i , тобто тільки i містить маркери. Виконання сервісу WS_i починається, коли маркер знаходиться в позиції i і закінчується, коли маркер досягає позиції o .

Побудова алгебри для моделювання взаємодії web-сервісів

Множина сервісів визначається за допомогою формальної граматики в нотации подібній нотации BN:

$$WS ::= \varepsilon \mid X \mid Seq(WS_1, WS_2) \mid Choise(WS_1, WS_2) \mid (Pr_Choise(WS_1, WS_2) \mid Loop(WS)) \mid Conc(WS_1, WS_2) \mid ArbSeq(WS_1, WS_2) \mid WS_1 \parallel_C WS_2 \mid (WS_1 \mid WS_2) \rightarrow WS_3 \mid [WS_1(p_1, q_1) : WS_2(p_n, q_n)] \mid [Pr_ (WS_1(p_1, q_1) : WS_n(p_n, q_n))]$$

де:

- ε являє собою порожній сервіс, тобто сервіс, який не виконує ніяких операцій.
- X являє собою сервіс-константу, яка визначає базовий сервіс.
- $(Seq(WS_1, WS_2))$ – представляє собою комбінований сервіс, який виконує завдання сервісу WS_1 , за яким виконується завдання сервісу WS_2 . $Seq(WS_1, WS_2)$ є оператором *послідовності*. Якщо нема різниці якіх сервіс виконати першим, а який потім, то такий порядок дій називається невпорядкованою послідовністю. На практиці

порядок вирішується за будь-якою умовою, так як порядок не важливий, то неупорядковану послідовність також можна розглядати як *послідовність*.

– (*Choise* (WS_1, WS_2)) – представляє собою комбінований сервіс, який виконується як сервіс WS_1 , або сервіс WS_2 . Коли один з сервісів обрано, тоді інший сервіс не розглядається. *Choise* (WS_1, WS_2) – це оператор *альтернативного вибору*. Вибір не є довільним і залежить від умов. Таким чином, умовну конструкцію з булевими варіантами можна розглядати як сервіс вибору.

– (*Pr_Choise* (WS_1, WS_2)) – представляє собою комбінований сервіс, який виконується як сервіс WS_1 , або сервіс WS_2 . Коли один з сервісів обрано, тоді інший сервіс не розглядається. *Pr_Choise* (WS_1, WS_2) – це оператор *пріоритетного вибору*. Вибір не є довільним і залежить від умов пріоритету. Таким чином, умовну конструкцію з булевими варіантами можна розглядати як сервіс пріоритетного вибору.

– (*Loop* (WS)) – є комбінованим сервісом, який виконує певну кількість разів сервіс WS . *Loop* (WS) є оператором *ітерації*.

– (*Conc* (WS_1, WS_2)) – це комбінований сервіс, який виконує завдання сервісів WS_1 і WS_2 незалежно та паралельно. Обидва сервіси включаються одночасно, і узагальнений сервіс очікує, поки обидва сервіси WS_1 і WS_2 не будуть завершені. *Conc* (WS_1, WS_2) є *паралельним* оператором.

– (*ArbSeq* (WS_1, WS_2)) – представляє собою комбінований сервіс, який виконує завдання сервісу WS_1 , за яким виконується завдання сервісу WS_2 . Або спочатку виконується сервіс WS_2 , а потім сервіс WS_1 . *ArbSeq* (WS_1, WS_2) є оператором *довільної послідовності*.

– ($WS_1 \parallel_C WS_2$) – є комбінованим сервісом, який виконує сервіси WS_1 та WS_2 незалежно один від одного, з можливостями комунікації над множиною C пар операцій, тобто $\parallel_C$ – є *паралельним* оператором зв'язку.

– ($((WS_1 \mid WS_2) \rightarrow WS_3)$) – є комбінованим сервісом, який очікує виконання одного з сервісів (або WS_1 , або WS_2), перш ніж активувати наступний сервіс WS_3 . $\rightarrow$ – оператор *дискримінатора*. Сервіси WS_1 та WS_2 не зв'язані між собою.

– [$WS_1(p_1, q_1) : WS_n(p_n, q_n)$] – *динамічний вибір*, є комбінованим сервісом, який динамічно вибирає один сервіс серед n доступних сервісів $WS_1, \dots, WS_n$ і виконує його. Він працює наступним чином: спочатку запит відсилається постачальнику послуг першого сервісу на їх вхідні точки доступу $p_1, \dots, p_n$. Тоді, виходячи з отриманих відповідей, з їх вихідних точок доступу $q_1, \dots, q_n$, та відповідно до певних критеріїв, обирається найкращий постачальник сервісу. Відповідно $[:]$ є оператором динамічного вибору.

– [$Pr_WS_1(p_1, q_1) : WS_n(p_n, q_n)$] – *пріоритетний динамічний вибір*, є комбінованим сервісом, який динамічно вибирає один сервіс серед n доступних сервісів $WS_1, \dots, WS_n$ і виконує його. Він працює наступним чином: спочатку запит відсилається постачальнику послуг першого сервісу на їх вхідні точки доступу $p_1, \dots, p_n$. Тоді, виходячи з отриманих відповідей, з їх вихідних точок доступу $q_1, \dots, q_n$, та відповідно до певних критеріїв пріоритету, обирається найкращий постачальник сервісу. Відповідно $[:]$ є оператором пріоритетного динамічного вибору.

Формальне визначення операторів взаємодії в термінах мереж Петрі наступне. Нехай $WS_i = (Name_{WS_i}, Desc_i, Loc_i, URL_i, CS_i, SN_i, Pws)$ для $SN_i = (P_i, T_i, W_i, i_i, o_i, l_i, k)$ для $i=1, \dots, n$, буде n web-сервісів таким щоб $P_i \cap P_j = \emptyset$, та $T_i \cap T_j = \emptyset$ для $i \neq j$.

Склад сервісів може застосовуватись до синтаксично різних складних інтелектуальних web-сервісів. Це пов'язано з тим, що для правильної взаємодії позицій та переходів сервісів вони не повинні перетинатися. Кожен сервіс містить дві окремі частини, одну частину для обробки запиту на обслуговування, а інша частина – повинна виконувати сервіс самостійно. Пара операцій *SendReqServ* і *SelectServ* є результатом обраної стратегії вибору. Запропонована алгебра перевіряє властивість замикання. Це гарантує, що кожен результат операції над сервісами є сервісом, до якого можна знову застосувати алгебраїчні оператори. Таким чином, можливо будувати більш складні та комбіновані сервіси шляхом агрегування і повторного використання існуючих сервісів за допомогою декларативних виразів сервісної алгебри.

Висновки

Представлено підхід до моделювання інтелектуальних сервісів та їх компонентів на основі алгебри сервісів. Семантика представленої алгебри може бути використана для доказу алгебраїчних властивостей конструкцій інтелектуальних web-сервісів. Це може бути тоді, коли створюється Інтелектуальний web-сервіс, на основі об'єднання існуючих web-сервісів з використанням операторів алгебри. Алгебраїчні властивості можуть потім використовуватися для перетворення і оптимізації комбінованих web-сервісів на основі наступних операційних показників web-сервісів, таких як вартість і тривалість. Використання сервісної алгебри з представленим набором операцій дозволяє використовувати її для моделювання web-сервісів більш складні типи мереж Петрі.

References

1. Aliiev T.I. Osnovy modeliuвання dyskretnykh system/T.I. Aliiev. – SPbHU YTMO, 2009. – 363 s.
2. Bodianskyi Ye.V., Kucherenko Ye.I., Mykhaliev O.I. Neiro-fazzi merezhi Petri u zadachakh modeliuвання skladnykh system: monohrafiia. Dnipropetrovsk : Systemni tekhnologii, 2005. 311 s.
3. Bokhan K.A. Modeli korporatyvnykh servisiv na osnovi iierarkhichnykh merezh Petri/K.A. Bokhan, M.S. Khudolei // Radioelektronni ta kompiuterni systemy. – 2010. – Vyp. 47. – С. 36-41.
4. Kotov, V.Ie. Merezhi Petri [Tekst]/V.Ie. Kotiv. - M.: Nauka. Holovna redaktsiia fizyko-matematychnoi literatury, 1984. – 160 s.
5. Piterson Dzh. Teoriia merezh Petri ta modeliuвання system: prov. z anhl. / Dzh. Piterson. Svit, 1984. – 264 s.
6. Bause, F. Stochastic Petri nets: an introduction to the theory [Text] / F. Bause, Pieter S. Kritzinger. – Friedrich Vieweg & Sohn Verlag, 2002. – 223 p.
7. Hamadi R. and Benatallah B. “A Petri net based-model for web service composition”, in proc. the 14th Australasian database conference, adelaide. Darlinghurst: Australian Computer Society, 2003, pp. 191-200.
8. Jensen K., Kristensen L.M., Wells L. Coloured Petri Nets and CPN Tools for Modelling and Validation of Concurrent Systems. Software Tools for Technology Transfer manuscript. 2007. 40 p.
9. Genrich H. J. Lautenbach K. “System modeling with high level Petri nets”, Theoretical Computer Science, Vol. 13, pp. 109-136, 1981.
10. Modelling with Generalized Stochastic Petri Nets [Text] / M. Ajmone Marsan, G. Balbo, G. Conte, S. Donatelli, G. Franceschinis. – John Wiley & Sons, 1995. – 324 p.
11. Murata T. “Petri nets: Properties, analysis and applications,” in Proc. of the IEEE, Vol. 77(4), 1989, pp. 541580.
12. Perkusich J. and J. C. A. De Figueiredo, “G-nets: A Petri net based approach for logical and timing analysis of complex software systems”, Journal of Systems and Software, Vol. 39, Issue 1, pp. 39-59, Oct 1997.

DOI 10.34132/mspc2025.01.14.06

Viktor Gozhij

APPROACH TO MODELING INTERACTION OF COMPONENTS OF INTELLECTUAL WEB SERVICES

An approach to modelling the interaction of components of intelligent web services based on algebra for modelling the interaction of components of intelligent web services is considered. The following basic constructions are used for algebra: sequence, alternative choice, priority choice, cycle and cycle with priority. In addition, combined constructions are defined, which are parallelism, arbitrary sequence, discriminator, dynamic choice and priority dynamic choice. The constructions were chosen to provide a general and extended set of operations for the interaction of a combination of web services. After defining each construction, the formal semantics of the operator in terms of Petri nets is given.

Keywords: Algebra, web service, web service interaction, Petri nets.

Відомості про автора / Information about the Author

Віктор Гожий, кандидат технічних наук, старший викладач кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: gozhyi.v@gmail.com, orcid: : <https://orcid.org/0000-0002-5341-0973>

Viktor Gozhij, Candidate of Technical Sciences, Senior Lecturer, Department of Intellectual Information Systems of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: gozhyi.v@gmail.com, orcid: <https://orcid.org/0000-0002-5341-0973>

DOI 10.34132/mspc2025.01.14.07

Олександр Гожий
Ірина Калініна

МЕТОДОЛОГІЯ АНАЛІЗУ ТА ПОПЕРЕДНЬОЇ ОБРОБКИ ДАНИХ ДЛЯ ВИРІШЕННЯ ЗАДАЧ МАШИННОГО НАВЧАННЯ

В статті описується та досліджується методологія аналізу та попередньої обробки даних для вирішення задач машинного навчання. Методологія об'єднує на основі системного підходу наступні групи методів: методи обробки пропусків в даних, методи обробки аномальних значень, методи генерації ознак, методи ідентифікації нелінійностей та нестационарностей та методи нормалізації. Досліджено метод генерації ознак. Методи відбору і генерації ознак поділяють на три основні групи: методи фільтрації, методи обгортки і вбудовані методи. Метод генерації ознак складається з п'яти кроків для ефективного вибору найбільш релевантних ознак набору даних. Він пропонує кращий підхід до відбору ознак. Це призводить до покращення продуктивності моделі з меншою кількістю ознак і скорочення обчислювальних витрат. Для експериментальної перевірки методології аналізу та попередньої обробки даних була розглянуто створення системи класифікації якості червоних вин. Для вирішення завдання класифікації якості вина було проведено моделювання з використанням різних алгоритмів. Було доведено ефективність методу для вирішення задач аналізу та попередньої обробки даних для вирішення задач класифікації.

Ключові слова, методологія аналізу та попередньої обробки даних, методи генерації ознак, методи обробки пропусків в даних, методи обробки аномальних значень, методи ідентифікації нелінійностей та нестационарностей та методи нормалізації.

Аналіз та попередня обробка даних, це початковий етап вирішення завдань машинного навчання. Від якісно проведеного аналізу та попередньої обробки даних залежить кінцевий результат вирішення завдання. На цьому етапі досліджуються реальні данні різної природи, їх властивості та поведінка, які впливають на процес вирішення завдання в цілому. Ігнорування цього етапу призводить до неможливості побудови коректних ймовірнісно-статистичних моделей для конкретного процесу та їх придатності для розв'язання задачі машинного навчання.

На етапі аналізу та попередньої обробки даних вирішуються наступні задачі: визначення набору даних для аналізу та попередньої обробки даних; ідентифікація та заповнення пропусків в даних; ідентифікація та підбір методів обробки аномальних значень в даних; визначення зовнішніх впливів і збурень та їх типу, та встановлення можливості їх ймовірнісно-статистичного опису за допомогою конкретних типів розподілів випадкових величин; визначення методів та алгоритмів фільтрації з метою виділення корисної складової при обробці даних; генерація ознак в наборі даних; аналіз та встановлення основних типів нелінійностей та нестационарностей даних; розробка підходу до застосування методів нормалізації та стандартизації даних.

Для ефективного вирішення задач аналізу та попередньої обробки даних в задачах машинного навчання, необхідно визначити головні властивості тих наборів даних, що досліджуються:

1. **Складність** – кількість і різноманітність факторів, які складають та будуть впливати на набір даних.
2. **Множина зв'язків** між факторами, які характеризують данні, тобто сила, з якою зміна одного параметра (фактора) впливає на зміну інших параметрів набору даних
3. **Динамічність** – швидкість, з якою відбуваються зміни у наборі даних.
4. **Невизначеність** Поява невизначеності зумовлена недостатністю або спотворенням інформації на будь-якому етапі згаданого процесу. Це може бути коротка вибірка даних, якої недостатньо для побудови адекватної моделі; спотворення вимірів шумовими складовими; некоректність встановлення типу розподілу даних і, як наслідок, невірний вибір методу оцінювання параметрів моделі; помилки дослідника при обчисленні оцінок прогнозів або альтернативних рішень на їх основі.

Виділення та обробка такого роду інформаційних характеристик при аналізі інформації для опису набору даних свідчить про те, що необхідно застосовувати *системний підхід* і розглядати набір даних як систему або сукупність систем, які впливають на кінцевий результат аналізу. Саме в рамках цього підходу прийнято представляти будь-які об'єкти у вигляді структурованої системи, виділяти елементи системи, взаємозв'язок між ними і динаміку розвитку елементів і всієї системи в цілому. Тому аналіз і попередня обробка даних, яка використовується для вирішення різних завдань машинного навчання, а також для накопичення необхідної інформації для подальшого використання на різних етапах її обробки, розглядається як необхідний компонент методології аналізу та попередньої обробки даних в завданнях машинного навчання.

Метою даної роботи є дослідження системного підходу до аналізу та попередньої обробки даних при розв'язанні задач машинного навчання. Для цього необхідно розробити та дослідити методологію аналізу та попередньої обробки даних.

Структура методології аналізу та попередньої обробки даних

Методологія аналізу та попередньої обробки даних побудована на основі системного підходу і об'єднує групи методів та методологічних підходів, які згруповані по функціональному призначенню в процесі аналізу та підготовки даних до моделювання. Структурна модель методології аналізу та попередньої обробки даних представлена на рисунку 1. Методологія передбачає системне використання наступних груп методів: методів обробки пропусків в даних, методів обробки аномальних значень, методів генерації ознак, методів ідентифікації та подолання нелінійностей та нестационарностей, методів нормалізації даних.

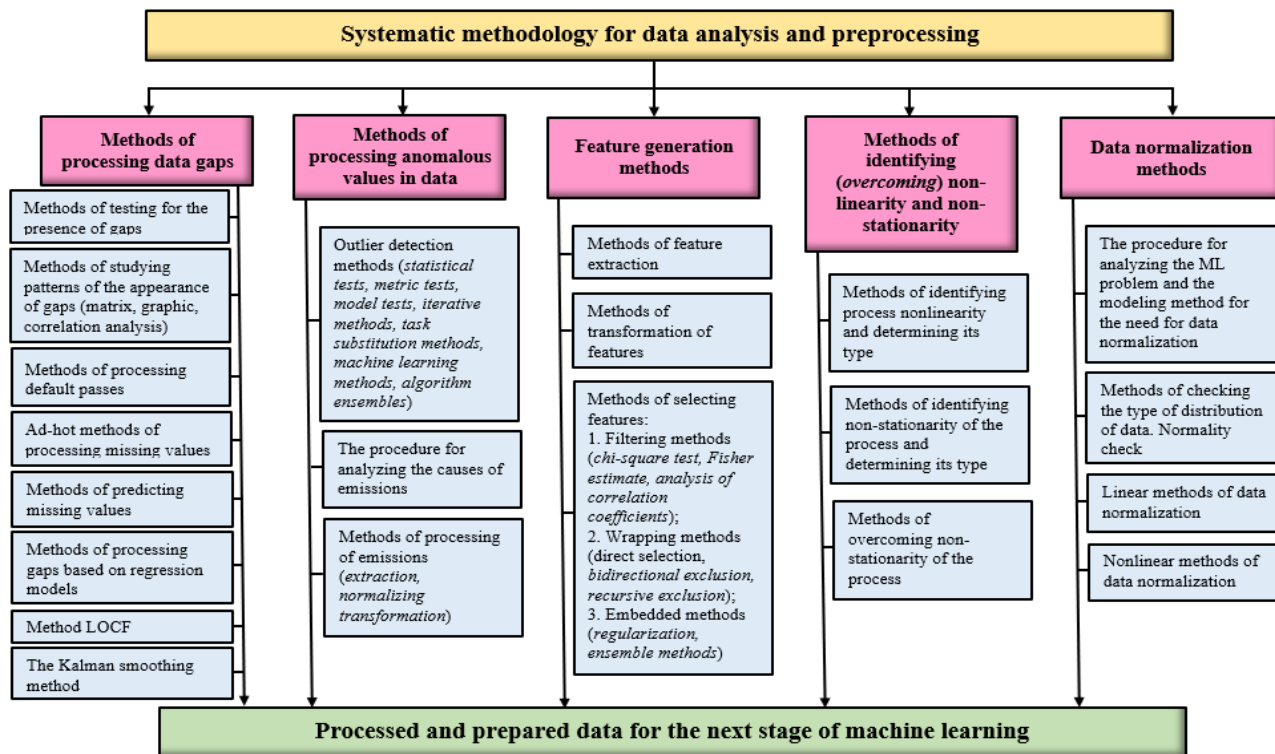


Рис. 1. Методологія аналізу та попередньої обробки даних

Джерело: сформовано автором

Методи обробки пропусків в даних. Методи обробки пропусків в даних об'єднують метод і використання тестів для перевірки на наявність пропусків, різновиди методів дослідження закономірностей появи пропусків, методи обробки пропусків за замовченням, ad-hot методи, методи прогнозування відсутніх значень, методи обробки пропусків на основі регресійних моделей, метод LOCF, методи згладжування на основі фільтра Калмана [1].

Методи обробки аномальних значень. Методи обробки аномальних значень об'єднують методи виявлення аномальних значень в наборі даних, процедуру аналізу причин їх появи та методи обробки викидів [2].

Методи генерації ознак. Методи генерації ознак об'єднують методи вилучення постійних, квазіпостійних та дублюючих ознак, методи перетворення ознак та три групи методів по відбору ознак.

Методи ідентифікації (подолання) нелінійностей та нестационарностей. Методи ідентифікації (подолання) нелінійностей та нестационарностей об'єднують методи ідентифікації нелінійностей процесу та визначення його типу, методи ідентифікації нестационарностей процесу та визначення його типу, методи подолання нестационарностей процесу [3].

Методи нормалізації даних. Методи нормалізації даних об'єднують процедуру аналізу задачі машинного навчання та метода моделювання на необхідність нормалізації даних, методи перевірки типу розподілу даних (перевірка на нормальність), лінійні та нелінійні методи нормування даних [4].

Результатом використання методів і методологічних підходів методології аналізу та попередньої обробки даних є оброблений та підготовлений до наступного етапу машинного навчання набір даних.

Методи генерації ознак

Відбір ознак – це процес вибору найбільш релевантних ознак моделі машинного навчання, адаптований до конкретної проблеми, яку намагаються вирішити. Це передбачає вибіркове включення чи виключення важливих ознак із збереженням їх незмінними. Таким чином, він ефективно усуває нерелевантний шум з даних та зменшує розмір та обсяг вхідного набору даних. Мета відбору ознак у машинному навчанні полягає у наступному: зменшити розмірність простору ознак; прискорити алгоритм навчання; покращити точність прогнозування алгоритму класифікації; покращити зрозумілість результатів навчання.

Комбінований метод відбору ознак. Оскільки відбір ознак грає вирішальну роль машинному навчанні, підвищуючи продуктивність моделі і знижуючи обчислювальні витрати, то роботі запропоновано комбінований метод відбору ознак, який включає поетапне використання і методів фільтрації і методів обгортки. На рисунку 2 представлено блок-схему запропонованого комбінованого методу.

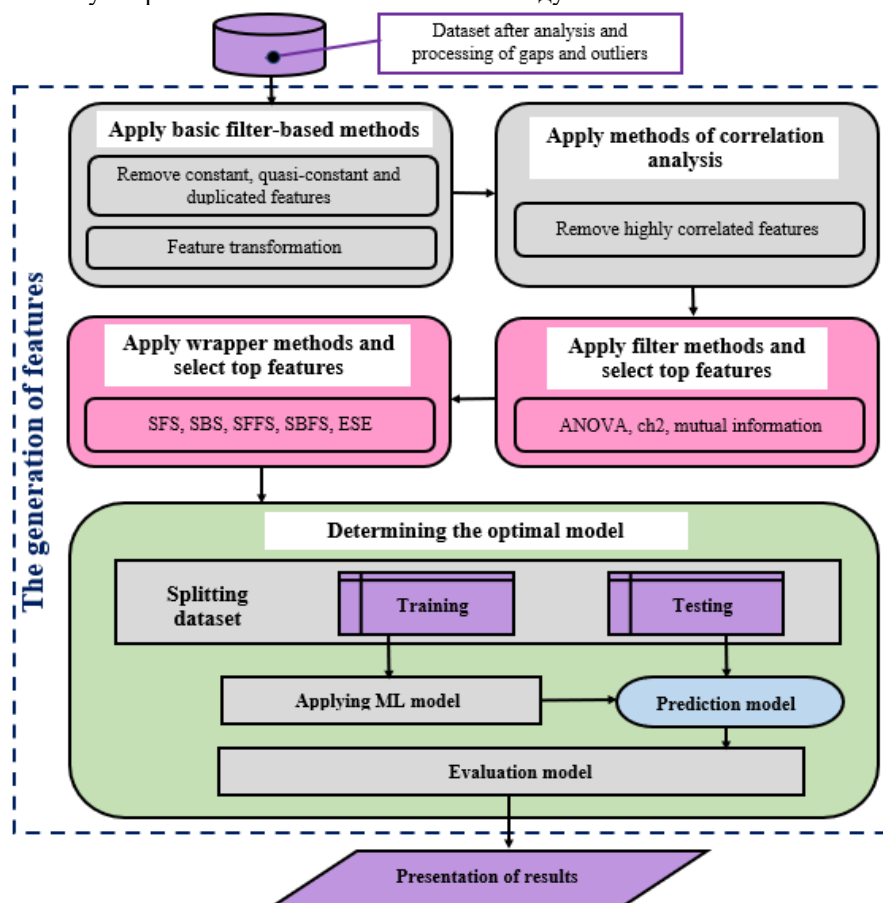


Рис. 2. Блок-схема комбінованого методу відбору ознак
 Джерело: сформовано автором

Метод складається з кількох кроків для ефективного вибору найбільш релевантних ознак набору даних.

Крок 1. Застосовування базових методів вибору ознак на основі фільтрів для виявлення та видалення постійних, квазіпостійних та дублюючих ознак.

Крок 2. Використання кореляційного аналізу для виключення ознак з високою кореляцією вище за певний поріг.

Крок 3. Використання методів фільтрації для відбору ознак, такі як ANOVA, хі-квадрат та взаємна інформація для вибору найбільш важливих ознак.

Крок 4. Використання методів обгортки, включаючи послідовний прямий вибір ознак (SFS), послідовний зворотний вибір ознак (SBS), послідовний плаваючий прямий вибір ознак (SFFS), послідовний плаваючий вибір ознак (SBFS) і вичерпний вибір ознак (EFS) для подальшого уточнення підмножини ознак.

Крок 5. Для відбору оптимальної моделі з кількох альтернативних використовуються методи *перевірочної вибірки* та *перехресної перевірки*. Для цього на кроці визначення оптимальної моделі набір даних поділяється на два, навчальний та тестовий набори з використанням вибраних ознак. Дали використовується алгоритми машинного навчання для навчання моделі з меншою кількістю ознак, виконується налаштування моделі та оцінка результатів моделювання.

Комбінований метод пропонує кращий підхід до відбору ознак. Це призводить до покращення продуктивності моделі з меншою кількістю ознак і скорочення обчислювальних витрат. Завдяки систематичному відбору найбільш інформативних ознак та використанню методів фільтрації та обгортки комбінований метод забезпечує більш ефективне та дієве навчання та розгортання моделі машинного навчання.

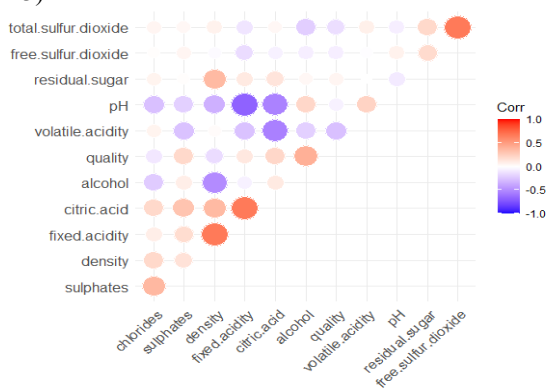
Експериментальна частина

Для вирішення задачі машинного навчання (класифікації) на основі методології аналізу та попередньої обробки даних було використано набір даних *Red Wine Quality*. Набір стосується червоних сортів португальського вина «Vinho Verde» [1]. Мета класифікації – виявити чинники, пов'язані з ризиком невчасної доставки попередніх замовлень клієнтам та інформацію про одержувачів товарів. Набір даних складається з 1599 записів і включає 11 атрибутів та мітку. Це набір числових ознак, що визначають такі характеристики вин, як: значення фіксованої та летючої кислотності, вміст лимонної кислоти, залишкового цукру, хлоридів, вільного діоксиду сірки, загального діоксиду сірки, а також значення густини, pH, сульфатів та алкоголю. Результуюча змінна, якість вина, категоріальна, що приймає значення від 0 (дуже погана якість) до 10 (відмінна якість).

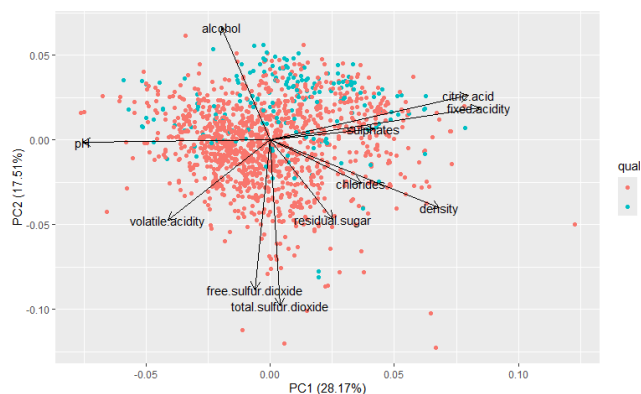
Аналіз та попередня обробка даних. Після завантаження дані були проаналізовані за допомогою звітів по зведеної статистики та побудована гістограма розподілу якості червоного вина. Графік показує що якісні вина значно переважають погані вина з великим відривом. Більшість вин є непоганими (оцінка 5 або 6 балів), але також є і вина поганої якості (3 або 4 бали). Переважна більшість якісних вин має рейтинг якості 7.

За результатами аналізу, визначено відсоток викидів в кожному атрибуті та їх кількість. Найбільша кількість викидів у атрибута «residual.sugar», 9,7% зі всіх значень. Враховано, що деякі з методів машинного навчання працюють погано за наявності викидів. Тому аномальні значення в наборі оброблені за допомогою метода **затискного перетворення**

Оскільки метою класифікації є прогноз якості вина, важливо дізнатися, яка зі змінних має найсильніший зв'язок з цільовою ознакою, якістю вина. З теплової карти, алкоголь має найсильнішу кореляцію з якістю вина (рис.3).



а) Теплова карта



б) Результати методу головних компонент

Рис. 3. Аналіз та попередня обробка даних

Джерело: сформовано автором

Для подолання мультіколеніарності використано метод головних компонент (рис.3). Таким чином, головними властивостями, що розрізняють якісне та посереднє вино, є вміст сульфатів та рівень алкоголю. Алкоголь має найсильніший зв'язок з якістю вина.

Для мінімізації «кореляційних плеяд» використано фактор інфляції дисперсії VIF. Для відбору оптимальної підмножини змінних в моделях машинного навчання використані методи обгортки (методи *покрокового включення* та *покрокового виключення*).

Перевірка розподілів в ознак в наборі даних та аналіз описової статистики виявили необхідність нормалізації. Нормалізацію числових змінних виконано з допомогою перетворення Йео-Джонсона, тому що набір містить нульові значення.

Результати роботи класифікаторів

Для вибору оптимальної моделі для вирішення завдання класифікації якості вина було проведено моделювання з використанням різних алгоритмів. Оцінка результатів класифікації за допомогою метрик якості зібрано у таблицю 1.

Таблиця 1.

Оцінка результатів класифікації

Тип моделі	Коефіцієнт успішності (accuracy)	Коефіцієнт помилок (error rate)	Каппа-статистика (Kappa)	Чутливість (sensitivity)	Специфічність (specificity)	Точність (precision)	Повнота (recall)	F-міра (F-measure)
LDA	0.8616	0.1384	0.4166	0.9922	0.3226	0.8581	0.9922	0.9203
QDA	0.8868	0.1132	0.5774	0.9766	0.5161	0.8929	0.9766	0.9329
Bagging	0.8742	0.1258	0.5166	0.9766	0.4516	0.8803	0.9766	0.9260

Оцінка кращої моделі була здійснена на основі аналізу ефективності кожної моделі, проте рішення про вибір було прийнято з урахуванням значень F-міри. Отже, найкраща модель – QDA зі значенням F-міри=0.9329, або 93%.

Висновки

В роботі досліджується методологія аналізу та попередньої обробки даних для вирішення задач машинного навчання. Методологія об'єднує наступні групи методів: методи обробки пропусків в даних, методи обробки аномальних значень, методи генерації ознак, методи ідентифікації нелінійностей та нестационарностей та методи нормалізації. Методи поєднано на основі системного підходу. Детально досліджено методи генерації ознак. Методи відбору і генерації ознак поділяють на три основні групи: методи фільтрації, методи обгортки і вбудовані методи. Запропоновано комбінований метод відбору ознак, який включає поетапне використання і методів фільтрації і методів обгортки. Метод складається з кількох кроків для ефективного вибору найбільш релевантних ознак набору даних. Він пропонує кращий підхід до відбору ознак. Це призводить до покращення продуктивності моделі з меншою кількістю ознак і скорочення обчислювальних витрат.

References

1. Roonizi, K. (2023) Kalman filter/smoothing-based design and implementation of digital IIR filters, *Signal Processing*, vol. 208, 108958. URL: <https://doi.org/10.1016/j.sigpro.2023.108958>
2. Bidyuk, P., Kalinina, I. & Gozhyj, A. (2022) An Approach to Identifying and Filling Data Gaps in Machine Learning Procedures. *Lecture Notes on Data Engineering and Communications Technologies*, vol. 77, (pp. 164–176). [in Switzerland]
3. Kiptoon, D. (2023) Understanding Feature Engineering in Machine Learning. [Online]. URL: <https://medium.com/@jdkiptoon/understanding-feature-engineering-in-machine-learning-59fc343a29c9> (application date: 15.08.2024).

4. Kalinina, I. & Gozhyj, A. (2022) Modeling and forecasting of nonlinear nonstationary processes based on the Bayesian structural time series. *Applied Aspects of Information Technology*, vol. 5, no. 3, (pp. 240-255). URL: <https://doi.org/10.15276/aait.05.2022.17>.

DOI 10.34132/mspc2025.01.14.07

Oleksandr Gozhyj,
Iryna Kalinina

METHODOLOGY OF DATA ANALYSIS AND PRE-PROCESSING FOR SOLVING MACHINE LEARNING PROBLEMS

The paper describes and investigates the methodology of data analysis and pre-processing for solving machine learning problems. The methodology combines the following groups of methods based on a systematic approach: methods for processing data gaps, methods for processing abnormal values, methods for generating features, methods for identifying nonlinearities and non-stationarities, and methods for normalization. The method of generating features is investigated. Methods for selecting and generating features are divided into three main groups: filtering methods, wrapper methods, and embedded methods. The feature generation method consists of five steps for effectively selecting the most relevant features of the data set. It offers a better approach to feature selection. This leads to improved model performance with fewer features and reduced computational costs. To experimentally verify the methodology of data analysis and pre-processing, the creation of a red wine quality classification system was considered. To solve the problem of wine quality classification, modeling was carried out using various algorithms. The effectiveness of the method for solving data analysis and preprocessing problems for solving classification problems was proven.

Key words, data analysis and preprocessing methodology, feature generation methods, data gap processing methods, outlier processing methods, nonlinearity and non-stationarity identification methods, and normalization methods.

Відомості про авторів/ Information about the Authors

Олександр Гожий, доктор технічних наук, професор кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: alex.gozhyj@gmail.com, orcid: <https://orcid.org/0000-0002-3517-580X>.

Oleksandr Gozhyj, Doctor of Technical Sciences, Professor of the Department of Intellectual Information Systems, of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: alex.gozhyj@gmail.com, orcid: <https://orcid.org/0000-0002-3517-580X>

Ірина Калініна, доктор технічних наук, професор кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: irina.kalinina1612@gmail.com, orcid: <https://orcid.org/0000-0001-8359-2045>

Iryna Kalinina, Doctor of Technical Sciences, Professor of the Department of Intellectual Information Systems, of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: irina.kalinina1612@gmail.com, orcid: <https://orcid.org/0000-0001-8359-2045>

DOI 10.34132/mspc2025.01.14.08

Олександр Єлісєєв,
Гліб Горбань

АВТОМАТИЧНА ГЕНЕРАЦІЯ ПРОЄКТІВ ЗА ДОПОМОГОЮ LLM-АГЕНТІВ

Тези досліджують використання LLM-агентів для автоматичної генерації програмних проєктів, аналізуючи їхні можливості, переваги та ризики порівняно з традиційними алгоритмами та простими LLM-запитами. Розглядається архітектура агентів, їх інтеграція з сучасними технологіями (Django, React, OpenAI API), а також сценарії застосування в бізнесі, освіті та фінансових технологіях. Розглядається система автоматичної генерації проєктів із залученням LLM-агентів. Наголошується на важливості обґрунтованого вибору між складними агентами та простими рішеннями залежно від вимог проєкту для забезпечення ефективності, безпеки та мінімізації технічного боргу.

Ключові слова: LLM-агенти, великі мовні моделі, автоматична генерація коду, штучний інтелект, розробка програмного забезпечення, Django, React, OpenAI API.

У прагненні до використання штучного інтелекту існує поширене переконання, що найсучасніше рішення – зазвичай це AI-агент – завжди є найкращим вибором. Однак у більшості випадків простіші підходи, такі як виклики великих мовних моделей (LLM) або традиційні алгоритми, не лише є достатніми, але й значно практичнішими та ефективнішими. Багато проєктів автоматично використовують AI-агенти, хоча простий виклик LLM або традиційний алгоритм можуть дати необхідний результат із меншими витратами ресурсів.

Останнім часом великі мовні моделі (LLM) стали неймовірно популярними. На сьогоднішній день вже спостерігається активний розвиток цієї галузі, що підтверджується численними дослідженнями, публікаціями, фреймворками та проєктами провідних компаній. Прикладом може бути GitHub CoPilot.

Розглянемо, що саме мається на увазі під поняттям LLM-агенти.

Агенти – це системи, які взаємодіють із динамічним середовищем, сприймають його та діють відповідно до закладених цілей або завдань. У цьому контексті LLM (великі мовні моделі) – це потужні моделі, здатні передбачати ймовірний наступний токен на основі вхідного тексту (де токен – це слово або його частина) [1].

Агенти, які працюють виключно з текстовою інформацією, можуть ефективно виконувати широкий спектр складних завдань, отримуючи дані про стан середовища у текстовому вигляді. Водночас їхні можливості істотно зростають, якщо вони здатні обробляти інші модальності, такі як аудіо чи зображення. Для цього важливо не лише перетворити кожну з таких модальностей у вектор (або послідовність векторів), але й забезпечити, щоб усі ці вектори належали до єдиного прихованого простору.

Автоматична генерація проєктів із використанням LLM-агентів докорінно змінює процес створення програмного забезпечення. Сучасні LLM-агенти, зокрема побудовані на GPT-4, Claude 3, Gemini та інших передових архітектурах, дозволяють переходити від трудомісткого, ручного написання коду до сценаріїв, де вся основна робота виконується штучним інтелектом. На практиці це виглядає так: замовнику або розробнику достатньо описати функціональність і вимоги зрозумілою мовою, після чого агент аналізує промпт, формує технічне завдання, створює репозиторій, ініціалізує структуру проєкту, генерує код для бекенду (наприклад, на Django чи Flask), одразу прописує тестові сценарії з pytest або unittest, формує docker-контейнери, налаштовує CI/CD на GitHub Actions або GitLab CI, дописує документацію в docstring та створює README, що пояснює використання, а потім автоматично розгортає базову версію проєкту на Heroku, Vercel чи іншому хмарному провайдері.

LLM-агенти не обмежуються звичайною генерацією коду. Вони здатні самостійно аналізувати помилки компіляції, запускати юніт-тести, знаходити й виправляти баги, отримувати фідбек від користувача через інтерактивний чат і далі адаптувати код або інтерфейс. Якщо, наприклад, агент отримує повідомлення про невідповідність API-ендпоінтів чи помилкову логіку автентифікації, він самостійно вносить зміни, фіксує їх у git та повторно тестує результати. Тобто це не просто генератор коду, а багатофункціональний автоматизований виконавець із можливістю ітеративного вдосконалення проєкту без постійної участі людини.

Яскравим прикладом може слугувати генерація повноцінного MVP SaaS-сервісу. Припустимо, що користувач хоче отримати вебзастосунок для бронювання готелів із системою авторизації, панеллю адміністратора, інтеграцією з поштою, автогенерацією інвойсів і мобільною адаптацією. LLM-агент самостійно розбиває задачу на етапи: генерує backend на Django (створює моделі для готелів, бронювань, користувачів, реалізує REST API через DRF), підключає простий frontend на Node.js або генерує React-компоненти, прописує Swagger/OpenAPI документацію, додає автоматичні тести, створює docker-compose для локального запуску, налаштовує email-інтеграцію через SMTP або сторонній сервіс, а потім моніторить логи й оптимізує код за результатами тестування. Весь цей процес може відбуватись у межах одного інтерактивного діалогу, де агент уточнює деталі (наприклад, зовнішній вигляд інтерфейсу чи специфіку платежів) і одразу втілює зміни.

Щоб зрозуміти, чим саме є LLM-агенти, варто розкласти їхню архітектуру на складові. Це багатоетапна система, що поєднує в собі LLM (яка відповідає за генерацію і аналіз текстів та коду), інтерфейс зв'язку з користувачем, спеціалізовані інструменти для роботи з файлами, системами контролю версій, інфраструктурою, а також модулі планування і координації завдань. Наприклад, популярний агент Open Interpreter здатен не тільки генерувати Python-скрипти для обробки CSV-файлів, а й автоматично завантажувати файли, виконувати скрипти на сервері, зчитувати отримані результати, будувати діаграми і генерувати нові інструкції на основі проміжних даних.

Конкретні типові сценарії вже застосовуються у сучасному бізнесі, наприклад у сфері фінансових технологій, де LLM-агенти генерують індивідуальні робочі кабінети для кожного клієнта автоматично після підписання договору – від створення бази даних до налаштування інтеграції з банківськими API. Також у онлайн торгівлі агенти пишуть промо-лендінги, налаштовують канали комунікації через Telegram Bot API, генерують сценарії масових email-розсилок, шаблони лендінгів та системи персоналізації. В освітніх проєктах агент формує динамічні навчальні платформи з реєстрацією, системою модулів, автоматичними тестами та завантаженням сертифікатів у PDF.

Проте при всій технологічності LLM-агентів їх треба застосовувати обґрунтовано. Задачі, де потрібна максимальна швидкість – наприклад, масовий парсинг даних, конвертація великих обсягів даних у різні формати, перетворення тексту чи структурованих даних – часто ефективніше вирішити через pipeline із класичних скриптів і одноразових LLM-запитів. Якщо бізнес-процес вже усталений, а специфікація чітка, дешевше й надійніше використати генерацію шаблонів через template-інструменти та традиційні build-системи. AI-агенти виправдані для реальних R&D-завдань, стартапів, виробництва прототипів, коли замовник або команда постійно коригують вимоги на ходу [2].

Серйозною перевагою LLM-агентів виступає їх здатність до автоматизації повного циклу розробки. Наприклад, агент може проаналізувати issue у JIRA, сформулювати план розробки з оцінкою часу, згенерувати основні класи і функції, налаштувати автоматичне тестування через pytest-cov, пройтись по checklist якості, запустити лінійтер, згенерувати pull request і самостійно зібрати звіт про покриття тестами. В окремих екосистемах агенти навіть інтегруються з MLOps-середовищем: після тренування моделі агент створює REST API для її розгортання, прописує dockerization, підключає Prometheus для моніторингу, тестує latency та додає алерти для DevOps-команди. Не менш важливо, що такі агенти можуть запускати бета-тестування, збирати поведінкові дані користувачів, автоматично аналізувати фідбек і пропонувати зміни у функціоналі та UI. Таким чином, до переваг використання агентів можна віднести наступне: 1) немає потреби в розробці складної системи AI; 2) здатні виконувати різноманітні завдання, пов'язані з обробкою текстової інформації; 3) надають швидкі відповіді, що робить їх оптимальними для використання в режимі реального часу.

Істотні ризики пов'язані з контролем якості: згенерований код часто містить дублювання, неефективні структури, пропущені валідації, а вразливості можуть бути внесені разом із шаблонними рішеннями. Є ризик появи технічного боргу, який важко виявити й усунути без залучення досвідчених фахівців. Крім того, LLM-агенти залежать від актуальності моделі й доступності API – збої або втрати зв'язку з платформою можуть зупинити весь процес розробки. Також слід зазначити питання безпеки: за неналежного налаштування агент може внести сторонній код, залишити тестові бекдори, або помилково розкрити чутливі дані у відкритому доступі. До недоліків використання агентів можна віднести: 1) невпевненість або нестійкість під час ведення тривалих діалогів; 2) спроможні генерувати різні варіанти відповідей на один і той самий запит; 3) позбавлені передбачуваності типових алгоритмів, що може ускладнити процес контролю; 4) часте звернення до LLM може бути затратним і призводити до затримок у роботі [3].

Для підвищення якості сучасні команди інтегрують у пайплайн додаткові перевірки: автоматичні статистики аналітики (наприклад, SonarQube), fuzzing-тести, верифікацію ліцензій сторонніх бібліотек, динамічне тестування продуктивності. Часто агенти комбінуються із внутрішніми репозитаріями шаблонів та best practices, що дозволяє формувати більш стандартизований і безпечний код. З боку бізнесу велику роль відіграє інтеграція агента з системами аналітики (BI/ETL), які одразу після релізу дозволяють збирати метрики використання й отримувати інсайти для наступних ітерацій розвитку.

Як приклад роботи LLM-агентів більш детально розглянемо систему автоматичної генерації проєктів із залученням LLM-агентів. Дана система має продуману структуру, чітку взаємодію компонентів і реалістичний підхід до розмежування відповідальностей між класами, агентами та користувачами. В основі архітектури в даному випадку лежить інтеграція Django як бекенду, React – як фронтенду та набору агентів, що реалізують різноманітну бізнес-логіку. Уся система організована пакетами, де кожен пакет представляє певний функціональний домен: користувачі, чати, генерація, налаштування, маршрутизація, компонентна логіка фронтенду тощо. Такий розподіл дозволяє швидко масштабувати та супроводжувати рішення, а також впроваджувати окремі поліпшення, не впливаючи на інші частини проєкту.

У пакеті Accounts визначається клас CustomUser, що містить не лише ідентифікаційні атрибути, а й зберігає абсолютний шлях до згенерованого проєкту користувача та статус активності. Це дозволяє системі ефективно відстежувати сесію, керувати правами доступу та виконувати персоналізацію бізнес-логіки для кожного користувача. Пакет Chats містить клас PlanningChat, який працює з історією взаємодії у JSON-форматі та має пряме відношення до користувача, що створює або володіє чатом. Це дозволяє фіксувати всю логіку спілкування з LLM-агентом, підтримувати чат-контекст і забезпечувати зручне повернення до історії завдань.

Найсильнішою стороною реалізації даної системи є структурована ієрархія агентів у пакеті Generation. Базовий клас BaseAgent відповідає за загальні функції взаємодії з OpenAI API, опрацювання шляхів до проєктів та виконання типових запитів. На основі базового агента будуються спеціалізовані агенти, кожен з яких має чітко окреслені обов'язки: FunctionalAgent реалізує запис технічних завдань та результатів генерації у файл; NonFunctionalAgent здійснює перевірку чи валідацію плану; DoerAgent запускає механізм безпосередньої генерації та реалізації плану; InspectorAgent фокусується на інспекції згенерованого коду; UpdaterAgent займається оновленням існуючих проєктів та завдань. Такий підхід забезпечує максимальну гнучкість і розширюваність, а також дозволяє швидко додавати нових агентів із унікальними сценаріями використання, не змінюючи фундаментальну логіку системи.

Пакет Django Project містить компоненти для налаштувань (Settings) та маршрутизації (URLs), що забезпечує централізоване управління параметрами застосунку та спрощує масштабування. На фронтенді, у пакеті React Project, виділені ключові компоненти: AuthContext є ядром логіки автентифікації, ChatBox відповідає за організацію чат-сесії з агентом, UpdatingBox дозволяє користувачу оперативно отримувати інформацію про процес оновлення або результатів генерації. Це дозволяє створити зрозумілий, швидкий і гнучкий клієнтський інтерфейс для роботи розробника чи кінцевого користувача.

Загалом вся система спроектована так, що користувачі взаємодіють із платформою через інтуїтивно зрозумілий вебінтерфейс. Типовий сценарій роботи може бути визначеним наступним чином: 1) користувач авторизується, 2) створює новий проєкт, 3) формулює вимоги, які передаються одним з агентів на бекенді, 4) формується план, 5) генерується код, 6) результати надходять у чат, де користувач може поставити додаткові запитання чи скорегувати завдання. Якщо потрібно оновити вже існуючий проєкт, UpdaterAgent приймає зміни, перевіряє актуальність, вносить корекції та повертає оновлений код на фронтенд. Вся взаємодія супроводжується чіткими статусами, повідомленнями про успішність чи помилки, а також підтримкою журналу подій для адміністрування.

З точки зору бази даних, дана структура побудована під максимальне розширення. Окрім типових таблиць для користувачів (CustomUser), управління сесіями (UserSession) та збереження проєктів (GeneratedProject), ведеться історія чатів (PlanningChat), що дозволяє швидко відновлювати контекст навіть при повторному запуску агента або після втрати з'єднання. Для забезпечення безпеки використовується сучасна система автентифікації на основі JWT, що дозволяє зручно масштабувати рішення з точки зору багатокористувацької роботи та захищеного обміну даними між фронтендом і бекендом.

Базова логіка агентів передбачає обробку помилок на стороні API, повторні спроби, ведення журналу використаних токенів для оптимізації витрат на AI-запити, а також підтримку декількох моделей для порівняння якості результатів генерації. Це дозволяє не лише автоматизовано створювати проєкти, а й оперативно реагувати на збої або обмеження з боку зовнішніх провайдерів.

Всі ключові сценарії використання системи деталізуються у вигляді usecase-описів та діаграм. Так, наприклад, сценарій створення нового проєкту охоплює весь ланцюг дій: від аутентифікації користувача, збору вимог і генерації плану – до перевірки коду й повідомлення про завершення. Оновлення проєкту передбачає як типові редагування, так і альтернативні сценарії з автоматичним виправленням помилок або відхиленням через недоступність зовнішніх сервісів. Система підтримує поділ на різні ролі користувачів: розробник зосереджений на генерації та редагуванні проєкту, адміністратор здійснює контроль, моніторинг і налаштування параметрів, а також аналізує логи взаємодій.

Щодо деталізації життєвих циклів об'єктів через state- і activity-діаграми, то вони чітко описують послідовність подій при генерації, оновленні чи управлінні проєктом. Це дозволяє виявляти вузькі місця, оптимізувати ключові процеси та гарантувати високий рівень якості кінцевого продукту, зменшуючи ризик виникнення помилок ще на етапі проєктування.

Всі взаємозв'язки в системі ілюструються через діаграми класів, компонентів, розгортання та пакетів. Це дає чітке уявлення про залежності між модулями, коректно організовану архітектуру та можливість швидкої заміни технологічних стеків у майбутньому, без втрати цілісності рішень.

Таким чином, розглянута структура системи демонструє сучасний підхід до автоматизації розробки через LLM-агентів. Використання чітких інтерфейсів, обґрунтоване розділення відповідальностей, підтримка розширення та гнучкості у взаємодії між компонентами роблять цю платформу максимально ефективною для задач генерації та супроводу програмних проєктів різної складності. Фокус на автоматизації рутинних та складних процесів, супроводжуваність через систематизовану документацію, чіткі діаграми та структуровану основу дозволяють не лише реалізувати сьогоденні бізнес-вимоги, а й гарантують легку адаптацію до майбутніх змін технологічного ландшафту.

Підсумовуючи, автоматична генерація проєктів за допомогою LLM-агентів найбільш ефективна там, де потрібна швидка розробка, гнучкість та інтеграція багатьох технологічних компонентів. Для задач із чіткими алгоритмами та фіксованими вимогами краще застосовувати прості підходи – вузькі LLM-запити або класичні скрипти – адже вони споживають менше ресурсів і забезпечують вищу передбачуваність. Вибір між агентом, LLM чи традиційним рішенням має базуватись на детальному аналізі задач і ресурсів, а не на загальному хайпі. Тільки тоді штучний інтелект стане не просто ще одним модним засобом, а справді продуктивним і безпечним інструментом.

References

1. Liu, J. (2024). Large Language Model-Based Agents for Software Engineering: A Survey. arXiv:2409.02977 [cs.SE].
2. Tan, S. H. (2024). Automatic Programming: Large Language Models and Beyond. arXiv:2405.02213 [cs.SE].
3. Sergii Cherbazhy. AI-агенти vs. LLM vs. Алгоритми. [DOUDay форум] Retrieved from: <https://dou.ua/forums/topic/52895/>

DOI 10.34132/mspc2025.01.14.08

Yelisieiev Oleksandr
Horban Hlib

AUTOMATIC PROJECT GENERATION WITH THE HELP OF LLM AGENTS

The thesis explores the use of LLM agents for the automatic generation of software projects, analyzing their capabilities, advantages, and risks compared to traditional algorithms and simple LLM queries. The author discusses the architecture of agents, their integration with modern technologies (Django, React, OpenAI API), and application scenarios in business, education, and financial technologies. The system of automatic project generation involving LLM agents is considered. The importance of making an informed choice between complex agents and simple solutions depending on project requirements to ensure efficiency, security, and minimize technical debt is emphasized.

Keywords: LLM agents, large language models, automatic code generation, artificial intelligence, software development, Django, React, OpenAI API.

Відомості про автора / Information about the Author

Єлісеєв Олександр, здобувач першого рівня вищої освіти (бакалавр) зі спеціальності 121 «Інженерія програмного забезпечення», Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: alexander0402112@gmail.com.

Yelisieiev Oleksandr, Student, Specialty 121 "Software Engineering", Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: alexander0402112@gmail.com.

Горбань Гліб, кандидат технічних наук, доцент, доцент кафедри інженерії програмного забезпечення, Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: hlib.horban@chmnu.edu.ua, ORCID: <https://orcid.org/0000-0002-6512-3576>

Horban Hlib, PhD, Associate Professor of Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: hlib.horban@chmnu.edu.ua, ORCID: <https://orcid.org/0000-0002-6512-3576>

DOI 10.34132/mspc2025.01.14.09

Ірина Калініна,
Олександр Гожий

ПОБУДОВА БАГАТОРІВНЕВИХ АНСАМБЛЕЙ В ЗАДАЧАХ МАШИННОГО НАВЧАННЯ

Анотація. Розглянуто підхід до розв'язання задач машинного навчання, побудований на основі багаторівневих гетерогенних ансамблів. Підхід дозволяє зменшувати помилки прогнозів за рахунок одночасного зменшення зміщення та дисперсії на основі використання багаторівневих гетерогенних ансамблів прогнозних моделей. Розглянуто основні особливості та передумови створення ансамблів. Проведено аналіз загальної помилки алгоритму машинного навчання. Вона складається з трьох компонентів: шум, зміщення та дисперсія. Досліджено та визначено складові цих компонентів. Розглянуто процес створення гетерогенного ансамблю. Проаналізовано найбільш поширені методи агрегування прогнозних значень: *bagging*, *boosting* і *staking*. Подано обґрунтування вибору типу моделей для створення гетерогенної багаторівневої ансамблевої структури. Запропоновано дворівневу архітектуру системи класифікації на основі методів *staking* та *bagging*. Для побудови багаторівневого ансамблю моделей з урахуванням різних базових методів розроблено алгоритм. Проаналізовано результати роботи ансамблів. Доведено ефективність багаторівневих гетерогенних ансамблів при вирішенні завдань класифікації.

Ключові слова: Машинне навчання, багаторівневі гетерогенні ансамблі, зміщення, дисперсія, *bagging*, *boosting*, *staking*, дворівнева архітектура ансамблю.

Вступ

Останнім часом спостерігається розвиток різних підходів до вирішення завдань машинного навчання, які призвели до дослідження та створення нових моделей у різних галузях, таких як охорона здоров'я, розпізнавання мовлення, класифікація зображень, прогнозування та багато інших. Один із підходів, який поєднує декілька різних моделей для створення остаточної моделі відомий, як ансамблеве навчання або ансамблева модель. Підхід на основі ансамблів моделей успішно застосовують у широкому спектрі завдань машинного навчання. Ефективність цих моделей пояснюється найкращим поданням ознак за допомогою багатопарових архітектур обробки даних. Ансамблеві моделі в основному використовуються для задач класифікації, регресії та кластеризації. Ансамбль – це набір класифікаторів/регресорів, які вивчають цільову функцію, та їх індивідуальні прогнози поєднуються для класифікації нових прикладів. Ансамбль поєднує в собі кілька моделей, певним чином, наприклад, усередненням або голосуванням, таким чином, щоб отримати нову модель, яка має кращі показники, ніж моделі, що складають ансамбль. Оскільки моделі машинного навчання є обчислювальними та великими за даними, отже, ансамблеві моделі вимагають особливої уваги щодо додаткової інформації для об'єднання кількох алгоритмів у єдиній структурі. Ансамблеві моделі машинного навчання повинні вирішувати кілька завдань, таких як викликати різноманітність серед базових моделей, як зберегти час навчання, а також зниження складності моделей для практичних додатків, як об'єднати передбачення на основі додаткових алгоритмів. Створення ефективних ансамблевих моделей дозволить суттєво підвищити точність розв'язання задач машинного навчання.

Постановка задачі. Особливий інтерес для досліджень представляє створення багаторівневих гетерогенних ансамблів, які здатні істотно підвищити ефективність вирішення завдань машинного навчання. Створення таких ансамблів розглядатиметься у цій роботі.

Методи та моделі

Ефективним підходом до підвищення якості прогнозних значень при розв'язанні завдань класифікації є підхід до прогнозування на основі ансамблевих методів. Він дозволяє зменшувати помилки прогнозів за рахунок компромісу між зміщенням та дисперсією на основі використання багаторівневих гетерогенних ансамблів прогнозних моделей.

При реалізації ансамблевого підходу слід враховувати такі особливості:

- Ансамблі, як правило, покращують продуктивність узагальнення набору лише окремих моделей у домені.
- Велика кількість відносно некорельованих моделей, які працюють як домен, перевершують будь-яку з окремих складових моделей.

Для використання ансамблів є ряд передумов: набір класифікаторів зі схожими результатами навчання може мати різні результати узагальнення, об'єднання вихідних даних кількох класифікаторів знижує ризик вибору погано працюючого (слабкого) класифікатора, якщо обсяг даних для аналізу занадто великий, один класифікатор може не впоратися з ним; необхідно навчати різні класифікатори різних розділах даних. Системи ансамблів також можна використовувати, коли даних замало, за допомогою методів повторної вибірки.

Помилка алгоритму машинного навчання складається з трьох компонентів: шум, зміщення та дисперсія:

$$Error(x) = Bias(x)^2 + Variance(x) + Noise(x),$$

де *Noise* – це непереборна помилка (випадкові помилки в даних, які неможливо усунути, наприклад, через пошкоджені вхідні дані, помилки введення даних і т. д.); *Bias* – це систематична помилка, яку, як очікується, зробить алгоритм навчання через, наприклад, архітектурний вибір або через недостатні/нерепрезентативні навчальні дані; *Variance* вимірює чутливість алгоритму до конкретного навчального набору та/або гіперпараметрів, що використовуються.

При добірї оптимальної прогнозової моделі необхідно враховувати компроміс між зміщенням і дисперсією. Техніка ансамблювання, тобто агрегування прогнозних значень різних базових моделей для створення однієї «оптимальної» є прийомом, що дозволяє пом'якшити компроміс між зміщенням та дисперсією. На рисунку 1 зображено умовний компроміс між зміщенням та дисперсією. Таким чином, ансамблі можуть допомогти зменшити як зміщення, так і дисперсію (за винятком шуму, який є непереборною помилкою).

На практиці ідея полягає в наступному: якщо використовуються «прості» моделі, навчені на менших вибірках, індивідуальна дисперсія кожної гіпотези може бути вищою, ніж для складнішого «учня», але все ж таки в більшості випадків дисперсія ансамблю нижче, ніж при використанні одного навченого складного.

Створення гетерогенного ансамблю

Застосування ансамблю моделей – це процес, у якому різні та незалежні моделі поєднуються для отримання кращого результату. При побудові ансамблю використовують кілька технік агрегування результатів базових моделей, кожен із яких забезпечує різну точність діагностики. Розглядалися три найбільш поширені методи агрегування прогнозних значень: *bagging*, *boosting* і *steking*. У таблиці 1 наведено обґрунтування вибору типу моделей для створення гетерогенної багаторівневої ансамблевої структури.

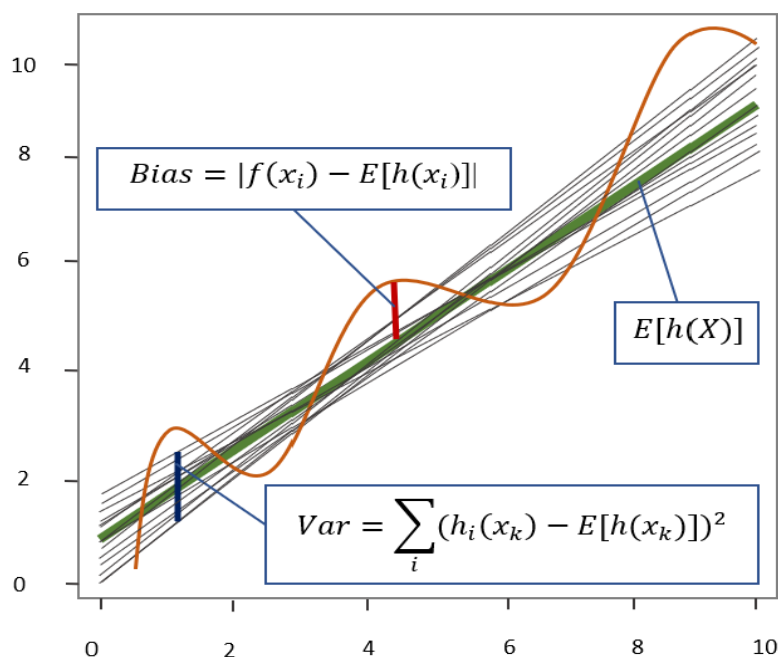


Рис.1 Компроміс між зміщенням та дисперсією.

Джерело: сформовано авторами

Таблиця 1

Обґрунтування вибору типу моделей для створення ансамблю

Тип ансамблю	Тип базових моделей	Характер помилки моделі	Результат агрегації
<i>Bagging</i> (модельне усереднення)	однорідні моделі	моделі з низьким зміщенням та високою дисперсією (перенавчені моделі)	зменшує дисперсію
<i>Boosting</i> (модельне підсилювання)	однорідні моделі	моделі з низькою дисперсією та високим зміщенням (недонавчені моделі)	зменшує зміщення
<i>Stacking</i> (модельне накладення)	різнорідні моделі	моделі з низькою дисперсією та високим зміщенням (недонавчені моделі)	зменшує зміщення, але залежно від вибору метамоделі може зменшувати дисперсію

Після експериментального підбору типу ансамблю розроблено дворівневу архітектуру системи класифікації на основі методів *steking* та *bagging* (рис. 2). Дворівневий класифікатор на обох рівнях агрегації використовує виважену суму. Ваги розраховуються на основі значення *F*-міри кожного з базових методів.

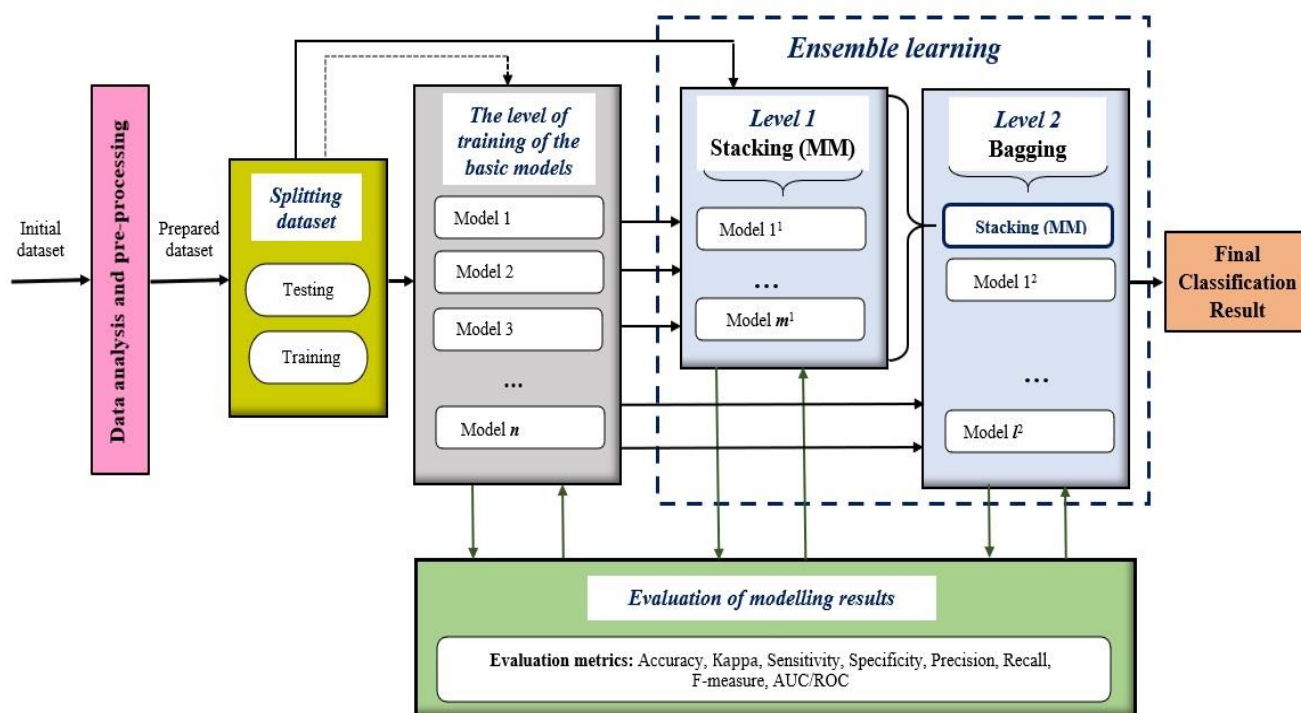


Рис.2. Дворівнева архітектура системи класифікації на основі методів *Steking* та *Bagging*.
 Джерело: сформовано авторами

Для побудови багаторівневого ансамблю моделей з урахуванням різних базових методів розроблено алгоритм. Він складається із 7 кроків:

Крок 1. Поділ підготовленого до моделювання набору даних на дві вибірки (навчальну та тестову).

Крок 2. Навчання *n* різнорідних базових моделей на певній навчальній вибірці. Оцінка якості моделей. Вибір оптимальних параметрів моделей кожного типу.

Крок 3. Розрахунок кожної з *n* моделей прогнозових значень. Оцінка якості прогнозів за допомогою тестової вибірки.

Крок 4. Обчислення (оцінка) дисперсії та усунення кожної з *n* базових моделей.

Крок 5. Поділ n базових моделей на дві групи (з високим усуненням – недоучені та з високою дисперсією – перевчені).

Крок 6. Формування з урахуванням $m(m < n)$ моделей першої групи 1-го рівня ансамблевої структури – *steking*.

Крок 7. Формування з урахуванням $l(l \leq n-m)$ моделей другої групи + результат методу *steking* 2-го рівня ансамблевої структури – *bagging*. Результат вважатиме остаточною прогнозом значенням.

Результати роботи класифікаторів

Головне з переваг дворівневого ансамблю полягає у системності використання методів агрегування та підбору базових моделей класифікації для кожного рівня ансамблевого навчання. Більшість моделей мають задовільну якість результату, але результуючі значення потребують покращення. Тому для підвищення якості результату будується дворівнева ансамблева модель. Як вхідні дані вибираються тестові оцінки якості базових моделей і додаються до тестового набору даних.

Розроблена система класифікації, заснована на багаторівневому ансамблевому підході, була перевірена на двох наборах даних. Перший набір даних містить інформацію про 748 донорів крові з Центру служби переливання крові в місті Синьчу, Тайвань. Метою класифікації є визначення кількості добровольців, які, швидше за все, знову здадуть кров. Другий набір даних (Indian Liver Patient Dataset) був підібрано для дослідження пацієнтів із захворюваннями печінки. Метою класифікації є виявлення потенційних донорів печінки, що сприяє більш ефективному та справедливому розподілу донорських органів серед пацієнтів, які очікують трансплантації.

У блоці дворівневого ансамблевого навчання послідовно використовуються *steking* та *bagging*. На першому рівні працює метод *steking* на основі моделі логістичної регресії. Модельний *steking* є ефективним методом ансамблювання. Прогнози, сформовані з допомогою базових моделей, використовуються як вхідні дані на навчання першому рівні.

Як базові моделі першого рівня були обрані: дерева рішень (DecisionTree); наївний Байєсов класифікатор (NB); лінійний дискримінантний аналіз (LDA); логістична регресія (LR); метод найближчого сусіда (KNN); штучні нейронні мережі (ANN).

Перший рівень: *Stacking*. Основні відомості ансамблю моделей першого шару: логістична регресія як метамодель, 10 предикторів, два класи результуючої змінної (*no* та *yes*).

На другому рівні працює метод *bagging* на основі алгоритму *Bagged CART*. Алгоритм створює N регресійних дерев, використовуючи M початкових навчальних наборів та усереднює отримані прогнози. Кожне дерево має високу дисперсію, але низьку зміщення. Усереднення дерев зменшує дисперсію. Прогнозованими значеннями спостережень є мода (класифікація) чи середнє значення (регресія) дерев. Одним із недоліків *Bagged Trees* є те, що невелика кількість додаткових навчальних спостережень може різко змінити ефективність передбачення навченого дерева. Як базові моделі другого шару вибрано: модель першого рівня (*Stacking LR*); модель випадкового лісу (*RF*); квадратичний дискримінантний аналіз (*QDA*).

Представлені значення показників ефективності першого та другого наборів даних. Ці дані враховують здатність моделей відрізнити один клас від інших (точність прогнозування, каппа-статистика, чутливість та специфічність, точність та повнота, F -міра). Використання дворівневої архітектури гетерогенного ансамблю моделей послідовно підвищує точність розв'язання задачі класифікації в машинному навчанні по всім метрикам якості.

Висновки

Для практичної реалізації багаторівневих гетерогенних ансамблів для вирішення задач класифікації розглянуто два набори даних: Blood Transfusion Service Center та ILPD (Indian Liver Patient Dataset). Для оцінки якості класифікаторів було використано систему показників якості. Проаналізовано результати роботи ансамблів. У таблицях представлені значення показників ефективності першого і другого наборів даних. Ці дані враховують здатність моделей відрізнити один клас від інших (точність прогнозування, каппа-статистика, чутливість та специфічність, точність та повнота, F -міра). Для першого набору даних на першому етапі в результаті об'єднання деяких базових моделей точність підвищена до 70%, а на другому етапі об'єднання деяких базових моделей та моделі *steking* в ансамбль моделей за допомогою *bagging* підвищило точність моделі до 88%. Для другого набору даних на першому етапі в результаті об'єднання базових моделей точність підвищена до 77%, а на другому етапі об'єднання деяких базових моделей та моделі *steking* в ансамбль моделей за допомогою *bagging* підвищило точність моделі до 82%. Таким чином використання гетерогенного дворівневого ансамблю значно підвищує ефективність класифікаційних моделей.

References

1. P. Bidyuk, I. Kalinina, O. Zhebko, A. Gozhyj and T. Hannichenko, "Classification System Based on Ensemble Methods for Solving Machine Learning Tasks". CEUR- WS. vol. 3426, 2023, pp. 1-11. CEUR-WS.org/Vol-3426/paper5.pdf.
2. V. Pandey, "Ensemble Methods in Practice: Combining the Strengths of Multiple Models and Making Decisions," 2023. [Online]. Available: <https://www.linkedin.com/pulse/ensemble-methods-practice-combining-strengths-multiple-pandey> (application date: 15.08.2024).

DOI 10.34132/mspc2025.01.14.09

Iryna Kalinina, Oleksandr Gozhyj

CONSTRUCTION OF MULTILEVEL ENSEMBLES IN MACHINE LEARNING PROBLEMS

The approach to solving the problems of machine learning and motivation on the basis of rich heterogeneous ensembles is examined. The approach allows you to change forecast estimates based on one-hour changes in variance and variance based on a variety of diverse heterogeneous ensembles of forecast models. The main features and changes in the creation of ensembles are examined. An analysis was carried out for the purpose of the machine learning algorithm. It consists of three components: noise, displacement and dispersion. The warehouse components have been tracked and identified. The process of creating a heterogeneous ensemble is examined. The most extensive methods of aggregation of forecast values have been analyzed: bagging, boosting and staking. The choice of type of models for the creation of a heterogeneous rich ensemble structure is given. The courtyard architecture of the classification system based on the methods of Staking and Bagging has been proposed. To create a rich ensemble of models based on various basic methods, an algorithm has been developed. The results of the robotic ensembles were analyzed. The effectiveness of a large number of heterogeneous ensembles with the highest specification of classification has been demonstrated.

Key words: Machine learning, rich heterogeneous ensembles, displacement, dispersion, bagging, boosting, staking, courtyard architecture ensemble.

Відомості про авторів/ Information about the Authors

Ірина Калініна, доктор технічних наук, професор кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: irina.kalinina1612@gmail.com, orcid: <https://orcid.org/0000-0001-8359-2045>

Iryna Kalinina, Doctor of Technical Sciences, Professor of the Department of Intellectual Information Systems, of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: irina.kalinina1612@gmail.com, orcid: <https://orcid.org/0000-0001-8359-2045>

Олександр Гожий, доктор технічних наук, професор кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: alex.gozhyj@gmail.com, orcid: <https://orcid.org/0000-0002-3517-580X>.

Oleksandr Gozhyj, Doctor of Technical Sciences, Professor of the Department of Intellectual Information Systems, of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: alex.gozhyj@gmail.com, orcid: <https://orcid.org/0000-0002-3517-580X>.

© I. Калініна, О. Гожий

Стаття отримана редакцією 19.05.2025.

DOI 10.34132/mspc2025.01.14.10

Назарій Кіріак
Гліб Горбань

ІНФОРМАЦІЙНА СИСТЕМА ЕЛЕКТРОННОЇ КОМЕРЦІЇ НА ОСНОВІ СУЧАСНИХ ВЕБТЕХНОЛОГІЙ

Тези присвячені теоретико-прикладному аналізу розробки вебзастосунків як ключового інструменту цифрової трансформації комерційного сектора в Україні. Електронна комерція розглядається як рушійна сила розвитку цифрової економіки, що забезпечує новий рівень взаємодії між бізнесом і споживачем, сприяє прозорості бізнес-процесів, автоматизації торгівлі та розширенню доступу до товарів і послуг.

Особлива увага приділяється технологічним аспектам реалізації вебплатформ, їх ролі в забезпеченні безпеки даних, масштабованості систем, інтеграції з платіжними сервісами та державними реєстрами. Підкреслюється, що впровадження електронної комерції значно підвищує ефективність підприємницької діяльності та сприяє формуванню конкурентоспроможного національного ринку в умовах глобалізації.

Ключові слова: електронна комерція, вебзастосунок, цифрова трансформація, інформаційні технології, інтернет-торгівля, UI/UX, платіжні системи, безпека даних.

Сучасні вебзастосунки у сфері електронної комерції відіграють ключову роль у цифровій трансформації торгівлі, дозволяючи компаніям оперативнo взаємодіяти з клієнтами, автоматизувати бізнес-процеси та розширювати ринки збуту. Вебзастосунки електронної комерції є комплексними програмними рішеннями, які об'єднують інтерфейс користувача, бізнес-логіку та інструменти керування товарами, замовленнями, платежами та доставкою.

Для реалізації сучасного вебзастосунку електронної комерції широко використовуються стек технологій, орієнтований на гнучкість, масштабованість і безпеку. Зокрема, фронтенд розробляється з використанням React – популярної бібліотеки JavaScript, яка забезпечує високу динамічність та інтерактивність інтерфейсу користувача [1]. Завдяки компонентному підходу та віртуальному DOM React дозволяє ефективно створювати швидкодіючі вебінтерфейси, які адаптуються під різні пристрої.

Як бекенд переважно обираються фреймворки на основі Java Spring або Node.js, які дозволяють побудувати надійне REST API з підтримкою авторизації, обробки платежів, ведення кошика, списків бажань тощо [2]. Крім того, використання WebSocket або Server-Sent Events забезпечує реальний час для таких функцій, як повідомлення про зміну статусу замовлення, підтвердження транзакції, онлайн-підтримка.

Одним із ключових елементів є база даних, яка виконує роль централізованого сховища для інформації про користувачів, товари, замовлення та транзакції. Для цього часто використовуються реляційні СУБД, як-от MySQL або PostgreSQL, або нереляційні рішення (наприклад, MongoDB) у разі потреби високої гнучкості в структурі даних [3].

Системи електронної комерції повинні відповідати сучасним вимогам до безпеки – захист персональних даних, шифрування з'єднань, підтримка безпечної авторизації через OAuth 2.0, інтеграція з платіжними шлюзами (LiqPay, Stripe, PayPal тощо), а також захист від SQL-ін'єкцій, CSRF і XSS-атак [4].

У сучасних умовах глобалізації електронна комерція є не лише зручною формою ведення бізнесу, а й потужним драйвером інноваційних змін в IT-сфері. За даними аналітичної компанії Statista, обсяг глобального ринку e-commerce у 2024 році перевищив 6 трильйонів доларів США, а очікуване зростання до 2027 року становить ще понад 50%. При цьому понад 70% онлайн-замовлень здійснюється через мобільні пристрої, що зумовлює впровадження mobile-first підходу до розробки вебзастосунків. Це стимулює створення інтерфейсів, які однаково добре працюють на смартфонах, планшетах та десктопах, з орієнтацією на зручність, швидкість і мінімалізм.

Ключовими трендами в розвитку вебзастосунків для електронної комерції є інтеграція штучного інтелекту (AI) та машинного навчання (ML) для персоналізації контенту, аналізу поведінки користувачів і автоматизації

маркетингу. Все ширше застосовуються чат-боти для обслуговування клієнтів у режимі 24/7, що значно знижує навантаження на операторів і підвищує якість сервісу.

Окрему нішу займають Progressive Web Apps (PWA) – прогресивні вебзастосунки, що поєднують зручність нативних мобільних додатків із гнучкістю звичайних сайтів. Вони забезпечують миттєве завантаження, роботу офлайн, push-сповіщення та кращу продуктивність, що особливо важливо для ринків із нестабільним інтернетом або обмеженими ресурсами пристроїв.

Також стрімко розвивається інтеграція вебзастосунків з хмарними платформами, такими як AWS, Google Cloud та Microsoft Azure, що дозволяє масштабувати рішення відповідно до потреб бізнесу та забезпечує високу доступність даних. Усе це формує нову цифрову реальність, де швидкість адаптації до технологічних змін визначає конкурентоспроможність компанії.

Перевагою вебзастосунків для e-commerce є їх доступність – як для малого бізнесу, так і для великих компаній. Вони дозволяють автоматизувати продажі 24/7, забезпечують аналітику поведінки користувачів, інтегруються з CRM, ERP та логістичними системами. Інструменти типу Google Analytics, Facebook Pixel, а також внутрішні панелі адміністратора допомагають відстежувати дії користувачів, керувати товарами, знижками, рекламними кампаніями.

Порівняно з платформами типу Shopify або Wix, власноручна розробка вебзастосунку дає більше контролю, розширюваність і налаштування, але потребує глибших технічних знань. З іншого боку, застосування відкритих рішень, таких як WooCommerce або Magento, дозволяє значно пришвидшити запуск проєкту, однак вимагає адаптації під специфіку бізнесу.

Суттєвим фактором успішності вебзастосунків у сфері електронної комерції є якість користувацького досвіду (UX) та інтерфейсу користувача (UI). Навіть за наявності функціонального функціоналу, незручний або перевантажений інтерфейс може відштовхнути потенційного покупця. Тому сучасні вебзастосунки розробляються з урахуванням принципів UX-дизайну: простоти навігації, логічного розташування елементів, швидкості завантаження сторінок, адаптивності до різних пристроїв і психологічного комфорту користувача. Крім того, використовується A/B тестування для визначення оптимальних рішень щодо структури та дизайну сторінок.

Значну роль у підтримці бізнес-процесів відіграє інтеграція з аналітичними сервісами (рис. 1). Такі інструменти, як Google Analytics, Hotjar, Facebook Pixel, дозволяють відстежувати поведінку користувачів, джерела трафіку, глибину перегляду, конверсії тощо. Це дає змогу формувати таргетовані маркетингові кампанії та оперативно реагувати на потреби ринку. На основі зібраних даних активно впроваджується персоналізація контенту – наприклад, система може рекомендувати товари на основі попередніх переглядів або купівельної історії користувача.

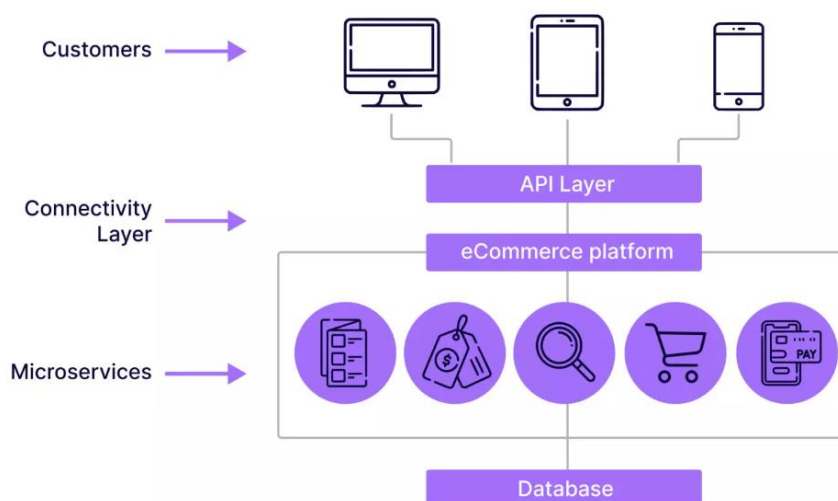


Рисунок 1 – Технологія інтеграції з аналітичними сервісами

Технічна інфраструктура вебзастосунків для e-commerce вимагає стабільної роботи під навантаженням, зокрема у періоди пікового попиту (розпродажі, свята, «чорна п'ятниця»). З цією метою широко впроваджуються хмарні сервіси для автоматичного масштабування, балансування навантаження та забезпечення відмовостійкості.

Окрім AWS, Google Cloud і Microsoft Azure, використовуються CDN-мережі (Content Delivery Networks), які прискорюють завантаження контенту незалежно від географії користувача.

Окрему увагу заслуговує впровадження технологій розширеної реальності (AR) у вебзастосунки (рис. 2). Це дозволяє користувачам віртуально "приміряти" товар, розмістити його у своєму інтер'єрі або оглянути в 3D, що підвищує довіру до онлайн-покупок і зменшує кількість повернень.

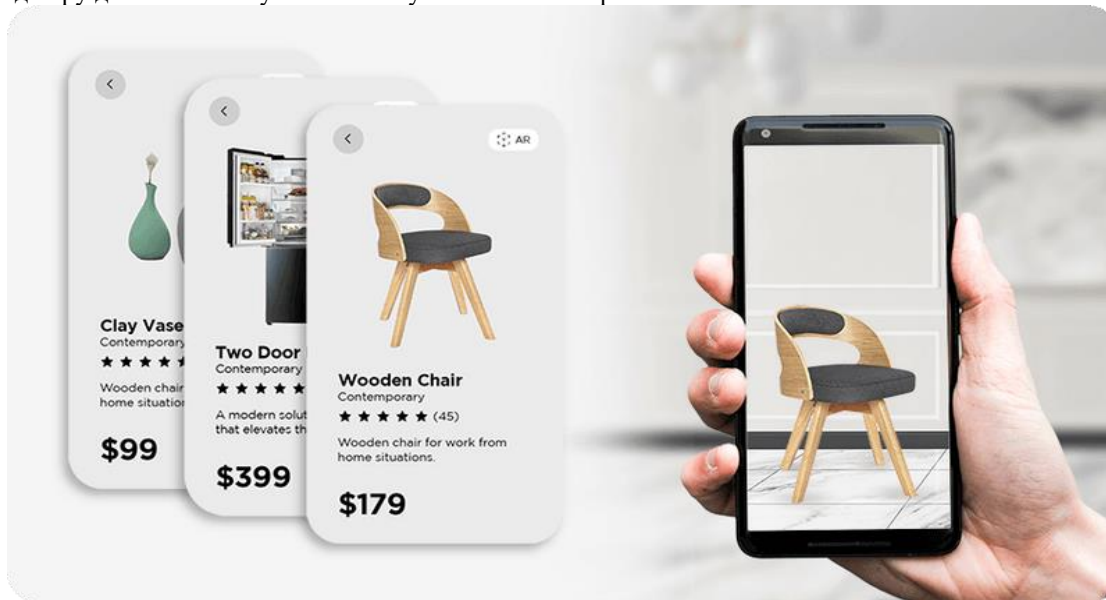


Рисунок 2 – Використання AR технологій розширеної реальності

Серед недоліків можна відзначити високу складність первинної реалізації: потрібно продумати логіку обробки кошика, повернень, промокодів, складів, обліку товарів, конфігурацій товарів (розмір, колір, модель тощо). Також підтримка та масштабування системи з часом потребують значних ресурсів, особливо за умови високого навантаження або сезонних піків.

Зважаючи на актуальність електронної комерції у постпандемічному світі, збільшення обсягу онлайн-продажів та цифрових платежів, створення сучасного вебзастосунку для e-commerce є перспективним напрямом як для бізнесу, так і для розробників. Це дозволяє створювати високофункціональні продукти, які поєднують інтуїтивно зрозумілий інтерфейс з потужною серверною частиною та аналітикою в реальному часі.

Таким чином, вебзастосунок електронної комерції є складною, проте необхідною інвестицією в цифрову стратегію бізнесу, яка забезпечує конкурентні переваги, зручність для клієнтів та аналітичну гнучкість для прийняття рішень. Його успішна реалізація вимагає міждисциплінарного підходу, поєднання технічних, UX/UI, бізнес- та безпекових рішень, що робить даний напрямок одним із ключових у сфері IT-розробки.

References

1. Grider S. Modern React with Redux. San Francisco: Udemy Publishing, 2022. 420 p.
2. Johnson R., Hoeller J. Professional Java Development with the Spring Framework. Indianapolis: Wrox Press, 2005. 720 p.
3. Chodorow K. MongoDB: The Definitive Guide. 3rd ed. Sebastopol: O'Reilly Media, 2019. 500 p.
4. Stuttard D., Pinto M. The Web Application Hacker's Handbook: Finding and Exploiting Security Flaws. 2nd ed. Indianapolis: Wiley Publishing, 2011. 912 p.

INFORMATION SYSTEM OF E-COMMERCE BASED ON MODERN WEB TECHNOLOGIES

The theses present a theoretical and applied analysis of the development of web applications as a key tool for the digital transformation of the commercial sector in Ukraine. E-commerce is viewed as a driving force behind the growth of the digital economy, enabling a new level of interaction between businesses and consumers, promoting transparency in business processes, automating trade, and expanding access to goods and services.

Particular attention is paid to the technological aspects of implementing web platforms, their role in ensuring data security, system scalability, integration with payment services and government registries. It is emphasized that the introduction of e-commerce significantly increases the efficiency of entrepreneurial activity and contributes to the formation of a competitive national market in the context of globalization.

Keywords: *e-commerce, web application, digital transformation, information technology, online trade, UI/UX, payment systems, data security.*

Відомості про автора / Information about the Author

Назарій Кіріак, здобувач першого рівня вищої освіти (бакалавр) зі спеціальності 121 «Інженерія програмного забезпечення», Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: nazar2468123@gmail.com.

Kiriak Nazarii, Student, Specialty 121 "Software Engineering", Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: nazar2468123@gmail.com.

Гліб Горбань, кандидат технічних наук, доцент, доцент кафедри інженерії програмного забезпечення, Чорноморський національний університет ім. Петра Могили, м. Миколаїв, Україна. E-mail: hlib.horban@chmnu.edu.ua, ORCID: <https://orcid.org/0000-0002-6512-3576>

Horban Hlib, PhD, Associate Professor of Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: hlib.horban@chmnu.edu.ua, ORCID: <https://orcid.org/0000-0002-6512-3576>

DOI 10.34132/mspc2025.01.14.11

Олександр Ковалів,
Євген Сіденко,
Галина Кондратенко

ПОРІВНЯЛЬНИЙ АНАЛІЗ РЕЗУЛЬТАТІВ ЗАСТОСУВАННЯ РІЗНИХ АРХІТЕКТУР НЕЙРОННИХ МЕРЕЖ ДЛЯ ПРОГНОЗУВАННЯ ІНФЕКЦІЙНИХ ЗАХВОРЮВАНЬ

У даній роботі представлено результати прогнозування часових рядів із використанням методів машинного навчання та мобільних технологій. Основними результатами цієї роботи є порівняння ефективності різних архітектур нейронних мереж у рамках прогнозування поширення інфекційних захворювань в Україні за період з грудня 2016 року по січень 2024 року.

Для досягнення мети було вирішено такі завдання: проаналізовано поточний стан прогнозування часових рядів; проаналізовано існуючі аналоги систем; обрано необхідні архітектури нейронних мереж, як один із методів машинного навчання; проаналізовано та нормалізовано набір даних про поширення інфекційних захворювань; проведено тестування. Розроблена система може бути ефективним інструментом для прийняття раціональних управлінських рішень щодо забезпечення епідеміологічного благополуччя та біобезпеки населення.

Ключові слова: Прогнозування, часові ряди, методи машинного навчання, нейронні мережі.

Серед надзвичайних ситуацій особливий негативний вплив на здоров'я людини надають війни та військові конфлікти. Війни завжди супроводжуються спалахами інфекційних хвороб, підвищенням рівня захворюваності та смертності від інфекцій. Історія війн продемонструвала, що втрати людських життів від інфекцій часто перевищували втрати від бойових травм. Незважаючи на досягнення сучасної медицини, інфекції продовжують супроводжувати війни у 21 столітті. Тому є важливою задача прогнозування розповсюдження інфекційних захворювань для можливості кращої підготовки до можливих спалахів захворювань [1].

У ході дослідження було проведено систематичний збір даних про поширення інфекційних захворювань в Україні з грудня 2016 року по січень 2024 року. Щомісячні звіти отримували з національних баз даних охорони здоров'я, а саме Центру громадського здоров'я МОЗ України. Дані, що збирає Центр громадського здоров'я, наведено згідно зі звітною формою №1; у звітах наведено абсолютні цифри та інтенсивні показники на 100 000 населення. Дані наведено у порівнянні за певний проміжок часу за поточний і минулий роки, а також за повний рік (12 місяців) [2]. Для врахування коливань чисельності населення протягом досліджуваного періоду дані були нормалізовані за оцінками чисельності населення Державної служби статистики України. Абсолютні значення були нормалізовані для відображення кількості захворювань за поточний місяць на 1 000 000 населення України за поточний рік. Для нормалізації кількості захворювань було використано наступну формулу:

$$\frac{\text{абс. кількість захворювань за поточний місяць} * 1\,000\,000}{\text{кількість населення за поточний рік}} \quad (1)$$

Ця нормалізація дозволила провести аналіз на 1 000 000 населення, полегшивши порівняння за різні місяці та роки. У випадках, коли з'являлися дублікати звітів або суперечливі дані, застосовувався консенсусний підхід, віддаючи пріоритет місяцям із найбільшою кількістю зареєстрованих випадків, щоб забезпечити точність і надійність набору даних. Ця ретельна обробка даних була важливою для забезпечення чіткого та точного представлення тенденцій інфекційних захворювань в Україні протягом зазначеного періоду часу. Всього в наборі даних присутні 57 інфекційних хвороб, зокрема холера, черевний тиф, паратиф А, В, С, шигельоз та інші.

Було побудовано та проаналізовано графіки захворювань з нормалізованими значеннями. Нижче наведено приклад графіків захворювань (рисунок 1).

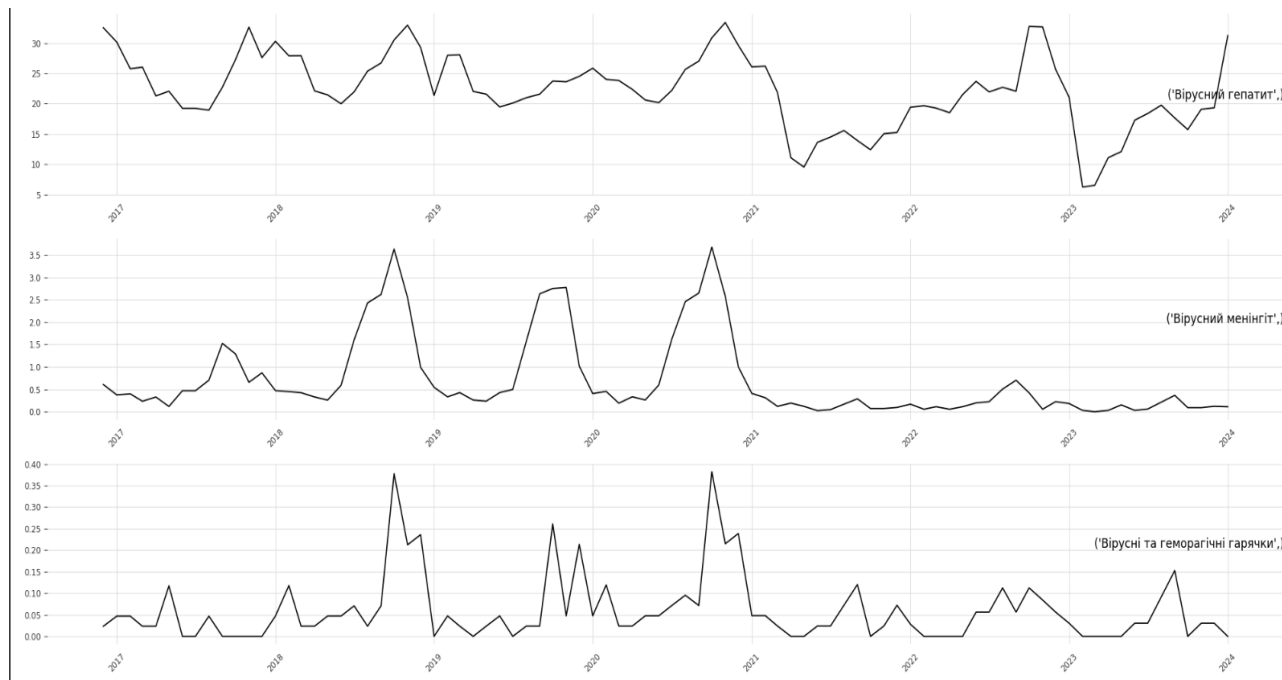


Рис. 1. Графіки розповсюдження інфекційних захворювань згідно датасету.

Джерело: сформовано автором на основі [2]

Серед 57 вибраних інфекційних захворювань спостерігалось збільшення кількості випадків у 22 випадках після початку війни в Україні (24.02.2022), наприклад, випадки лямбліозу із середнім показником до війни 4,53 збільшилися до 10,06, що становить збільшення на 121,91%, кишкові протозойні захворювання із середнім показником до війни 5,06 збільшилися до 11,19, що становить збільшення на 120,9%.

Загалом спостерігалось зменшення кількості випадків у 35 інфекційних захворюваннях із середнім зменшенням кількості випадків на 35,52%, та у 22 захворюваннях із середнім збільшенням кількості випадків на 62,38%.

З огляду на ці цифри, вплив надзвичайних ситуацій очевидний. Масштабна агресія Росії в Україні призвела до значної надзвичайної ситуації, яка вплинула на всі аспекти життя та спричинила додаткові надзвичайні ситуації, такі як повені та посухи внаслідок пошкодження Каховської гідроелектростанції. Інші ризики включають потенційне забруднення радіонуклідами внаслідок аварій на Чорнобильській та Запорізькій АЕС, лісові пожежі на Кінбурнській косі, які спустошили понад 1500 гектарів лісу, та викиди хімічних речовин, включаючи хлор та аміак. У контексті цих надзвичайних ситуацій динаміка та прояви епідемій інфекційних захворювань можуть відрізнятися від їх природного розвитку. Ключовими факторами ризику, що впливають на епідемічну ситуацію на тлі каскаду надзвичайних ситуацій, спричинених війною, є значна міграція населення, переповненість бомбосховищ та міграційних шляхів, підвищений рівень стресу, що призводить до підвищеної вразливості до інфекцій, порушення водо- та енергопостачання, різке зростання популяцій гризунів та пов'язані з цим епізоотії, забруднення харчових продуктів, стікання різних хімічних речовин у водойми, затоплення природних екосистем, активація механізмів передачі інфекційних агентів, збільшення кількості безпритульних тварин та їхньої взаємодії з дикими тваринами, а також забруднення територій внаслідок ракетних та артилерійських ударів та вибухів наземних мін [1].

У ході дослідження для прогнозування та аналізу часових рядів було проаналізовано наступні архітектури нейронних мереж: RNN, Block RNN, N-BEATS, TCN, Transformer Model, N-HITS.

Оскільки нейронна мережа з архітектурою Transformer Model є передовою у вирішенні проблеми прогнозування часових рядів, її було обрано основною для дослідження [3]. На рисунку 2 зображено архітектуру нейронної мережі Transformer Model.

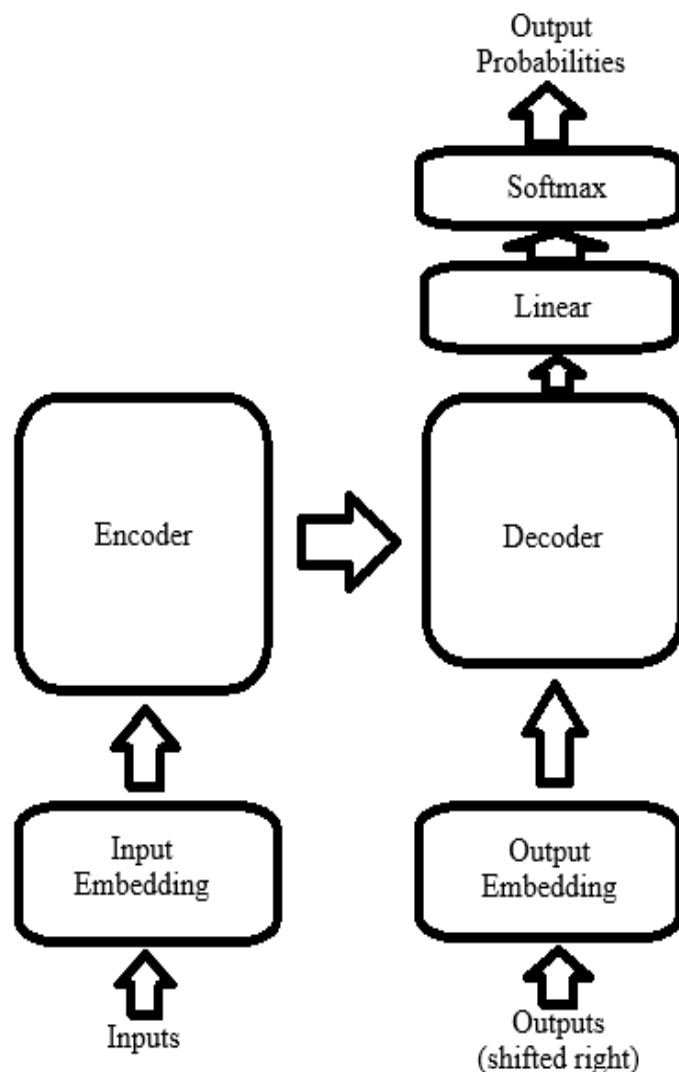


Рис. 2. Архітектура нейронної мережі Transformer Model.
 Джерело: сформовано автором

Нормалізований набір даних про випадки інфекційних захворювань в Україні за період з грудня 2016 року по січень 2024 року послужив основою для навчання нейронних мереж. Ці дані також були нормалізовані в межах від 0 до 1 для покращення навчання Transformer Model. Нейронна мережа була навчена двом окремим підходам: по-перше, з використанням всіх часових рядів, що охоплюють усі інфекційні захворювання, присутні в наборі даних; і по-друге, зосередження на одному часовому ряді для конкретного інфекційного захворювання.

Для обох конфігурацій навчання модель використовувала історичні дані за попередні 24 місяці для прогнозування випадків на наступні шість місяців. Ці часові рамки мали на меті вловити тенденції та закономірності, які могли б підвищити точність прогнозування.

Для визначення результатів прогнозування була використана середня абсолютна процентна похибка (MAPE), також відома як середнє абсолютне процентне відхилення (MAPD) [4]. Варто відзначити, що більшість інфекційних захворювань не мають деякого періоду коливань, що відображається в передбаченнях нейронної мережі. Transformer Model дала кращі показники передбачень на часових рядах, які коливаються відносно стало (наприклад, кількість захворювань збільшується влітку і зменшується взимку). На рисунку 3 показано передбачення нейронних мереж на часових рядах, які коливаються відносно стало.

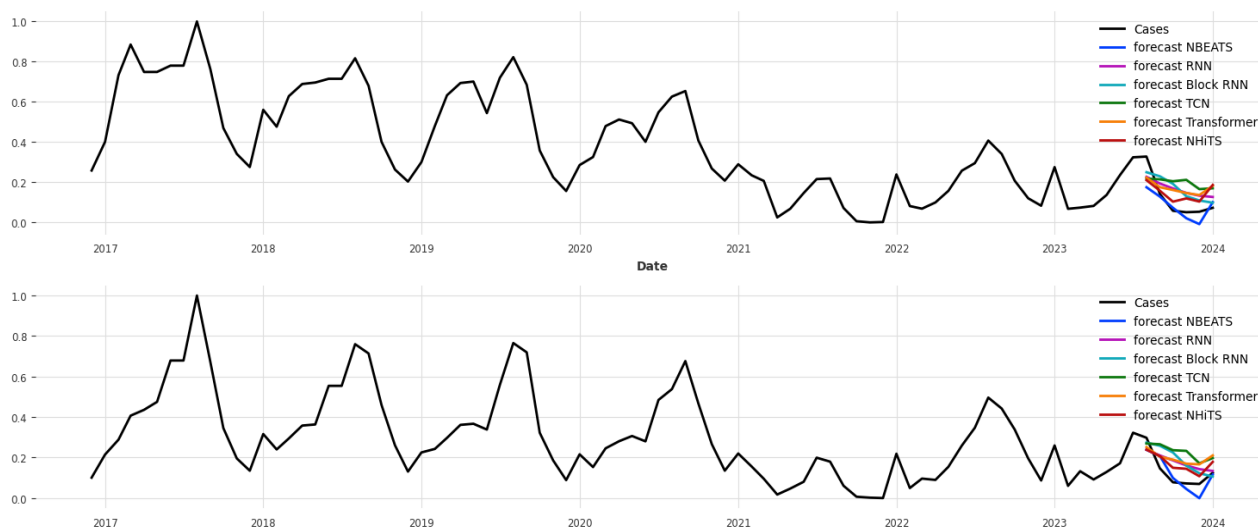


Рис. 3. Приклад прогнозування нейронних мереж.

Для проведення аналізу точності передбачення нейронних мереж, було використано інтервал впевненості в 95%. Таким чином, середні значення MAPE передбачення Transformer, яка навчалась з використанням всіх часових рядів, що охоплюють усі інфекційні захворювання в 95% випадків є 67.99%. В таблиці 1 представлено прогнози архітектури нейронної мережі, яка була навчена на всіх часових рядах.

Таблиця 1.

Прогнози нейронної мережі, яка були навчена на всіх часових рядах

Neural Network	Average MAPE, %
Transformer	67.99

Примітно, що нейронна мережа, навчена на кількох часових рядах, продемонструвала кращу продуктивність порівняно з тією, що було навчено на окремих часових рядах. Середнє значення MAPE для нейронної мережі Transformer Model, яка була навчена на окремих часових рядах склало 122.72%. Це покращення можна пояснити здатністю моделі використовувати спільну інформацію про різні захворювання, що збагатило процес навчання та сприяло більш надійним прогнозам. Результати підкреслюють потенціал підходів із багаточасовими рядами для підвищення точності прогнозування інфекційних захворювань за допомогою нейронних мереж.

Це відкриття також свідчить про те, що інфекційні захворювання можуть впливати одна на одну та на їх поширення, що вказує на складну взаємодію між різними збудниками.

В результаті даної роботи проведено аналіз застосування нейронних мереж для вирішення задачі прогнозування часових рядів, зокрема, прогнозування розповсюдження інфекційних захворювань. Проведено аналіз та нормалізація датасету для тренування. Досліджено архітектуру нейронної мережі Transformer Model. Нейронна мережа з архітектурою Transformer в 95% випадків має 67.99% середню похибку. Відкрито, що інфекційні захворювання можуть впливати одна на одну та на їх поширення.

References

- Goniewicz, K., Burkle, F.M., Horne, S., Borowska-Stefańska, M., Wiśniewski, S., Khorram-Manesh, A. (2021) The Influence of War and Conflict on Infectious Disease: A Rapid Review of Historical Lessons We Have Yet to Learn. Sustainability, 13, 10783. Retrieved from: <https://doi.org/10.3390/su131910783>.
- Public Health Center of the Ministry of Health of Ukraine (2024) Infectious disease in the population of Ukraine, Retrieved from: <https://phc.org.ua/kontrol-zakhvoryuvan/inshi-infekciyni-zakhvoryuvannya/infekciyna-zakhvoryuvannist-naselennya-ukraini>.
- Vaswani A, Shazeer N., Parmar N., Uszkoreit J., Jones L., Gomez A. N., Kaiser L., Polosukhin I., (2017) Attention Is All You Need, *Advances in Neural Information Processing Systems*, 2023, 3-4,. Retrieved from: <https://arxiv.org/abs/1706.03762>
- Myttenaere, Golden, B., Grand, B. Le., Rossi, F. (2016) Mean absolute percentage error for regression models, *Neurocomputing*. Retrieved from: <https://arxiv.org/abs/1605.02541>.

COMPARATIVE ANALYSIS OF THE RESULTS OF THE APPLICATION OF DIFFERENT NEURAL NETWORK ARCHITECTURES FOR THE PREDICTION OF INFECTIOUS DISEASES

This paper presents the results of time series forecasting using machine learning methods and mobile technologies. The main results of this work are a comparison of the effectiveness of different neural network architectures in predicting the spread of infectious diseases in Ukraine for the period from December 2016 to January 2024.

To achieve the goal, the following tasks were solved: the current state of the time series forecasting was analyzed; existing analogs of the systems were analyzed; the necessary neural network architectures were selected as one of the machine learning methods; a dataset for infectious diseases spread was analyzed and normalized; conducted testing. The developed system can be an effective tool for rational management decisions to ensure the epidemiological well-being and biosecurity of the population.

Keywords: Forecasting, time series, machine learning methods, neural networks.

Відомості про автора / Information about the Author

Олександр Ковалів, аспірант кафедри інтелектуальних інформаційних систем в Чорноморському національному університеті імені Петра Могили, м. Миколаїв, Україна. E-mail: kizutoamazed@gmail.com, orcid: <https://orcid.org/0009-0001-2047-6319>.

Oleksandr Kovaliv, PhD student of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University (PMBSNU), Mykolaiv, Ukraine. E-mail: kizutoamazed@gmail.com, orcid: <https://orcid.org/0009-0001-2047-6319>.

Євген Сіденко, канд. техн. наук, доцент кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Ievgen Sidenko, PhD, Associate Professor of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Галина Кондратенко, канд. техн. наук, доцент кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: halyna.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-8446-5096>.

Galyna Kondratenko, PhD, Associate Professor of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: halyna.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-8446-5096>.

DOI 10.34132/mspc2025.01.14.12

Роман Кошовий
Катерина Кірей

ГЕЙМИФІКАЦІЯ НАВЧАЛЬНОГО ПРОЦЕСУ ПІДГОТОВКИ МАЙБУТНІХ ІТ-ФАХІВЦІВ

Тези представляють теоретичний та практичний аналіз актуальності, доцільності й перспектив використання гейміфікації, як інноваційного підходу щодо оптимізації навчального процесу та підвищення його результативності. Зокрема, дослідження ґрунтується на досвіді викладання фахових дисциплін у процесі підготовки здобувачів за спеціальністю 121 «Інженерія програмного забезпечення» в Чорноморському національному університеті ім. Петра Могили.

Теоретичний аналіз проілюстровано прикладом практичної реалізації гейміфікації. Авторами розроблено комп'ютерну гру для відпрацювання практичних навичок та контролю знань із дисципліни «Алгоритми та структури даних». У тезі наведено детальний опис цього ігрового проєкту з акцентом на ключових елементах, спрямованих на покращення якості подання навчального матеріалу та оцінювання навчальних досягнень студентів.

Результати аналізу підтверджують значний потенціал гейміфікації в підвищенні залученості та мотивації здобувачів освіти, а також у покращенні засвоєння складних технічних дисциплін. Запропонований ігровий проєкт може слугувати моделлю для розроблення подібних інструментів для інших дисциплін ІТ-спеціальностей.

Ключові слова: гейміфікація, навчальний процес, інформаційні технології, підвищення результативності, алгоритми та структури даних.

Для ІТ-індустрії дуже важливо постійно розвиватися для швидкого реагування на різноманітні виклики. Це неможливо без висококваліфікованих фахівців із глибоким знанням різних аспектів розробки програмних засобів. Так, зокрема, треба приділити увагу щодо якості підготовки майбутніх ІТ-фахівців. Саме досвідчені фахівці та експерти, які здатні йти в ногу із часом є його рушієм та невід'ємною частиною для розвитку більшої частини всіх інших сфер життя людини. Для появи та становлення таких спеціалістів навчальні заклади грають ключову роль.

У доповіді розглянуто один із перспективних напрямів вдосконалення навчального процесу підготовки майбутніх ІТ-фахівців – гейміфікація навчального процесу.

У рамках студентоцентрованого навчання, що зараз запроваджено у вітчизняних ЗВО, врахування інтересів здобувачів є важливою складовою. Це, зокрема, має знаходити своє відображення у виборі форми подання навчального матеріалу. Для здобувачів, які здебільшого є молодими людьми та «тинейджерами», у сферу інтересів входять фільми, короткі відео, ігри – останній пункт не тільки є цікавою формою пізнання нового, але і способом власне впливати на контент, бути його частиною, діяти та поринути в події. Важливо враховувати й додаткові чинники, які привертають увагу, зацікавлюють та затягують гравців – компонент змагання, можливість розвитку, потреба докладання певних зусиль, які потім супроводжуються кращим результатом або перемогою, стимулюючи прихід дофаміну до мозку й додаючи ще більше стимулу продовжувати. Використовуючи це як інструмент, можна надати можливість здобувачам засвоювати важливі для них знання та навіть використовувати їх на практиці через тип контенту, який їм цікавий. Результат же проходження гри або її конкретних рівнів може слугувати валідним показником рівня знань здобувача в тій чи іншій темі або навіть дисципліні – можливості для оцінювання гравців у такий спосіб, а також різні шляхи стимулу до покращення результатів нині є одним із рішень, яке можна використовувати та поширювати між навчальними закладами.

В Україні вже є досвід гейміфікації освітнього процесу в молодших класах – і він досить широко відомий, навіть поточним здобувачам вищої освіти, на прикладі проєкту «Сходінки до інформатики» [1], який покриває обширну сферу в алгоритмуванні початкового рівня, такого актуального для школярів до 7-го класу. Окремо

треба зазначити наявність низки головоломок та ігрових проєктів для покращення навичок програмування [2; 3] – виконані в різних жанрах та використовуючи різні підходи, такі проєкти надають гравцям можливість розвивати навички алгоритмізації або створення невеликих програмних рішень задля досягнення мети, паралельно покращуючи навички вирішення задач, використання певних мов програмування та інструментів для розробки, автоматизації та оптимізації процесів. Особливо оптимізація є ключовим питанням у процесі розроблення, оскільки іноді невмотивовані спеціалісти можуть вважати справу виконаною, якщо рішення працює, але саме удосконалення продуктивності через оновлення програмних рішень на різних рівнях створює покращені версії наявних або абсолютно нові продукти.

Впровадження гейміфікації в освітній процес ЗВО дає змогу зробити процес засвоєння навчального матеріалу більш простим та цікавим для здобувачів. Це також є своєрідним викликом для освітньої системи та проєктування навчально-ігрових інформаційних систем. У разі оновлення навчальних програм та устаткування, деякі елементи системи гейміфікованого викладання можуть застарівати або потребувати змін, але вдало спроектована система може бути гнучкою та легко піддаватися змінам, потребуючи невеликих зусиль та вкладень із погляду розроблення.

Гейміфікації піддається більшість дисциплін, які, зазвичай, викладаються в університеті – дисципліни з програмування (для перших курсів уже виглядає як гра, тому для фіналізації достатньо візуалізувати), дисципліни з математики, тестування ПЗ тощо. Надалі в доповіді розглянуто підходи та особливості гейміфікації навчального процесу на прикладі дисципліни «Алгоритми та структури даних». Розуміння одиниць даних, їхнє взаємовідношення на фізичному рівні та їхня обробка програмним кодом, складності та оптимізації алгоритмів стає набагато простіше для здобувача тоді, коли це можна «помацати», подивитися зблизька. Для цієї дисципліни, як і для багатьох інших, мова програмування не є чітко визначеною, тому найкраще використовувати абстрактний підхід до завдань/рівнів гри, результуючи в універсальному навчальному рішенні.

Варіантів реалізації та надання доступу до сервісу багато – це може бути централізована клієнт-серверна система з управлінням в одному або декількох серверних центрах у межах країни (але не обмежуючись однією країною), або така, яка впроваджується окремо для кожного університету (аналогічно Moodle), будучи децентралізованою в аспекті використання, але підпорядкована оновленням ПЗ від одного джерела, контрольованим державними або іншими органами регулювання освіти.

Для апробації запропонованого рішення реалізовано демонстративну версію ігрового застосунку із використанням React – продуктивне та легковагове кінцеве ПЗ із невеликими зусиллями розгортається в локальній мережі, даючи змогу навчальному закладу обмежити область використання та активно впроваджувати нові зміни. В основу розробленої нами гри покладено імітацію сортування масиву з погляду апаратного забезпечення – гравцю пропонується взяти на себе роль обчислювального комплексу та самостійно виконувати кожен крок для досягнення однієї мети – відсортованого масиву.

Гра має багато унікальних особливостей, виконаних за аналогією з фактичною реалізацією роботи з даними:

- гравець не знає значень у комірках масиву, але може бачити початкові індекси елементів;
- початковий масив випадкового розміру та містить елементи у випадковому порядку;
- гравець може порівнювати значення двох комірок між собою, використовуючи доступні в більшості мов програмування операції порівняння;
- гравець може копіювати значення з однієї комірки в іншу, що підтверджується перенесенням індексу початкового елемента для наочності;
- копіювання та порівняння коштує очок, які додаються до рахунку гравця, імітуючи використані ресурси (чим менше – тим краще);
- для використання окрім початкового масиву доступна ще одна комірка пам'яті, яка є абсолютним мінімумом для можливості переставити 2 значення (swap), а також доступний окремий масив (ідентичний початковому, але на початку рівня заповнений порожніми значеннями) – перше використання цього масиву також додає очки до рахунку, де їхня кількість залежить від розміру масиву, що імітує використання пам'яті;
- для зручності гравець також може заблокувати комірку подвійним натисканням для запобігання випадковим операціям («miss-click») та виконати маркування комірок, наприклад, для більш наочного планування стратегії досягнення цілі;
- гравець сам має зазначити, коли масив є відсортованим і рівень можна вважати виконаним, у такому разі рівень буде вважатися пройденим, навіть якщо очікуваний відсортований масив міститься в додатковій пам'яті;
- помилкові операції гравця, такі як хибне зазначення відсортованого масиву або випадкова спроба знищення останнього елемента певного індексу, накладають штрафні бали на рахунок, змушуючи не поспішати та краще обдумувати кожен хід.

Усе це разом надає здобувачу можливість не просто виконати сортування на нижчому рівні, полегшуючи розуміння роботи з даними, але і простір для фантазії. Гра не обмежує користувача в можливих методах вирішення, а також не має «програшу» – що би не робив гравець, рівень можна пройти, адже будь-який вхідний масив випадково перемішаний і може бути відсортованим із будь-якого свого стану (рис. 1). Тим не менш, визначена стратегія, будучи ідентичною наявному алгоритму сортування або унікальним методам гравця, вирішуватиме кінцевий результат – рахунок, який можна також вважати абстрактною кількістю витрачених ресурсів на операцію, до яких входить як кількість операцій, так і використана пам'ять (використання додаткового масиву не є обов'язковим).

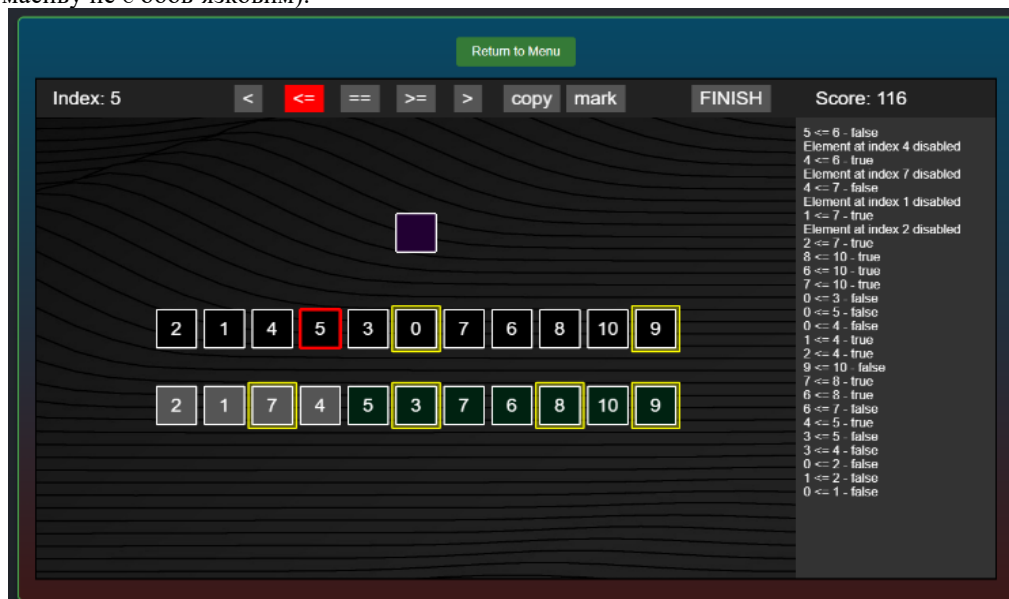


Рисунок 1 – Демонстрація ігрового процесу

За результатами сортування та залежно від довжини масиву вираховується кінцевий рейтинг, додаючи елемент змагання до гри – здобувачам буде цікаво отримати найкращий результат та потрапити в таблицю лідерів (рис. 2). Оптимізація гравцями свого алгоритму сортування та безперечне слідування ньому покращуватиме їхній результат, водночас даючи здобувачам простір для набуття навичок роботи з даними.

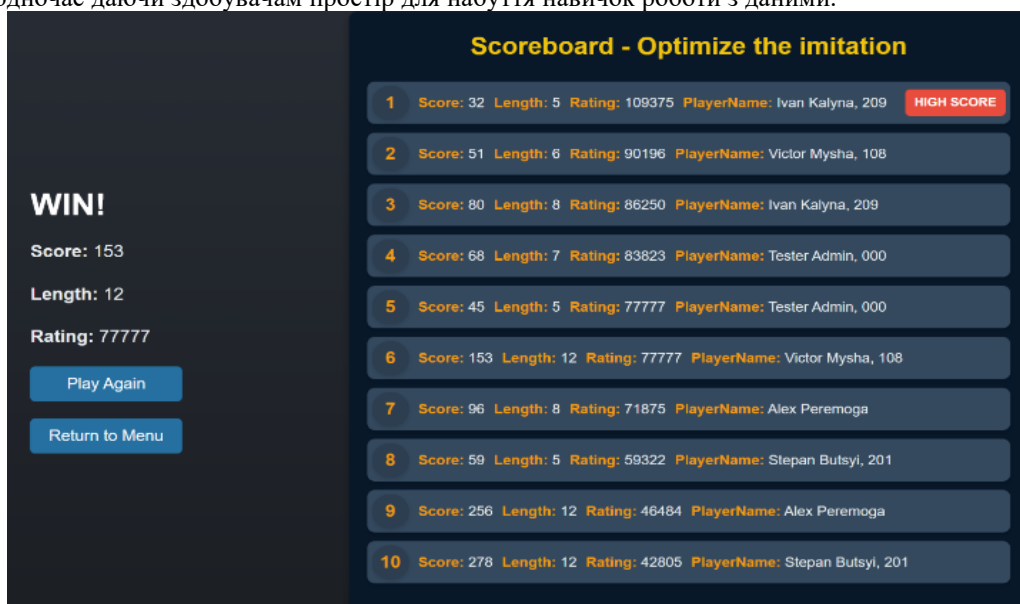


Рисунок 2 – Таблиця лідерів для ігрового рівня

Грунтуючись на запропонованому підході можна спроектувати та реалізувати масштабні та гнучкі системи гейміфікованого навчання, які можна використовувати як здобувачам для закріплення знань та напрацюванню навичок за темами, так і викладачам для оцінювання знань, наприклад, через демонстрацію здобувачем проходження рівня наживо.

Як висновок, проведене дослідження переконливо демонструє значний потенціал гейміфікації як ефективного інструменту для якісної трансформації навчального процесу в сфері інформаційних технологій. Впровадження ігрових механік сприяє не лише підвищенню рівня залученості та внутрішньої мотивації здобувачів освіти, але й забезпечує більш глибоке та усвідомлене засвоєння складного теоретичного матеріалу. Практична реалізація гейміфікованого підходу на прикладі навчальної гри з дисципліни «Алгоритми та структури даних» підтверджує можливість створення інтерактивного та цікавого освітнього середовища, що позитивно впливає на результативність навчання та якість підготовки майбутніх ІТ-фахівців. Запропонований ігровий проєкт може слугувати цінним прецедентом та методологічною основою для розробки аналогічних інструментів в інших технічних дисциплінах або більш глобального освітнього проєкту для ЗВО, сприяючи подальшому розвитку інноваційних педагогічних практик у вищій освіті.

References

1. Official project website. Steps to Computer Science. URL: <http://dvsvit.com.ua/cxodunku/> (access date: 04.26.2025).
2. CodeCombat - Coding games to learn Python and JavaScript. CodeCombat. URL: <https://codecombat.com/> (access date: 04.26.2025).
3. Tomorrow Corporation. 7 Billion Humans. Tomorrow Corporation, 2018. URL: https://store.steampowered.com/app/792100/7_Billion_Humans/ (access date: 04.26.2025).

DOI 10.34132/mspc2025.01.14.12

Roman Koshovyi,
Kateryna Kirei,

GAMIFICATION OF THE EDUCATIONAL PROCESS OF FUTURE IT SPECIALISTS

The theses present a theoretical and practical analysis of the relevance, feasibility, and prospects of using gamification as an innovative approach to optimizing the educational process and enhancing its effectiveness. The study is particularly based on the experience of teaching professional disciplines in the training of students majoring in 121 "Software Engineering" at Petro Mohyla Black Sea National University.

The theoretical analysis is illustrated by a practical example of gamification implementation. The authors developed a computer game for practicing practical skills and assessing knowledge in the discipline "Algorithms and Data Structures." The thesis provides a detailed description of this game project, emphasizing the key elements aimed at improving the quality of educational content delivery and evaluating students' academic achievements.

The analysis results confirm the significant potential of gamification in increasing student engagement and motivation, as well as improving the assimilation of complex technical subjects. The proposed game project can serve as a model for the development of similar tools for other IT disciplines.

Keywords: gamification, educational process, information technology, improving performance, algorithms and data structures.

Відомості про автора / Information about the Author

Роман Кошовий, здобувач 1-го курсу магістратури інженерії програмного забезпечення Чорноморського Національного Університету імені Петра Могили, м. Миколаїв, Україна. E-mail: romanrockel@gmail.com, orcid: <https://orcid.org/0009-0005-2951-6574>

Roman Koshovyi, 1st year Master's student in Software Engineering at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: romanrockel@gmail.com, orcid: <https://orcid.org/0009-0005-2951-6574>

Катерина Кірей, канд. пед. наук, доцент, доцент кафедри інженерії програмного забезпечення Чорноморського Національного Університету ім. Петра Могили, м. Миколаїв, Україна. E-mail: kateryna.kirey@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-9338-2380>

Kateryna Kirei, Candidate of Pedagogical Sciences, Associate Professor, Associate Professor of the Department of Software Engineering at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: kateryna.kirey@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-9338-2380>

DOI 10.34132/mspc2025.01.14.13

Олександр Обушко
Інесса Кулаковська

СИСТЕМА РОЗПІЗНАВАННЯ І КЛАСИФІКАЦІЇ ФІНАНСОВИХ ДОКУМЕНТІВ

Анотація. Дослідження присвячене автоматизації обробки документів, зокрема товарно-транспортних накладних (ТТН), для цифровізації підприємств. Основною метою є створення гнучкої програмної системи, яка автоматично витягує реквізити з документів різних форматів із високою точністю. Система використовує класичні методи обробки зображень, оптичне розпізнавання символів (OCR) через бібліотеку Nicomsoft OCR, а також шаблонний підхід до семантичного аналізу тексту. Дані проходять нормалізацію, забезпечуючи стандартизоване зберігання інформації. Реалізація виконана мовою Java, із використанням Swing для UI, Apache POI для роботи з документами та Firebird SQL для збереження даних. Завдяки модульній архітектурі система легко інтегрується в корпоративні інформаційні платформи, підвищує ефективність обробки документів та може бути масштабована для інших типів документації, що сприяє цифровій трансформації підприємств.

Ключові слова: автоматизація обробки документів, оптичне розпізнавання символів (OCR), семантичний аналіз тексту, цифрова трансформація підприємств, модульна архітектура системи.

Цифрова трансформація підприємств значно змінює методи ведення документації, особливо в сферах торгівлі та логістики. Незважаючи на активне впровадження електронного документообігу, значна частина процесів все ще пов'язана із паперовими накладними, що вимагає ручного внесення даних до комп'ютерних систем. Це призводить до значних витрат часу, підвищеного ризику помилок та зростання витрат на адміністрування.

Автоматизація обробки товарно-транспортних накладних (ТТН) потребує інтеграції сучасних методів комп'ютерного бачення та оптичного розпізнавання символів (OCR). Враховуючи різноманітність форматів ТТН, ключовим завданням є розробка адаптивної системи, здатної коригувати викривлення та нерівномірність освітлення під час сканування або фотографування документів. Застосування методів попередньої обробки, таких як коригування контрасту, видалення артефактів та покращення роздільної здатності, дозволяє підвищити точність подальшого текстового аналізу.

Традиційні OCR-системи ефективні у розпізнаванні стандартних шрифтів і чітких документів, проте в умовах змінних форматів ТТН можуть виникати труднощі з інтерпретацією нестандартних символів, розпізнаванням числових полів або коректним відокремленням розділових знаків. Для усунення таких проблем доцільним є застосування додаткових алгоритмів семантичного аналізу тексту та валідації числових реквізитів.

Запропонована система інтегрує технологію OCR з шаблонним підходом до обробки даних, що дозволяє швидко адаптувати розпізнавання під різні формати документів. Реалізація на базі Java та Firebird SQL забезпечує високу продуктивність, гнучкість інтеграції з корпоративними інформаційними системами та масштабованість рішення. Впровадження автоматизованої обробки ТТН сприяє оптимізації процесів управління поставаннями, скороченню витрат та підвищенню точності бухгалтерських розрахунків, що є важливим аспектом цифрової трансформації підприємств.

Автоматизація обробки документів є важливим завданням для цифровізації сучасних підприємств. Особливої актуальності це набуває в тих галузях, в яких присутній великий обсяг паперової документації, зокрема це галузі логістик, торгівлі і виробництва. Ручне введення даних із накладних займає багато часу, супроводжується помилками та вимагає залучення додаткового персоналу. Основною метою цього стартапу є створення гнучкої програмної системи, здатної автоматично витягувати реквізити з товарно-транспортних накладних (ТТН) різних форматів з високою точністю і зберігати результати у структурованому вигляді для подальшої обробки.

Основною частиною проекту є класичні етапи обробки вхідного зображення: вирівнювання, покращення контрастності, фільтрація шумів, видалення фону та зайвих артефактів. Данні операції покращують якість зображення для підвищення точності подальшого розпізнавання тексту. Система повинна бути орієнтованою на обробку зображень знятих на мобільні пристрої. Враховуючи велику різноманітність форматів надходження

документів, система повинна бути здатна самостійно адаптувати зображення до прийнятного формату для аналізу. Саме тому було реалізовано механізми автоматичного визначення орієнтації документа, а також коригування спотворень зображень, знятих під кутом, які часто виникають при ручній зйомці.

Після етапу попередньої обробки зображення застосовується технологія оптичного розпізнавання символів (OCR). Обраною бібліотекою для даної задачі стала Nicomsoft OCR, яка забезпечує стабільне розпізнавання кирилиці, зокрема української мови. Важливою перевагою є можливість локального використання, що дозволяє забезпечити збереження конфіденційності оброблених даних без потреби у зовнішніх сервісах. Бібліотека має зручний інтерфейс для інтеграції в Java-застосунки та підтримує налаштування зон розпізнавання, що дозволяє визначати, які саме ділянки документа потрібно аналізувати. Інструменти Nicomsoft OCR дозволяють точно контролювати процес розпізнавання: задавати мови, формати символів, зональні обмеження, що критично важливо при роботі з багатосторінковими та неструктурованими документами. Висока гнучкість API цієї бібліотеки дозволяє зменшити час розробки та підвищити адаптивність до нестандартних форматів документів. На Рис.1. показано концептуальну архітектуру програми Nicomsoft, яка складається з трьох основних компонентів: тексту, зображення та OCR.

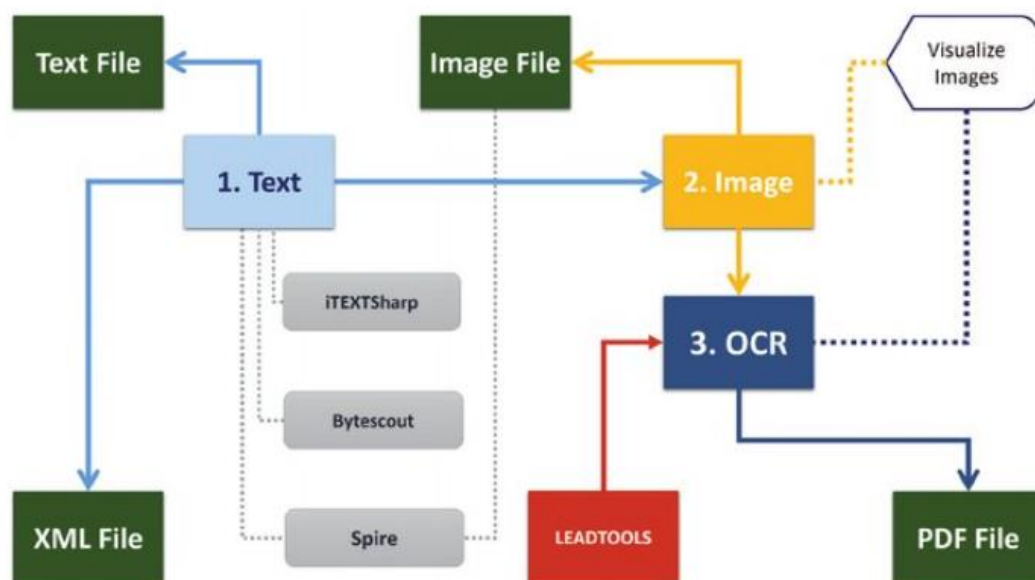


Рис. 1. Концептуальна архітектура OCR

Наступним етапом є семантичний аналіз розпізнаного тексту. Для реалізації цієї частини було запроваджено шаблонний підхід: для кожного нового постачальника створюється шаблон із вказівкою координат ключових полів документа. Завдяки такому рішенню при повторному надходженні схожого документа система автоматично застосовує відповідний шаблон і витягує необхідні дані. Це дозволяє значно підвищити швидкість і точність обробки. У майбутньому така система може бути розширена у бік машинного навчання, тобто дозволити їй самостійно створювати шаблони на основі попереднього досвіду без участі користувача. Враховуючи обмеження поточного етапу реалізації, шаблонний метод виступає перехідним етапом між повністю ручною і повністю автоматизованою системою.

Дані, які отримали внаслідок розпізнавання, проходять додаткову нормалізацію. Тут мається на увазі приведення числових значень, дат, одиниць виміру до єдиного формату. Наприклад, ціна «1 000,00 грн» та «1000.00 UAH» повинні зберігатися у базі даних у єдиному форматі. Також проводиться обробка розділових знаків, пробілів, символів валюти, форматів дати. У разі помилкового розпізнавання форматів система надає оператору можливість перевірити відскановані дані та виправити їх вручну. Але для подальшого вдосконалення система повинна мати змогу попереджати користувача, або автоматично виправляти типові неточності спираючись на вбудовані словники та шаблони даних. Стандартизація запису ключових реквізитів відіграє вирішальну роль при подальшому аналітичному аналізі, звітуванні та контролі бухгалтерського обліку. Також у

нормалізації використовуються правила перетворення реєстрів тексту, стандартизація позначень одиниць виміру (наприклад, шт. → штуки) та валют (грн, UAH), що дозволяє уникати неоднозначності в подальшому аналізі.

1 прим. - вантажовідправнику
 2 прим. - вантажоодержувачу
 3 і 4 прим. - перевізнику

ТОВАРНО-ТРАНСПОРТНА НАКЛАДНА

Форм. № 1-ТН затверджено наказом Мінінфраструктури
 України 05.12.2013 р. №983

№ 241901 від 20 р.

Автомобіль _____ / Прізвище/ім'я/прізвище _____ Вид перевезення _____
(марка, модель, рік, реєстраційний номер) (марка, модель, рік, реєстраційний номер)

Автомобільний перевізник _____ Водій _____
(найменування П.І.Б.) (П.І.Б., номер свідоцтва водія)

Замовник _____
(найменування П.І.Б.)

Вантажовідправник _____
(повне найменування, місцеве/загальнодержавне П.І.Б., місце проживання)

Вантажоодержувач _____
(повне найменування, місцеве/загальнодержавне П.І.Б., місце проживання)

Пункт навантаження _____ Пункт розвантаження _____
(найменування) (найменування)

Передаресування вантажу _____
(найменування, місцеве/загальнодержавне П.І.Б., місце проживання вантажоодержувача, П.І.Б., посвідчення водія)

Відпуск за довіреністю вантажоодержувача: серія _____ № _____ від _____ р., наділено _____

Вантаж наданий для перевезення у стані, що _____
(найменування вантажу) , отримав водій/експедитор _____
(П.І.Б., посвідчення водія)

кількість місць _____, масою бруто, т _____
(обчислюється за формулою)

Бухгалтер (відповідальна особа вантажовідправника) _____ Відпуск дозволив _____
(П.І.Б., посада, підпис, печатка) (П.І.Б., посада, підпис, печатка)

Усього відпущено на загальну суму _____, у т.ч. ПДВ _____
(сума, у зразковому ПДВ)

Супровідні документи на вантаж _____
 Транспортні послуги, які надаються автомобільним перевізником: _____

ВІДОМОСТІ ПРО ВАНТАЖ

№ з/п	Найменування вантажу (номер контейнера), у разі перевезення небезпечних вантажів, клас небезпечного речовини, де жодна з них не є вантажем	Одиниця виміру	Кількість місць	Маса без ПДВ за відпуском, грн	Загальна сума з ПДВ, грн	Вид транспортування	Документи і вантажівка	Маса бруто, т
1	2	3	4	5	6	7	8	9
Всього:								

Зак (відповідає особа вантажовідправника) _____
(П.І.Б., посада, підпис, печатка)

Прийняв водій/експедитор _____
(П.І.Б., посада, підпис)

Зак (відповідає особа) _____
(П.І.Б., посада, підпис, печатка)

Прийняв (відповідає особа вантажоодержувача) _____
(П.І.Б., посада, підпис, печатка)

ВАНТАЖНО-РОВАНТАЖУВАЛЬНІ ОПЕРАЦІЇ

Операція	Маса бруто, т	Час (год., хв.)		Підпис відповідальної особи	
		прибуття	вибуття		
10	11	12	13	14	15
Навантаження					
Розвантаження					

Рис.2. Приклад ТТН

Програмна реалізація системи виконана мовою Java. Такий вибір пояснюється її платформеною незалежністю, підтримкою багатопотоковості, високою продуктивністю та великою кількістю бібліотек. Бібліотека Swing використовується для створення графічного інтерфейсу користувача. Вона дозволяє будувати зручні вікна з формами, кнопками, полями введення та іншим функціоналом. Для обробки файлів є бібліотека Apache Commons IO, вона спрощує роботу з каталогами, потоками та файлами. Apache POI використовується для формування таблиць та звітів, вони можуть бути експортовані у формат Excel для подальшого аналізу. Log4j забезпечує системне логування, що необхідне для відстеження помилок і подій у програмі. Завдяки такому пакету бібліотек, кожен етап обробки, від отримання документа до виводу обробленої інформації, реалізується прозоро та стабільно.

Оброблені та нормалізовані дані зберігаються у базі даних Firebird SQL. Ця СУБД з відкритим кодом, з низькими системними вимогами, швидкою роботою на малих об'ємах і простотою адміністрування. Вона не потребує окремого сервера, що робить її придатною для вбудованих або локальних рішень. База підтримує тригери, збережені процедури, механізми забезпечення цілісності даних. Також дозволяє реалізовувати складні логічні залежності між таблицями за допомогою SQL-запитів. Такий підхід дозволяє спростити логіку роботи з даними, покращити контроль за змінами у БД та забезпечити надійність довготривалого зберігання інформації. Система підтримує транзакції – будь-яка група операцій виконується як єдине ціле, що особливо важливо при збереженні змін з декількох джерел одночасно або під час аварійного завершення роботи програми.

Проаналізувавши вимоги та функціонал майбутньої системи було прийнято рішення впровадити роботу системи у двох режимах: інтерфейсний - використовується для навчання та створення шаблонів бланків, та

автоматичний - створений для автоматичного розпізнавання реквізитів у документах на основі існуючих шаблонів бланків для них. Створений шаблон зберігається у базі даних. Розпізнавання інших документів, його типу бланку вже буде доступно у автоматичному режимі. Щодо прецеденту «перенавчання» бланку - воно проходить абсолютно так само як і навчання. Програма сама оновить дані щодо шаблону у базі даних. Таку процедуру навчання потрібно провести для усіх типів бланків, які будуть проходити розпізнавання у автоматичному режимі.

Алгоритм роботи у інтерфейсному режимі представлений у Рис. 3.

Алгоритм роботи у автоматичному режимі представлений на Рис. 4.



Рис.3. Алгоритм роботи у інтерфейсному режимі

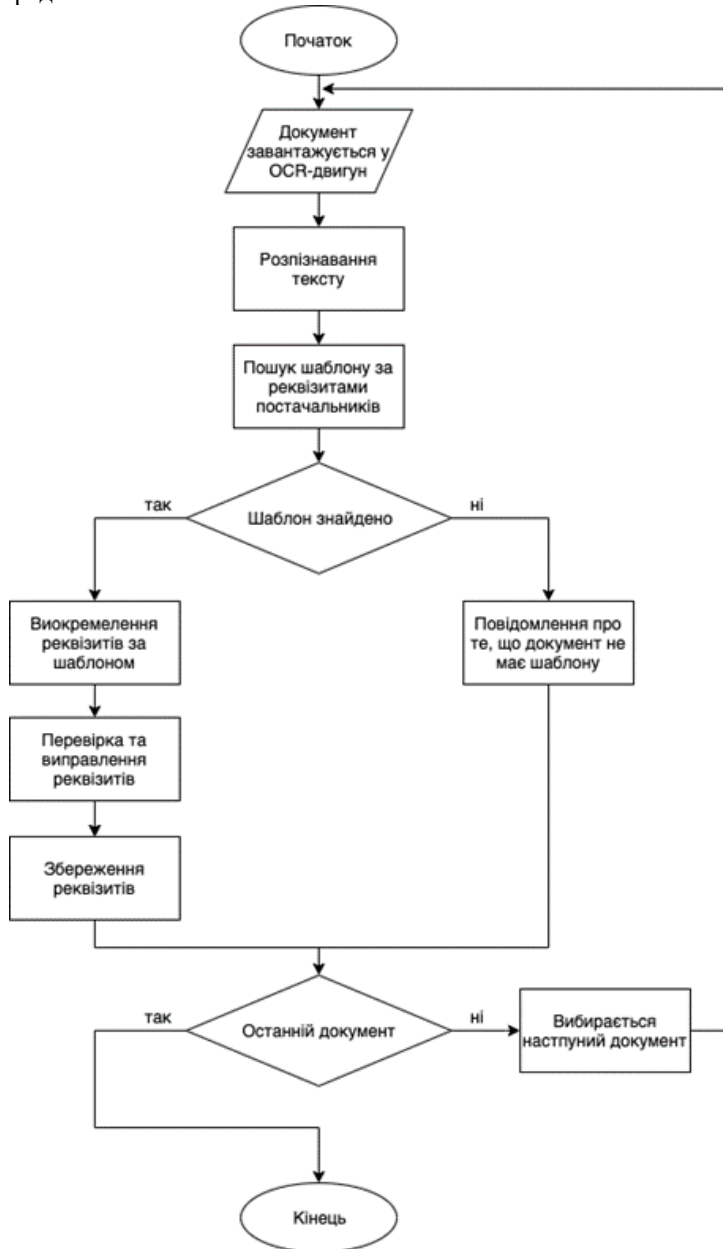


Рис.4. Алгоритм роботи у автоматичному режимі

Під час роботи у автоматичному режимі, на початку програма сама сканує базу даних на предмет нерозпізнаних документів. Або відразу завантажує потрібний документ, якщо є параметр `doc_pi_import_id`. Після цього починає їх розпізнавати. Після розпізнавання тексту, потрібно знайти шаблон документу. Для цього у базі

даних є дані усіх постачальників. У розпізаному документі шукають атрибути одного із збережених постачальників. Такими атрибутами є електронна пошта, номер телефону та назва. Після того як шаблон знайдений, за його допомогою виокремлюються потрібні реквізити. Перед збереженням, реквізити перевіряються. Наприклад, що реквізит суми не містить в собі літери. А всі дати переводяться до одного формату. Всі перевірені реквізити програма відразу зберігає у базі даних. При помилках чи не вдалому розпізнаванні, наприклад не розізнався певний реквізит чи не було знайдено шаблону для документа, програма позначає це у Log-файлі.

У роботі в автоматичному режимі відсутній інтерфейс, програма лише записує результати у Log-файл та відображає їх у консолі. Саме можливість запускати програму та задавати їй усі параметри із консолі дозволяють їй бути незалежною та легко-вбудованою у інші системи.

Таким чином, система, побудована на основі технологій Java, Nicomsoft OCR, Swing, Firebird SQL, Apache Commons IO та Apache POI, створює гнучке та ефективне рішення для автоматизації обробки ТТН. Вона легко адаптується до різних форматів документів, забезпечує точність розпізнавання і високу продуктивність роботи, а її модульна структура відкриває можливості для подальшого розвитку та масштабування. Перевагою цієї архітектури є те, що вона дозволяє створити не просто застосунок, а платформу, яка може бути адаптована для інших типів документів, наприклад, актів виконаних робіт, рахунків-фактур, сертифікатів та чеків. Завдяки цьому система набуває універсальності і може бути впроваджена у підприємствах з різними формами документообігу. Крім цього, програмний код, побудований на відкритих технологіях, дозволяє зменшити вартість впровадження та спрощує модифікацію у разі потреби.

Висновок. Автоматизація обробки документів є важливим етапом цифровізації підприємств, спрямованим на підвищення точності обробки даних та зниження витрат часу на ручне введення інформації. Запропонована система базується на сучасних методах комп'ютерного бачення та оптичного розпізнавання символів (OCR) з використанням бібліотеки Nicomsoft OCR, що забезпечує ефективне витягування реквізитів з товарно-транспортних накладних. Семантичний аналіз тексту із застосуванням шаблонного підходу сприяє адаптації алгоритмів до різних форматів документів. Використання мов програмування Java та бази даних Firebird SQL забезпечує високу продуктивність, надійність зберігання даних та гнучкість інтеграції з корпоративними інформаційними системами. Модульна архітектура дозволяє розширювати функціонал системи, інтегрувати її з зовнішніми сервісами та застосовувати для обробки інших типів документів. Ці технологічні рішення сприяють підвищенню ефективності управління документацією та впровадженню сучасних підходів до цифрової трансформації підприємств.

References

1. Shahar Markovitch and Paul Willmott/McKinsey – Accelerating the digitalization of business processes[Електронний ресурс] - Режим доступу: <https://www.mckinsey.com/business-functions/mckinsey-digital/our-insights/accelerating-the-digitization-of-business-processes>
2. Sayallar S., Sayar A., Babalik N.: An OCR Engine for Printed Receipt Images using Deep Learning Techniques. International Journal of Advanced Computer Science and Applications, Vol. 14, No. 2, 2023. URL: <https://avesis.kocaeli.edu.tr/yayin/aa68b0f4-568a-46da-b680-98d60f7e561c/an-ocr-engine-for-printed-receipt-images-using-deep-learning-techniques> (дата звернення: 22.04.2025)
3. Yue A. Automated Receipt Image Identification, Cropping, and Parsing, 2018. URL: https://www.cs.princeton.edu/courses/archive/spring18/cos598B/public/projects/System/COS598B_spr2018_ReceiptParsing.pdf (дата звернення: 23.04.2025)
4. Drobac S., Linden K.: Optical character recognition with neural networks and post-correction with finite state methods. International Journal on Document Analysis and Recognition (IJDAR) (2020) 23:279–295 URL: <https://link.springer.com/article/10.1007/s10032-020-00359-9> (дата звернення: 23.04.2025)

DOI 10.34132/mspc2025.01.14.13

Oleksandr Obushko
Inessa Kulakovska

SYSTEM FOR RECOGNIZING AND CLASSIFYING FINANCIAL DOCUMENTS

Annotation: This study focuses on the automation of document processing, specifically the extraction of requisites from waybills (TTN) to facilitate digital transformation in enterprises. The main objective is to develop a flexible software system capable of accurately extracting data from documents of various formats. The system incorporates image preprocessing techniques, optical character recognition (OCR) using Nicomsoft OCR, and a template-based semantic

text analysis approach. Extracted data undergoes normalization to ensure standardized storage. The implementation is based on Java, utilizing Swing for the user interface, Apache POI for document processing, and Firebird SQL for data storage. The modular architecture enables seamless integration with corporate platforms and scalability for processing different document types. By enhancing document management efficiency, this solution supports enterprise digitalization, optimizing data processing speed and accuracy while reducing operational costs. The adaptability of the system makes it a valuable tool for modernizing administrative workflows across various industries.

Key words: document processing automation, optical character recognition (OCR), semantic text analysis, digital transformation of enterprises, modular system architecture.

Відомості про автора / Information about the Author

Олександр Обушко, здобувач ступеню бакалавра, факультет комп'ютерних наук, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна, 54003. E-mail: vebvongolla@gmail.com, orcid: <https://orcid.org/0009-0000-7421-5993>

Oleksandr Obushko, Bachelor's Degree Candidate, Department of Intelligent Information Systems, Faculty of Computer Science, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine, 54003, E-mail: vebvongolla@gmail.com, ORCID: <https://orcid.org/0009-0000-7421-5993>

Інесса Кулаковська, кандидат фізико-математичних наук, доцент кафедри Інтелектуальних інформаційних систем, факультет комп'ютерних наук, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: inessa.kulakovska@chmnu.edu.ua, orcid: <http://orcid.org/0000-0002-8432-1850>

Inessa Kulakovska, PhD of Physical and Mathematical Sciences, Associate Professor, Department of Intelligent Information Systems, Faculty of Computer Science, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine, 54003, E-mail: Inessa.Kulakovska@chmnu.edu.ua, ORCID: <http://orcid.org/0000-0002-8432-1850>

DOI 10.34132/mspc2025.01.14.14

Віктор Раленко,
Євгеній Стоєв

ЗАСТОСУВАННЯ FIREBASE STUDIO В РОЗРОБЦІ REACT ЗАСТОСУНКІВ

Тези присвячені *Firebase Studio*, хмарній платформі Google, представлений 9 квітня 2025 року, для створення веб- та мобільних застосунків, зокрема на основі *React*. Платформа вирішує ключові проблеми розробки: відсутність віддаленого робочого місця, недостатню потужність обладнання, платну інтеграцію ШІ-сервісів та обмежену кількість тестових пристроїв. *Firebase Studio* підтримує імпорт проєктів із систем контролю версій, пропонує шаблони для популярних фреймворків (*React*, *Next.js*, *Angular*, *Flutter*), швидке прототипування через *Gemini* з мультимодальними підказками, інтеграцію з *Firebase* та *Google Cloud*, а також настроюване середовище на базі *Code OSS* із підтримкою *Nix*.

Тези демонструють створення *React*-застосунку з *API People Data Generator*, підкреслюючи обмеження у підтримці сторонніх бібліотек, як-от *Ant Design*, через експериментальний статус платформи. *Firebase Studio* має значний потенціал для розширення функціоналу, спрощуючи розробку та тестування застосунків у хмарі, що сприяє ефективності та масштабованості проєктів.

Ключові слова: *Firebase Studio*, *React*, хмарна розробка, *Google Cloud*, *Gemini*, прототипування, ШІ, *Code OSS*, *API*, *Ant Design*.

Firebase Studio – новий продукт, представлений компанією Google 9 квітня 2025 року [1] для використання *Google Platform* в якості середовища розробки різних видів застосунків, таких як веб-застосунки (як фронт-енд так і фулстек) та мобільні застосунки (кросплатформові та нативні)

Використання *Firebase Studio* дозволяє вирішувати такі проблеми як:

- відсутність віддаленого робочого місця для роботи з проєктами з будь-якої точки світу
- недостатню потужність персонального комп'ютера розробника для роботи з ресурсоемними технологіями (такими як нативна розробка для мобільних пристроїв з використанням симуляції (або емуляції в випадку *Android* пристрою

- платну інтеграцію сервісів ШІ розробки програмного продукту;

- відсутність достатньої кількості пристроїв тестування програмного забезпечення.

Проблема відсутності віддаленого робочого місця для роботи з проєктами з будь-якої точки світу залишається актуальною, оскільки не всі можуть дозволити собі оренду віртуального серверу. Проблема недостатньої потужності теж є актуальною, так як з часом технології розробки, тестування та емуляції роботи програмного забезпечення модернізуються, отримують нові функціональні особливості, в зв'язку з чим вони споживають більше ресурсів як з серверного, як з розробницького так і з клієнтського пристроїв.

Проблема платної інтеграції зумовлена ціновою політикою компаній- власників систем ШІ, таких як *OpenAI*. Всі ці та багато інших проблем вирішує *Firebase Studio*.

Firebase Studio – це хмарна платформа, яка надає наступні функціональні можливості:

- імпорт проєктів з системи керування версіями або локального архіву: перенесіть власні програми до *Firebase Studio*, імпортувавши локальний архів або підключивши публічний чи приватний репозиторій на *GitHub*, *GitLab* або *Bitbucket*.

- швидке налаштування проєкту за допомогою вбудованих шаблонів та зразків: *Firebase Studio* надає розширену підтримку фреймворків та мов з великою бібліотекою шаблонів та зразків програм, включаючи популярні мови, такі як *Go*, *Java*, *.NET*, *Node.js* та *Python Flask*, а також фреймворки, такі як *Next.js*, *React*, *Angular*, *Vue.js*, *Android*, *Flutter* та інші. Почніть із шаблону або зразка програми з галереї шаблонів та/або створіть власний

шаблон для спільного використання.

– швидке прототипування природною мовою: Використовуйте Gemini у Firebase для прототипування та публікації повноцінних веб-застосунків за допомогою агента App Prototyping. Створюйте цілі застосунки з мультимодальними підказками, включаючи природну мову, зображення та малюнки.

– завжди доступна допомога зі штучним інтелектом від Gemini у Firebase: використовуйте допомогу з кодування штучним інтелектом від Gemini у Firebase на всіх поверхнях розробки: інтерактивний чат, генерація коду, запуск інструментів та пропозиції вбудованого коду. Gemini у Firebase може допомогти вам писати код і документацію, виправляти помилки, писати та запускати модульні тести, керувати залежностями та вирішувати їх, працювати з контейнерами Docker тощо.

– знайоме та високо настроюване середовище розробки: Firebase Studio побудовано на популярному проєкті Code OSS та працює на повноцінній віртуальній машині (VM) на базі Google Cloud. Ви можете налаштувати майже кожен аспект вашого онлайн-середовища розробки за допомогою Nix, включаючи системні пакети, мовні інструменти, конфігурації IDE, попередній перегляд програм та конфігурацію IDE, а також поділитися проєктом та всією його конфігурацією середовища розробки за допомогою власного шаблону.

– вбудовані інструменти, емулятори та методи розгортання з глибокою інтеграцією з Firebase та Google Cloud: переглядайте свої веб- та Android-додатки прямо у браузері та скористайтеся перевагами вбудованих сервісів та інструментів середовища виконання для емуляції, тестування та налагодження. Firebase Studio легко інтегрується з сервісами Firebase та Google Cloud. Наприклад, ви можете використовувати Firebase Local Emulator Suite безпосередньо з Firebase Studio для ретельного тестування сервісів Firebase та Google Cloud, таких як Firebase Authentication, Cloud Functions, Cloud Firestore, Cloud Storage, Firebase App Hosting та Firebase Hosting, перш ніж публікувати свій додаток.

Firebase Studio підтримує кілька режимів для різних стилів розробки:

– кодування з повним контролем: Працюйте безпосередньо в середовищі розробки на основі Code OSS, де ви можете імпортувати існуючі репозиторії або запускати нові проєкти, а також використовувати розширення з реєстру Open VSX. Gemini у Firebase надає допомогу штучного інтелекту з урахуванням робочого простору з автодоповненням коду, генерацією коду, тестуванням, запуском інструментів та документацією. Ви можете повністю налаштувати свої робочі простори, підхід до розгортання та цільове середовище виконання з підтримкою розширюваної конфігурації за допомогою Nix.

– підказки без кодування: Агент прототипування додатків, також відомий як Prototyper, дозволяє створювати нові робочі простори для прототипування та вдосконалення ідей додатків за допомогою Gemini у Firebase – без написання будь-якого коду. Працюйте з агентом, використовуючи мультимодальні підказки, щоб ітеративно розробляти повноцінний додаток (наразі працює для веб-додатків), тестувати та налагоджувати, а також ділитися своєю роботою з іншими безпосередньо з браузера. Ви можете негайно скасувати зміни, якщо потрібно, додавати нові функції, тестувати, публікувати на Firebase App Hosting та контролювати продуктивність вашого додатка за допомогою вбудованої спостережливості.

Для тесту створено середовище для кодування з повним контролем, при цьому в Gemini введено запит з наступним текстом: «Згенеруй React застосунок який буде мати інтерфейс з використанням AntD та працювати з запитом згідно документації до API: People Data Generator API A FastAPI-based service for generating random person data including names, birthdate, emails and passwords.» Результат запиту показаний на рисунку 1.

В якості результату запиту отримуємо проєкт з попереднім переглядом в браузері на відкритій сторінці студії. Проєкт за замовчування використовує мову програмування TypeScript. Серед недоліків такого використання генерації коду можна виділити те, що в запиті було вказано використання бібліотеки AntDesign в якості React компонентів, але під час генерації система створила лише власні компоненти без під'єднання сторонньої бібліотеки. Так як Firebase Studio все ще знаходиться в experimental стані [1] то в майбутньому можливо буде додана підтримка сторонніх бібліотек.

Підсумовуючи поточний стан продукту та його необхідність можна вважати що Firebase Studio отримуватиме набагато більший функціонал ніж наявний на момент публікації тез.

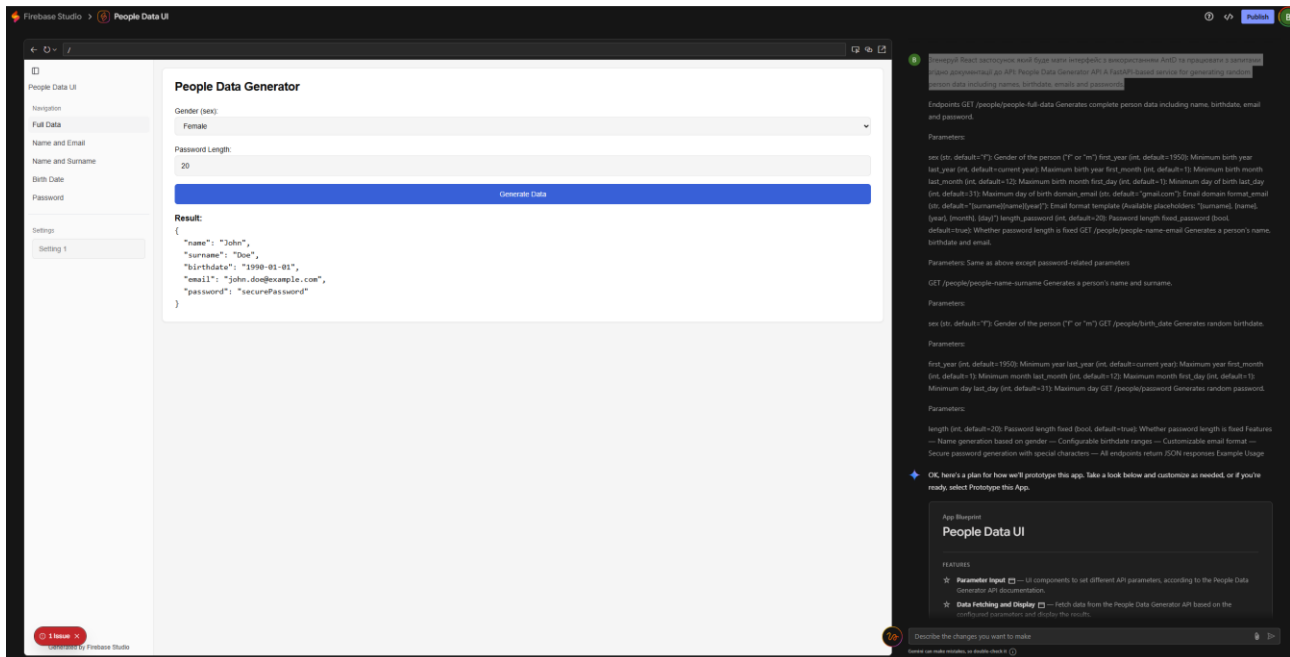


Рис. 1. – Результат запиту до Firebase Studio

References

1. Release Notes | Firebase. URL: <https://firebase.google.com/support/releases> (date of access 28.04.2025)
2. Firebase Studio URL: <https://firebase.google.com/docs/studio> (date of access 28.04.2025)

DOI 10.34132/mspc2025.01.14.14

Viktor Ralenko
Yevhenii Stoiev

USING FIREBASE STUDIO IN DEVELOPING REACT APPLICATIONS

The theses focuses on Firebase Studio, Google's cloud-based platform for building the web and mobile applications, specifically based on React, announced on April 9, 2025. The platform addresses key development challenges: remote workspace availability, limited hardware, paid integration of AI services, and limited number of test devices. Firebase Studio supports importing projects with version control, offers templates for popular frameworks (React, Next.js, Angular, Flutter), rapid prototyping via Gemini with multimodal prompts, integration with Firebase and Google Cloud, and setting up an environment based on Code OSS with Nix support.

The theses demonstrates building a React application with the People Data Generator API, highlighting the limitations in supporting third-party libraries such as Ant Design due to the experimental status of the platform. Firebase Studio has significant potential for expanding functionality, facilitating the development and testing of applications in the cloud, which increases the efficiency and scalability of projects.

Key words: Firebase Studio, React, Cloud development, Google Cloud, Gemini, prototyping, AI, Code OSS, API, Ant Design.

Відомості про авторів / Information about the Authors

Віктор Раленко, викладач кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

Viktor Ralenko, Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

Євгеній Стоєв, викладач кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: stoiev.yevhenii@chmnu.edu.ua, orcid: <https://orcid.org/0009-0008-8999-2267>.

Yevhenii Stoiev, Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: stoiev.yevhenii@chmnu.edu.ua, orcid: <https://orcid.org/0009-0008-8999-2267>.

DOI 10.34132/mspc2025.01.14.15

*Максим Салютін
Інесса Кулаковська*

АСПЕКТИ ЗАСТОСУВАННЯ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ НА ПРИКЛАДІ СКРИПТІВ JAVASCRIPT НА ПЛАТФОРМІ MOODLE

Анотація: Дослідження присвячене застосуванню інформаційно-комунікаційних технологій у навчальному процесі, зокрема використанню скриптів JavaScript на платформі Moodle для викладання дискретної математики. Впровадження цифрових технологій у навчання сприяє інтерактивності, адаптивності та автоматизації оцінювання. Особливу увагу приділено алгоритму Форда-Фалкерсона для знаходження максимального потоку в графах та його практичному застосуванню у моделюванні транспортних і комунікаційних мереж. Використання бібліотек JavaScript, таких як Vis.js та Chart.js, допомагає ефективно візуалізувати графові структури. Moodle надає можливість інтегрувати кастомні плагіни для алгоритмічних розрахунків, що підвищує якість навчання. Впровадження JavaScript у навчальний процес не лише економить час викладачів і студентів, а й сприяє глибшому засвоєнню складних математичних концепцій та підвищує рівень цифрової грамотності студентів.

Ключові слова: інформаційно-комунікаційні технології, JavaScript, Moodle, дискретна математика, візуалізація графів, алгоритм Форда-Фалкерсона, цифрова грамотність

Інформаційно-комунікаційні технології суттєво впливають на всі аспекти суспільного життя, включаючи освіту, бізнес, науку та повсякденне спілкування. Їх активний розвиток не лише змінює способи обробки й передачі інформації, а й сприяє трансформації навчальних методів, які використовуються у вищих навчальних закладах.

Цифрові технології дозволяють студентам та викладачам миттєво отримувати доступ до великої кількості освітніх матеріалів, включаючи наукові статті, навчальні відео, інтерактивні курси та онлайн-бібліотеки. З використанням ІКТ можна створювати адаптивні навчальні програми, які підлаштовуються під потреби студентів, їхні здібності та рівень знань. Сучасні технології дозволяють краще пояснювати складні концепції за допомогою мультимедійних презентацій, симуляцій, віртуальної реальності та доповненої реальності. Навчальні платформи допомагають викладачам автоматично перевіряти завдання, вести оцінювання та зворотний зв'язок у режимі реального часу. Розширюються можливості дистанційної освіти – студенти можуть навчатися онлайн без фізичної прив'язки до місця навчання, що особливо важливо для міжнародних програм та під час кризових ситуацій.

Moodle — це потужна платформа для управління навчанням (LMS), яка широко використовується у закладах освіти для організації дистанційного та змішаного навчання. Вона дозволяє викладачам створювати інтерактивні курси, проводити тестування, організовувати дискусії та відстежувати успішність студентів.

Використання скриптів JavaScript у викладанні дискретної математики з використанням платформи Moodle може покращити освітній процес, зробивши його більш інтерактивним та практичним. JavaScript дозволяє створювати веб-додатки для візуалізації математичних концепцій, автоматизації перевірки завдань. JavaScript чудово підходить для створення автоматизованих систем оцінювання та перевірки завдань. Використання JavaScript допомагає викладачам і студентам економити час, підвищує об'єктивність оцінювання та знижує ймовірність людських помилок. Автоматична перевірка дозволяє миттєво отримувати результати. Об'єктивність – уникнення суб'єктивних факторів в оцінюванні. Масштабованість – можливість перевіряти велику кількість студентських робіт одночасно.

Таблиця 1

Застосування JavaScript у візуалізації

Генерація графів	Побудова таблиць істинності	Візуалізація дерев та мереж
Використання бібліотеки Vis.js для малювання графічних об'єктів.	Автоматичне генерування логічних таблиць істинності на основі введених логічних виразів.	JavaScript може зчитувати графові структури з бази даних і передавати їх до JavaScript для рендерингу.
Використання vis.min.css-бібліотеки, таких як Chart.js або D3.js, у зв'язці з JavaScript для створення інтерактивних графів	Використання HTML і CSS для стилізації та зручного представлення даних	Використання бібліотеки Vis.js для автоматичної побудови графів і дерев.

Для реалізації алгоритмів та візуального супроводження були використані 1) для графів бібліотеки "vis.min.js" - для відображення графів, 2) "vis-network.min.css" для задання візуалу графам. Другою перевагою скриптів JavaScript в Moodle є візуалізація математичних структур – за допомогою JavaScript можна генерувати графи, таблиці істинності та інші математичні об'єкти. JavaScript дозволяє створювати та візуалізувати математичні структури, що особливо корисно при викладанні дискретної математики. За допомогою бібліотек та інструментів можна генерувати графи, діаграми, таблиці істинності та інші математичні об'єкти. Крім того завдяки Vis.js можна задавати координати вузлів вручну (наприклад, через x і y), додавати мітки (labels) для пропускової спроможності та налаштовувати вигляд ребер (стрілки, криві лінії тощо). Це забезпечує інтерактивність, наприклад, масштабування чи перетягування вузлів, що полегшує сприйняття графа студентами. Стосовно досвіду застосувань скриптів до теми алгоритм Форда - Фалкерсона пошуку максимального потоку в графі.

Moodle дозволяє додавати кастомні плагіни для інтеграції алгоритмічних розрахунків. Використання зовнішніх програмних пакетів для автоматизованої перевірки кодів студентів.

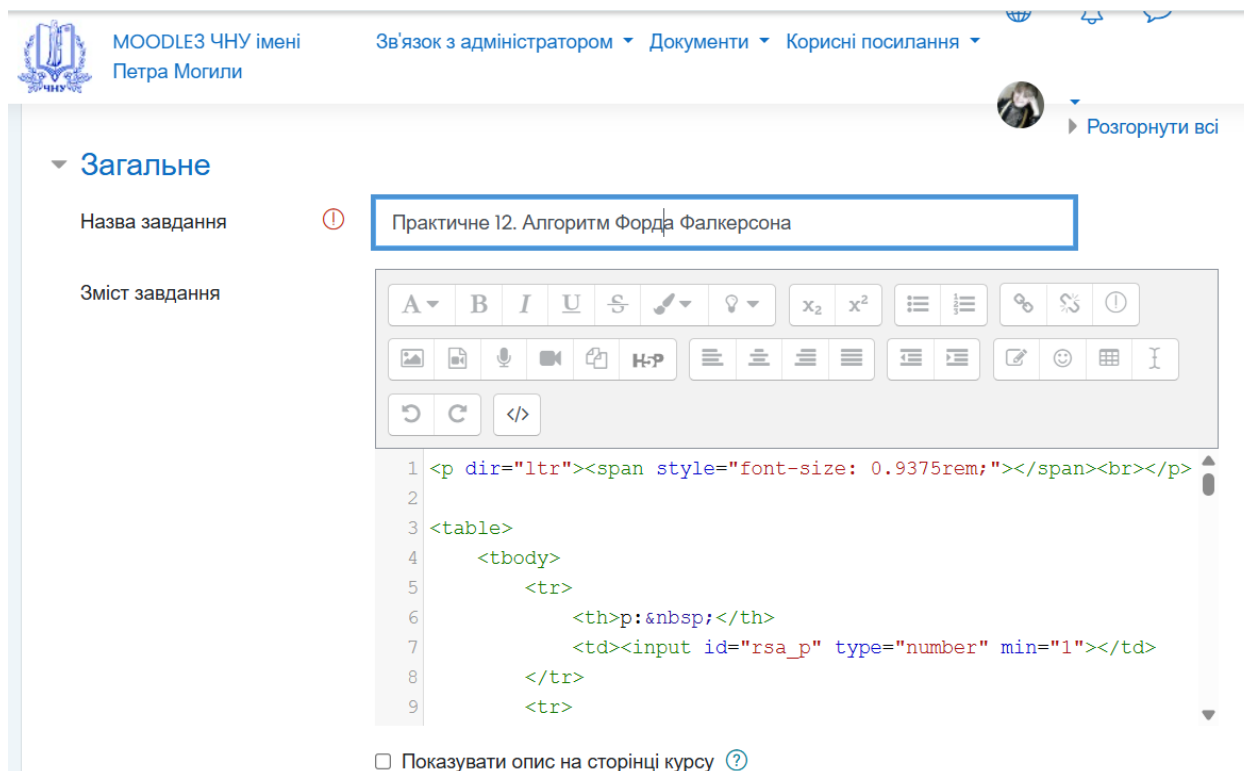


Рис. 3. Подання на платформі Moodle

Потокові графи широко використовуються для розв'язування практичних задач. Однією (дуже поширеною!) задачею є задача про максимальний потік: у поточковому графі знайти максимальний потік, який може бути переданий з виток до стока. Прикладом такої задачі є знаходження максимальної пропускної спроможності комп'ютерної мережі. Задача про максимальний потік розв'язується на основі теореми Форда - Фалкерсона теорему про максимальний потік у графі: величина максимального потоку у графі дорівнює величині пропускної спроможності його мінімального розрізу.

Транспортна мережа – це технічний об'єкт, яким розповсюджуються потоки. Прикладами таких мереж є водогінна мережа, комп'ютерна мережа, мережа автомобільних доріг тощо. Потік - це кількість речовини, енергії або інформації, яка проходить через переріз за одиницю часу.

Моделлю транспортної мережі є зважений граф, у якому задані потоки $v(x,y)$ і пропускні спроможності ребер $c(x,y)$ - максимальні кількості потоку, які можуть проходити через ребра від виток до стока, де x - початкова вершина ребра, y - кінцева вершина. Значення потоків визначаються функцією потоку $f(x,y)$, причому: $f(x,y) \leq c(x,y)$.

Приклад мережі показаний на рис. 1. На рисунку кожне ребро охарактеризоване пропускною здатністю (перша цифра) і величиною потоку (друга цифра). Наведена мережа має один виток, один сток і 4 проміжних вузлів.

Для розуміння прикладних застосувань у завданні є різні типи умов.

Завдання по алгоритму Форда-Фалкерсона

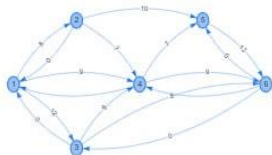
Завдання 1: Система автодоріг

Система автодоріг, що проходять через область, може забезпечити пропускні спроможності, показані на графі (тисячі автомашин на годину).

Питання:

1. Який максимальний потік через цю систему (тис. автомашин на годину)?
2. Скільки автомашин на годину повинно проїхати по дорозі 5-6, щоб забезпечити максимальний потік?

Номер варіанту (0-30):



Максимальний потік (тис. автомашин/годину): Потік дорогою 5-6 (тис. автомашин/годину):

Вірно! Обидві відповіді вірні.

Рис. 1. Приклад мережі для завдань

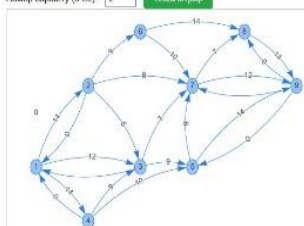
В мережі існує три типи вузлів:

- джерело(виток) S , з якого виходить більше потоку, ніж входить в нього;
- сток T , в який входить більше потоку, ніж виходить з нього;
- проміжні вузли, в які скільки виходить потоку, стільки ж і входить.

Розглянемо *метод Форда-Фалкерсона* визначення максимального потоку, йдеться про метод, оскільки існує декілька алгоритмів, які його реалізують.

Таблиця 2

Метод Форда-Фалкерсона визначення максимального потоку

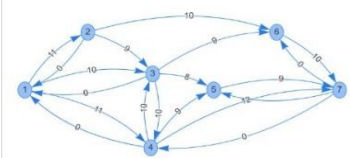
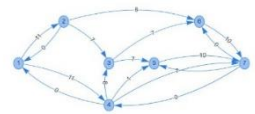
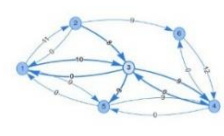
Крок 0. Для кожної дуги орграфа $(k,i) \in E$ встановлюють $x_{k,i} = 0$. Крок 1. Визначають збільшувальний шлях в орграфі. Крок 2. Якщо збільшувального шляху не існує, то потік оптимальний (завершення роботи). У протилежному випадку визначають резерви збільшення/зменшення потоків дугами, які утворюють збільшувальний шлях. Крок 3. Обчислюють значення D, яке дорівнює мінімальному з резервів збільшення/зменшення потоків дугами. Крок 4. Збільшують потік прямими дугами збільшувального шляху на значення D. Крок 5. Зменшують потік зворотними дугами збільшувального шляху на значення D. Крок 6. Переходять на перший крок.	Завдання 2 Завдання 3 Завдання 4 Завдання 5 Завдання 6 Завдання 6: Мережа хімічних труб Хімічний завод має мережу труб, призначених для керування рідинним зв'язком з однієї частини підприємства в іншу. Мережа труб і пропускна спроможність (тис.) показані на малюнку. Питання: 1. Який максимальний потік для системи, якщо завод збирається перекачати з вузла 1 до вузла 9 стільки рідин хімікатів, скільки це можливо? 2. Скільки хімікатів буде надіслати через сполучення 3-5? Номер варіанту (0-30): 0 Оформити граф  Максимальний потік (тис.м³): 10 Потік по сполученню 3-5 (тис.м³): 5 Параметри Вірно! Обидві відповіді вірно.
--	--

Модель мережі з природними умовами для алгоритму Форда-Фалкерсона — це підхід до розв'язання задачі максимального потоку в мережах з урахуванням реальних обмежень, таких як змінна пропускна здатність, втрати на маршруті та інші природні фактори. Алгоритм Форда-Фалкерсона, класично використовуваний для знаходження максимального потоку в графах, може бути модифікований для роботи з реалістичними умовами мережевого середовища. Недоліком алгоритму є залежність від вибору збільшувального шляху, потенційно довгий час роботи.

Ця тема є актуальною для аналізу транспортних, комунікаційних та водних мереж, а також для оптимізації потоків у складних реальних системах.

Таблиця 3

Моделі мережі з природними умовами

Завдання 2 Завдання 3 Завдання 4 Завдання 5 Завдання 6 Завдання 2: Система автодоріг Система автодоріг, що проходить через область, має забезпечити пропускну спроможність, показану на малюнку (кількість автомобілів на годину). Питання: 1. Який максимальний потік через цю систему (кількість автомобілів на годину)? 2. Скільки автомобілів на годину повинно проїхати по дорозі 5-6 щоб забезпечити максимальний потік? Номер варіанту (0-30): 0 Оформити граф  Максимальний потік (тис. автомобілів/годину): 28 Потік по дорозі 5-6 (тис. автомобілів/годину): 12 Параметри Вірно! Обидві відповіді вірно.	Завдання 3: Телефонна мережа Телефонна компанія використовує підземну кабельну мережу для зв'язку між абонентами. Мережа зв'язку між вузлами (вузла 1 і 7 вузла). Переговори здійснюються через серію кабельних ліній і відкривають їх вузли мережі, не це показано на малюнку. На малюнку показано число каналів телефонних переговорів (тис.). Як довго можна відмовитися від однієї з будь-якої ліній? Питання: 1. Яка максимальна кількість телефонних переговорів між двома вузлами, яка може бути додана одночасно (тис. ліній)? 2. Як часто телефонних переговорів має забезпечуватися кабелем 4-7? Номер варіанту (0-30): 0 Оформити граф  Максимальний потік (тис. переговорів): 31 Потік по кабелю 4-7 (тис. переговорів): 12 Параметри Вірно! Обидві відповіді вірно.
Завдання 2 Завдання 3 Завдання 4 Завдання 5 Завдання 6 Завдання 4: Мережа нафтопроводів Нафтова компанія експлуатує мережу нафтопроводів, через яку нафта перекачується від родовища до нафтозаводів. Часова ціна мережі передана на малюнку (пропускна здатність нафтопроводів показана в тис. т/год). Питання: 1. Якщо компанія хоче поставити нафту в споживачі 7 і повністю використувати пропускну здатність системи, то скільки часу займе поставка 10 тис. т. нафти? 2. Який максимальний потік може забезпечити ліній 2-3? Номер варіанту (0-30): 4 Оформити граф  Час поставки 10 тис. т. нафти: 0.45 Потік по лінії 2-3 (тис. т/год): 3 Параметри Вірно! Обидві відповіді вірно.	Завдання 5: Система автодоріг Система автодоріг, що проходить через область, має пропускну спроможність, показану на графі (тис. автомобілів на час). Питання: 1. Чому дорівнює максимальний потік автомобілів (кількість автомобілів на годину) для системи автодоріг, представленої на малюнку? 2. Який максимальний потік може забезпечити дорога 3-4? Номер варіанту (0-30): 8 Оформити граф  Максимальний потік (тис. автомобілів/годину): 27 Потік по дорозі 3-4 (тис. автомобілів/годину): 4 Параметри Вірно! Обидві відповіді вірно.

Для кожного з 30 варіантів генерується індивідуальна мережа, я відповіді перевіряються одразу, що допомагає студентам правильно оцінити отримані навички.

Вибір алгоритму максимального потоку безпосередньо впливає на швидкість обчислень, використання пам'яті та точність отриманих результатів у складних мережах. Тому, на лекціях пояснюється, що алгоритм Форда-Фалкерсона є одним із класичних методів знаходження максимального потоку, але існує кілька інших алгоритмів, які можуть бути ефективнішими в певних ситуаціях. Ось коротке порівняння.

Таблиця 4

Порівняння алгоритмів максимального потоку

	Форд-Фалкерсон	Edmonds-Karp	Дініц	Push-Relabel
Підхід:	Заснований на пошуку збільшувальних шляхів у графі.	Використовує пошук у ширину (BFS) для вибору шляхів.	Використовує рівневий граф і пошук в глибину (DFS).	Працює без використання збільшувальних шляхів, використовує операції підйому і перенесення
Часова складність:	Залежить від вибору шляху (може бути експоненційною в гіршому випадку).	$O(VE^2)$.	$O(V^2E)$..	$O(V^3)$.
Переваги	Простота реалізації, застосовність до будь-яких графів.	Детермінована поведінка та стабільні часові характеристики.	Ефективніший на щільних графах, має кращу оптимізацію.	Підходить для паралельної реалізації та великих графів.

Для конкретної задачі, вибір алгоритму залежить від розміру графа, структури зв'язків і обмежень реального середовища. Якщо потрібно працювати з динамічними потоками, слід враховувати модифікації алгоритмів для змінних пропускних спроможностей. Якщо граф змінюється під час виконання алгоритму (наприклад, у транспортних або комунікаційних мережах), можливо, доведеться застосовувати інші методи, такі як модифікований Push-Relabel.

Для реалізації використовується **function fordFulkerson(taskNumber, n)**

```
function fordFulkerson(taskNumber, n) {
    const edgesTemplate = edgesTemplates[taskNumber];
    const numVertices = taskNumber === 1 ? 7 :
    (taskNumber === 2 || taskNumber === 3 ? 8 :
    (taskNumber === 4 ? 7 : 10)); // Вершини + 1 для
    індексації з 1
    const source = 1;
    const sink = taskNumber === 1 ? 6 :
    (taskNumber === 2 || taskNumber === 3 ? 7 :
    (taskNumber === 4 ? 4 : 9));
    const graph = Array(numVertices).fill(0).map(0) =>
    Array(numVertices).fill(0);
    edgesTemplate.forEach(edge => {
        let capacity = edge.capacity;
        if (capacity.includes('n')) {
            capacity = eval(capacity.replace('n', n));
        }
        graph[edge.from][edge.to] = capacity;
        if (edge.reverse) {
            let reverseCapacity = edge.reverse;
            if (reverseCapacity.includes('n')) {
                reverseCapacity =
                eval(reverseCapacity.replace('n', n));
            }
            graph[edge.to][edge.from] =
            reverseCapacity;
        }
    });

    const residualGraph = graph.map(row => [...row]);
    const parent = new Array(numVertices).fill(-1);
    let maxFlow = 0;
    const flow = Array(numVertices).fill(0).map(0) =>
    Array(numVertices).fill(0);
    while (bfs(residualGraph, source, sink, parent)) {
        let pathFlow = Infinity;
        for (let v = sink; v !== source; v = parent[v]) {
            const u = parent[v];
            pathFlow = Math.min(pathFlow, residualGraph[u][v]);
        }
        for (let v = sink; v !== source; v = parent[v]) {
            const u = parent[v];
            residualGraph[u][v] -= pathFlow;
            residualGraph[v][u] += pathFlow;
            flow[u][v] += pathFlow;
            flow[v][u] -= pathFlow;
        }
        maxFlow += pathFlow;
    }
    const result = taskNumber === 3 ? 10 /
    maxFlow : maxFlow;
    const specificFlow = taskNumber === 2 ?
    flow[4][7] : (taskNumber === 3 ? flow[2][3] :
    (taskNumber === 4 ? flow[3][4] : (taskNumber ===
    5 ? flow[3][5] : flow[5][6])));
    return { result, specificFlow };
}
```

Рис.4. Лістинг алгоритму Форда-Фалкерсона

Висновок: застосування скриптів JavaScript у викладанні дискретної математики значно покращує освітній процес, роблячи його більш інтерактивним, адаптивним і ефективним. Використання JavaScript у платформі Moodle сприяє автоматизації перевірки завдань, підвищенню об'єктивності оцінювання та візуалізації математичних концепцій. Це не лише економить час викладачів і студентів, але й сприяє глибшому розумінню складних тем, таких як алгоритм Форда-Фалкерсона для пошуку максимального потоку в графі.

Завдяки технологічним можливостям JavaScript студенти можуть взаємодіяти з алгоритмами, аналізувати великі обсяги математичних даних та отримувати миттєвий JavaScript зв'язок щодо своїх розрахунків. У результаті впровадження JavaScript у навчальний процес створює умови для ефективного засвоєння знань, підвищуючи рівень цифрової грамотності студентів та сприяючи інтеграції інноваційних підходів у викладанні математичних дисциплін.

References

1. Hong, Z., Xi, D., Zhang, K., & Yan, T. (2024, June). Research on the Water Level Optimization of the Great Lakes Based on Ford-Fulkerson Algorithm. In *Proceedings of the 2024 4th International Conference on Artificial Intelligence, Big Data and Algorithms* (pp. 824-828).
2. Wahyuningsih, S., Octoviana, L. T., & Qohar, A. (2024, September). Performance of the maximum flow algorithm and its application to traffic density. In *AIP Conference Proceedings* (Vol. 3235, No. 1). AIP Publishing.
3. Астаф'єва, М. М., Глушак, О. М., & Литвин, О. С. (2024). Using STACK to support adaptive mathematics learning in LMS Moodle. In *3L-Person 2024: IX International Workshop on Professional Retraining and Life-Long Learning using ICT: Person-oriented Approach, co-located with the 19th International Conference on ICT in Education, Research, and Industrial Applications (ICTERI 2024)* (pp. 30-41). Lviv, Ukraine.

DOI 10.34132/mspc2025.01.14.15

Maksym Salyutin
Inessa Kulakovska

ASPECTS OF THE APPLICATION OF INFORMATION AND COMMUNICATION TECHNOLOGIES USING THE EXAMPLE OF JAVASCRIPT SCRIPTS ON THE MOODLE PLATFORM

Annotation: This research explores the application of information and communication technologies in education, specifically the use of JavaScript scripts on the Moodle platform for teaching discrete mathematics. Digital technologies enhance interactivity, adaptability, and automation in assessment. JavaScript enables the creation of web applications for visualizing mathematical concepts, automating assignment evaluation, and developing educational platforms. Special attention is given to the Ford-Fulkerson algorithm for finding the maximum flow in graphs and its practical implementation in modeling transport and communication networks. JavaScript libraries such as Vis.js and Chart.js effectively visualize graph structures and truth tables. Moodle allows for the integration of custom plugins for algorithmic computations, improving the quality of education. Implementing JavaScript in education not only saves time for instructors and students but also facilitates deeper comprehension of complex mathematical concepts and enhances students' digital literacy. By leveraging JavaScript, students can interact with algorithms, analyze large sets of mathematical data, and receive instant feedback on their calculations, fostering a more engaging and efficient learning experience.

Key words: information and communication technologies, JavaScript, Moodle, discrete mathematics, graph visualization, Ford-Fulkerson algorithm, digital literacy

Відомості про автора / Information about the Author

Максим Салютін, аспірант, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: salyutin1@gmail.com, ORCID: <https://orcid.org/0009-0003-0788-0012>

Maksym Salyutin, postgraduate student, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine, 54003, E-mail: salyutin1@gmail.com, ORCID: <https://orcid.org/0009-0003-0788-0012>

Інесса Кулаковська, кандидат фізико-математичних наук, доцент кафедри Інтелектуальних інформаційних систем, факультет комп'ютерних наук, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: inessa.kulakovska@chmnu.edu.ua, orcid: <http://orcid.org/0000-0002-8432-1850>

Inessa Kulakovska, PhD of Physical and Mathematical Sciences, Associate Professor, Department of Intelligent Information Systems, Faculty of Computer Science, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine, 54003, E-mail: Inessa.Kulakovska@chmnu.edu.ua, ORCID: <http://orcid.org/0000-0002-8432-1850>

DOI 10.34132/mspc2025.01.14.16

Максим Салютін,
Євген Сіденко,
Юрій Кондратенко

СУЧАСНІ ПІДХОДИ ДО РОЗПІЗНАВАННЯ ВІЙСЬКОВИХ ОБ'ЄКТІВ ІЗ ВИКОРИСТАННЯМ МОДЕЛІ YOLOV8

У роботі розглядаються сучасні підходи до розпізнавання військових об'єктів із використанням нейромережевої моделі YOLOv8. Основну увагу приділено модифікаціям *nano* та *small*, які забезпечують оптимальний баланс між швидкістю та точністю. Навчання здійснюється протягом 50 епох на основі підготовленого набору даних, що містить 12 класів різноманітних військових і цивільних об'єктів. Проведено розмітку зображень та аналіз ефективності моделей в умовах наближених до реальних бойових сценаріїв. Отримані результати свідчать про високу придатність легких версій YOLOv8 для виконання оборонних і розвідувальних завдань у режимі реального часу. Окрему увагу приділено можливості інтеграції навчених моделей у безпілотні літальні апарати (БпЛА) для автоматичного виявлення цілей.

Ключові слова: YOLOv8, нейронні мережі, Python, набір даних, розмітка, військові об'єкти.

У контексті зростання масштабів військових конфліктів актуальним завданням є впровадження сучасних технологій в оборонну та розвідувальну діяльність. Одним із перспективних напрямів є використання систем комп'ютерного зору на базі глибокого навчання для автоматизованого виявлення об'єктів. У межах даного дослідження обрано підхід із використанням моделі YOLOv8 – одного із нових поколінь алгоритмів виявлення об'єктів, розробленого компанією Ultralytics.

YOLOv8 відзначається підвищеною точністю та швидкістю порівняно з попередніми версіями, а також підтримкою уніфікованої архітектури для трьох основних задач комп'ютерного зору: виявлення об'єктів, сегментації екземплярів і класифікації зображень. Моделі цього сімейства варіюються за розміром і продуктивністю – від легкої та швидкої YOLOv8n (*nano*) до потужної, проте повільнішої YOLOv8x (*extra large*).

З урахуванням потреб реального часу та обмежень обчислювальних ресурсів для експериментів було обрано варіанти YOLOv8n та YOLOv8s, які поєднують високу швидкість роботи з прийнятною точністю.

Таблиця 1

Порівняння моделей YOLOv8 *nano* та *small* [1]

Model	size (pixels)	mAPval 50-95	Speed CPU ONNX (ms)	Speed A100 TensorRT (ms)	params (M)	FLOPs (B)
YOLOv8n	640	37.3	80.4	0.99	3.2	8.7
YOLOv8s	640	44.9	128.4	1.20	11.2	28.6

Після визначення архітектури моделі та її конфігурації було обрано відкритий датасет для навчання – *Military Assets Dataset* (12 Classes – YOLOv8 Format) <https://www.kaggle.com/datasets/rawsi18/military-assets-dataset-12-classes-yolo8-format>, розміщений на платформі Kaggle. Набір призначений для задач виявлення та класифікації об'єктів у військовому середовищі. Він містить 26 315 розмічених зображень, розподілених на 12 класів та поділених на підмножини:

- навчальна вибірка: 21 978 зображень;
- валідаційна вибірка: 2 941 зображення;
- тестова вибірка: 1 396 зображень.

Список класів:

- *camouflage_soldier* – солдати в маскувальному спорядженні;

- weapon – стрілецька та інша зброя;
- military_tank – бронетехніка з важким озброєнням;
- military_truck – вантажні військові автомобілі;
- military_vehicle – інші військові транспортні засоби;
- civilian – неозброєні цивільні особи;
- soldier – військовослужбовці без камуфляжу;
- civilian_vehicle – цивільні автомобілі;
- military_artillery – великокаліберні артилерійські системи;
- trench – бойові оборонні укріплення;
- military_aircraft – літаки та гелікоптери військового призначення;
- military_warship – бойові кораблі.

Цей набір даних є репрезентативною основою для навчання моделей штучного інтелекту, орієнтованих на автоматизоване виявлення загроз, аналіз ситуацій та підтримку розвідувальних операцій у режимі реального часу.

Після завантаження набору даних було виконано статистичний аналіз міток (рис. 1). На діаграмі розподілу класів (верхній лівий графік) чітко видно сильний дисбаланс: клас military_tank домінує з понад 17 000 зображень, тоді як класи civilian, civilian_vehicle та military_artillery значно непередставлені – менш ніж 1 000 зразків кожен.

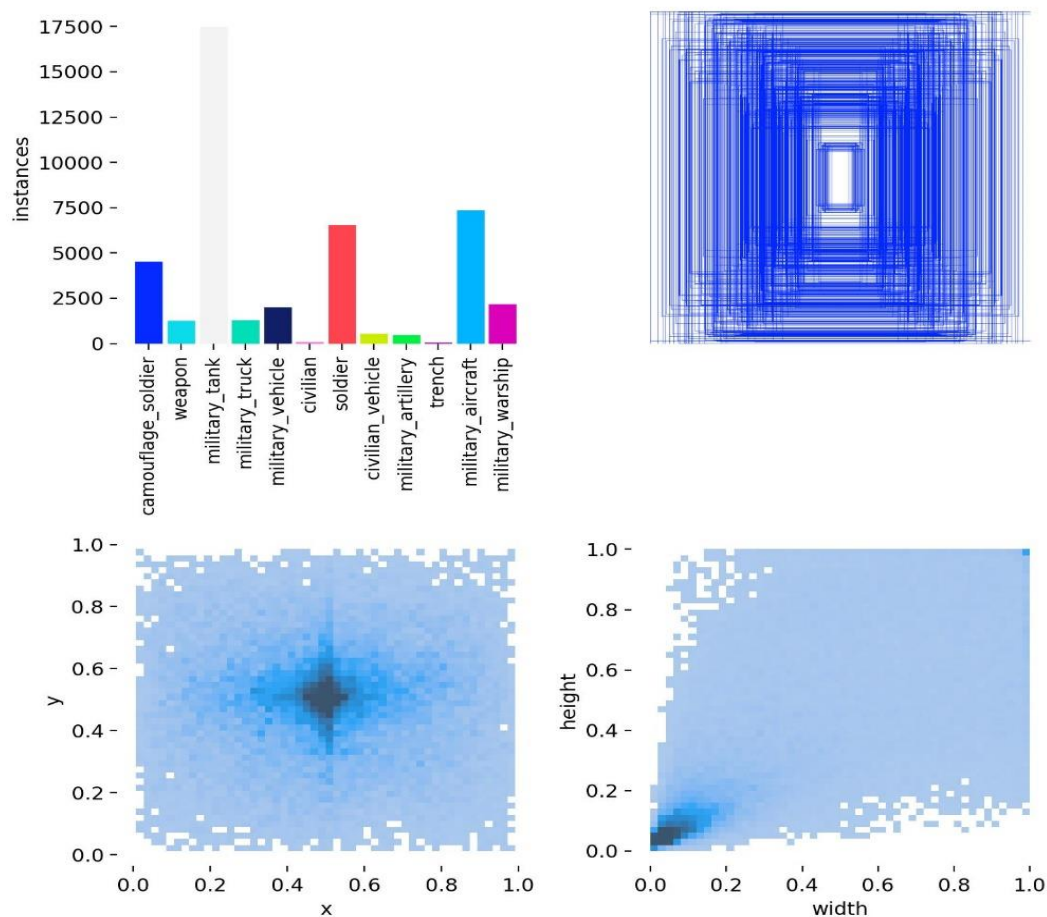


Рис. 1. Інформація про набір даних Military
 Джерело: сформовано автором

Таке співвідношення може спричинити перебіс у навчанні моделі, зниження точності виявлення рідкісних класів та погіршення генералізації.

На графіку в правому верхньому куті візуалізовано накладання всіх об'єктів у наборі даних, що демонструє велику кількість варіантів розташування і розмірів боксів – однак переважають центрально розміщені об'єкти зі стандартними пропорціями.

Два нижні графіки на рис. 1 деталізують просторові характеристики боксів. Графік x/y-координат центру об'єктів свідчить про високу щільність об'єктів поблизу центру зображення ($x \approx 0.5$, $y \approx 0.5$) [2]. Графік ширини/висоти bounding box'ів вказує, що більшість виявлених об'єктів є відносно компактними, з шириною та висотою в межах 0.1-0.2 нормалізованих одиниць.

Під час навчання необхідно враховувати дисбаланс класів та геометричних характеристик об'єктів під час навчання моделей YOLOv8, особливо при виборі функцій втрат і аугментації.

На рис. 2 наведено вміст набору даних з лейблами по закінченню процесу навчання, де номер – це об'єкт класу.

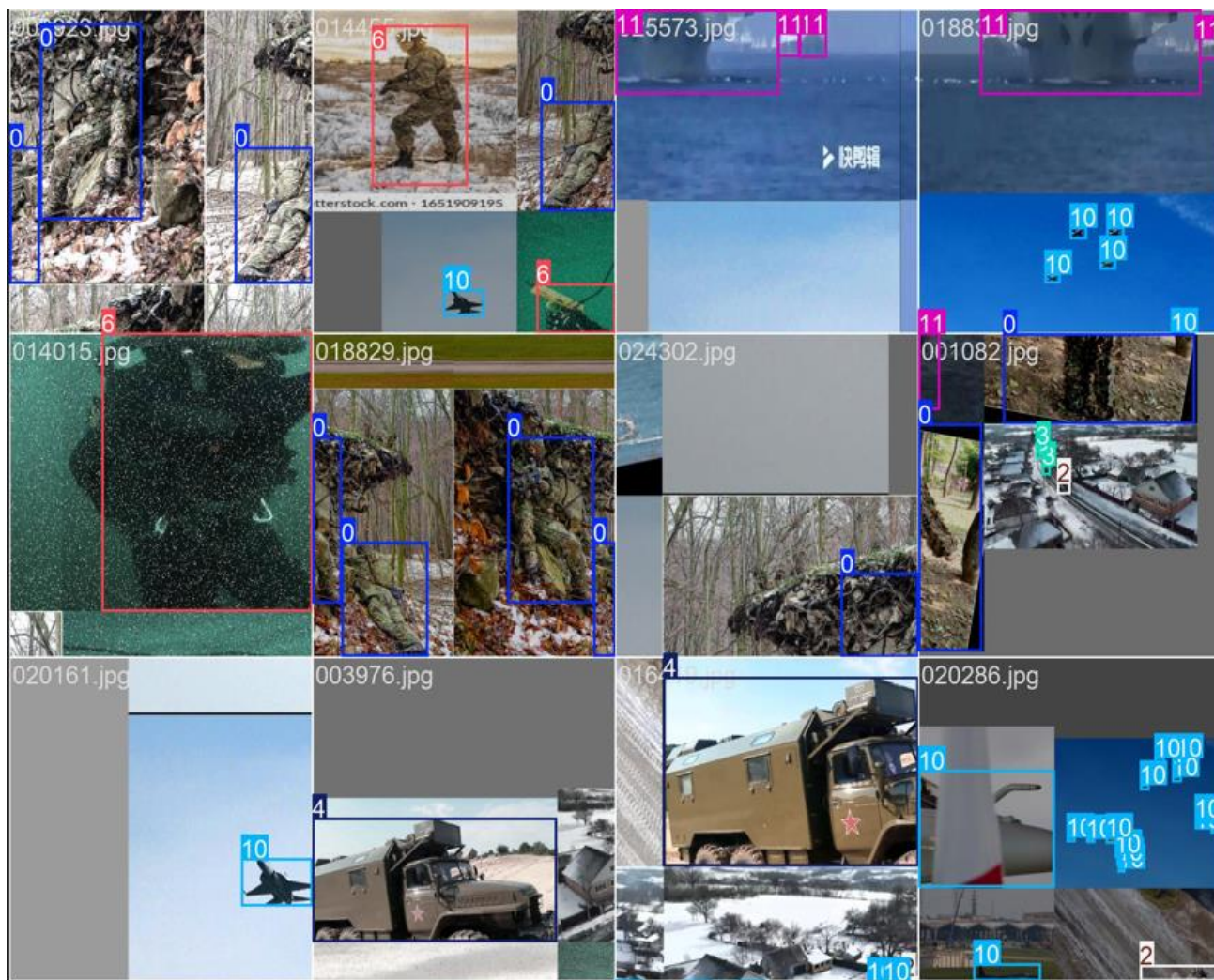


Рис 2. Вміст набору даних з лейблами
Джерело: сформовано автором

Першим кроком стала побудова базової моделі на основі архітектури YOLOv8 nano (YOLOv8n), однієї з найлегших модифікацій, оптимізованих для високої швидкодії з мінімальними вимогами до ресурсів. Навчання моделі здійснювалося на підготовленому наборі даних із 12 класами впродовж 50 епох, що є достатнім для базової первинної конвергенції без перенавчання (рис. 3).

Як оптимізатор використовувався SGD (Stochastic Gradient Descent) зі стандартними параметрами за замовчуванням, передбаченими реалізацією в YOLOv8 [3]. Такий вибір забезпечує стабільне оновлення вагів, хоча й може поступатися адаптивним методам (типу Adam) за швидкістю збіжності на складніших задачах.

Під час навчання здійснювався моніторинг метрик точності (precision, recall, mAP), а також втрат (losses) по класах та по об'єктах, що дозволило сформувати базову оцінку ефективності моделі та виявити потенційні проблеми, пов'язані з дисбалансом даних.

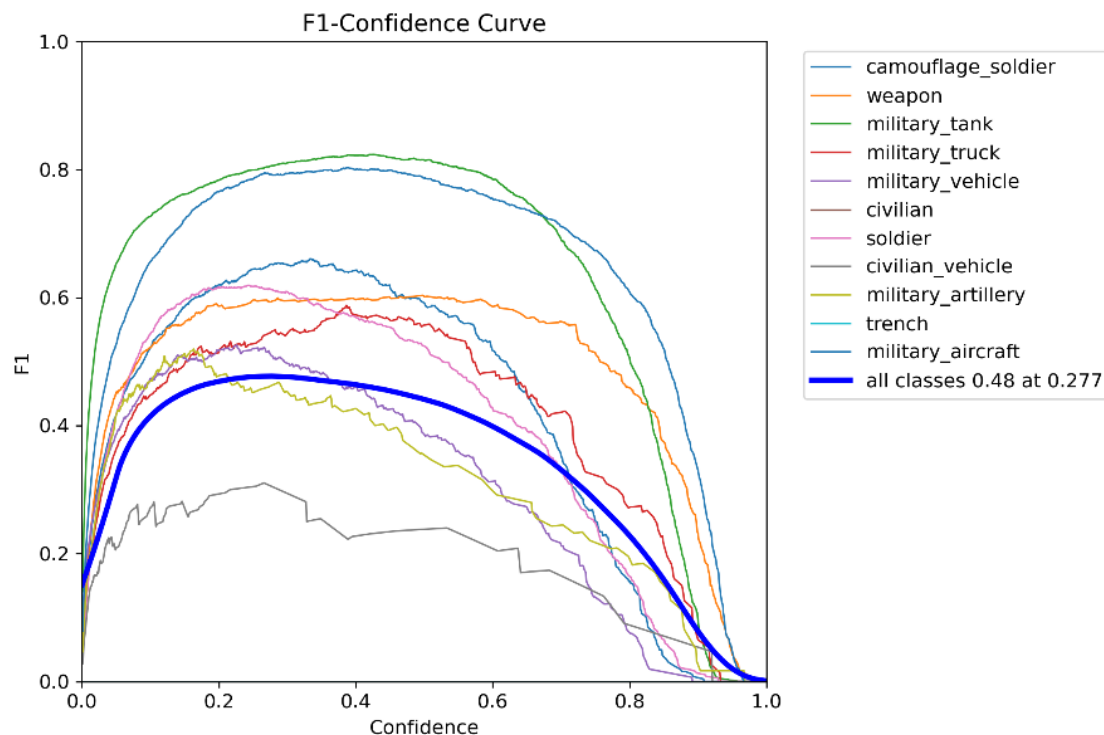


Рис. 3. Результат навчання YOLOv8 nano, 50 епох
Джерело: сформовано автором

F1-Confidence Curve демонструє як змінюється F1-метрика (гармонічне середнє точності та повноти) залежно від порогу впевненості (confidence threshold) [4] моделі для кожного з 12 класів. Найкращі класифікаційні результати модель досягла на класах military_tank, camouflage_soldier та military_aircraft. Слабка ефективність на класах civilian_vehicle і military_artillery – через недостатню кількість зображень у цих класах. Ідеальний поріг для усіх класів – близько 0.28, при якому модель досягає глобального максимуму F1. Загальна F1-метрика 0.48 свідчить про середню якість розпізнавання: модель виявляє об'єкти, але помилки залишаються значними. Цього недостатньо для високоточної автоматизованої системи, особливо в критичних сценаріях (наприклад, у бойових умовах).

Для порівняння було обрано модель YOLOv8 small з аналогічною тривалістю навчання у 50 епох (рис. 4). У результаті F1-score знизився до 0.36 при оптимальному порозі довіри 0.298, що свідчить про деяке погіршення якості розпізнавання об'єктів порівняно з моделлю YOLOv8 nano.

Таке зниження продуктивності вказує на важливість збалансованості класів у наборі даних. Наявність переважання одного класу над іншими, зокрема military_tank, значно впливає на здатність моделі навчатися розпізнаванню менш представлених об'єктів. Окрім цього, збереження сталої кількості епох у поєднанні з використанням оптимізатора SGD за замовчуванням може обмежувати потенціал моделі YOLOv8 small, яка має більшу кількість параметрів і потребує більш гнучких стратегій навчання. Підключення іншого оптимізатора, наприклад AdamW, а також збільшення кількості епох потенційно дозволило б досягти кращих результатів.

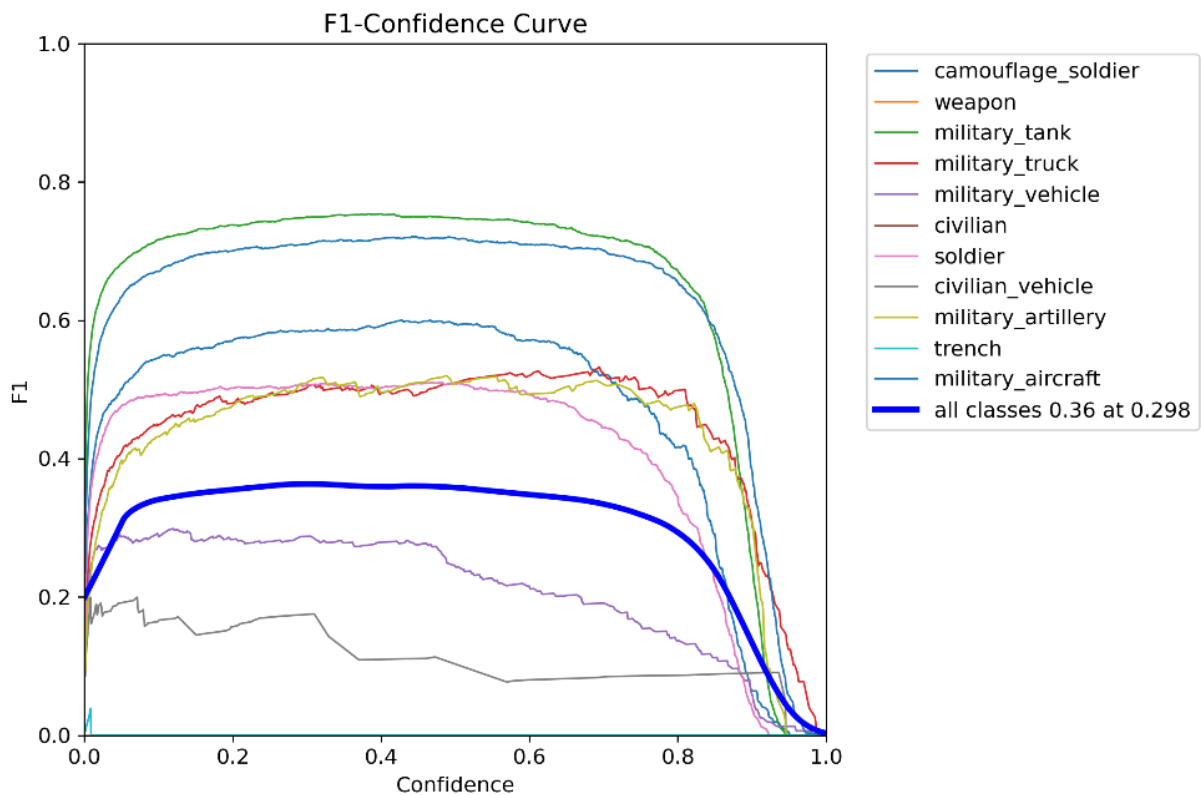


Рис.4. Результат навчання YOLOv8 small, 50 епох
 Джерело: сформовано автором

Результат роботи YOLOv8 папо на тестових даних наведено на рис. 5.



Рис. 5. Результат роботи YOLOv8 папо на тестових даних
 Джерело: сформовано автором

Втім, незважаючи на зниження точності, модель залишається придатною для практичних експериментів. Зокрема, її можна успішно застосовувати у тестових вильотах безпілотних літальних апаратів, де ключовим є не лише абсолютна точність, а й стабільність роботи та здатність працювати в умовах реального часу. Таким чином, навіть із поточними показниками, YOLOv8 nano демонструє перспективність для створення прототипу та подальшого вдосконалення в рамках розробки систем реального часу для виявлення військових об'єктів.

References

1. Alawi, M., & Mohammed, A. (2024). The Role of YOLOv8 in Enhancing Strategic Military Equipment Detection. ResearchGate.
2. Singh, S., & Ratna, G. N. (2024). Military Based Object Detection in Satellite Imagery by Optimising YOLOv8. IEEE Space, Aerospace and Defence Conference.
3. Alawi, M., & Mohammed, A. (2024). Empowering Military Vehicle Detection and Classification with YOLOv8 Model. ResearchGate.
4. Zhao, B., Zhou, Y., Song, R. et al. (2024) Modular YOLOv8 optimization for real-time UAV maritime rescue object detection. Sci Rep 14, 24492.

DOI 10.34132/mspc2025.01.14.16

*Maksym Salyutin,
Ievgen Sidenko,
Yuriy Kondratenko,*

MODERN APPROACHES TO MILITARY OBJECT RECOGNITION USING THE YOLOV8 MODEL

The paper discusses modern approaches to military object recognition using the YOLOv8 neural network model. The main attention is paid to nano and small modifications, which provide an optimal balance between speed and accuracy. Training is carried out over 50 epochs based on a prepared dataset containing 12 classes of various military and civilian objects. The images were labeled and the models' performance was analyzed under conditions close to real combat scenarios. The results obtained indicate the high suitability of lightweight versions of YOLOv8 for real-time defense and reconnaissance tasks. Special attention is paid to the possibility of integrating trained models into unmanned aerial vehicles (UAVs) for automatic target detection.

Keywords: YOLOv8, neural networks, Python, dataset, markup, military objects.

Відомості про автора / Information about the Author

Максим Салютін, аспірант кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: salyutin1@gmail.com, orcid: <https://orcid.org/0009-0003-0788-0012>.

Maksym Salyutin, PhD student of the Intelligent Information Systems Department of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: salyutin1@gmail.com, orcid: <https://orcid.org/0009-0003-0788-0012>.

Євген Сіденко, канд. техн. наук, доцент кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Ievgen Sidenko, PhD, Associate Professor of the Intelligent Information Systems Department of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Юрій Кондратенко, д-р техн. наук, професор, завідувач кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: yuriy.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-7736-883X>.

Yuriy Kondratenko, Doctor of Science, Professor, Head of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: yuriy.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-7736-883X>.

© М. Салютін, Є. Сіденко, Ю. Кондратенко

Стаття отримана редакцією 19.05.2025.

РОЗРОБКА ЛЮДИНО-МАШИННОГО ІНТЕРФЕЙСУ ДЛЯ ПОШУКУ ЗНИКЛИХ ТУРИСТІВ ГРУПОЮ БПЛА

Тези представляють теоретико-методологічний аналіз розробки людино-машинного інтерфейсу (ЛМІ) для пошуку зниклих туристів із застосуванням групи безпілотних літальних апаратів (БПЛА). Людино-машинна взаємодія розглядається як ключовий компонент підвищення ефективності та безпеки під час виконання пошуково-рятувальних операцій у складних географічних, погодних та часових умовах.

У центрі дослідження знаходиться сучасний підхід до створення ЛМІ, орієнтованого на когнітивні можливості оператора, який повинен обробляти великий обсяг інформації в умовах обмеженого часу. Такий підхід включає в себе принципи адаптивного дизайну, ситуаційної обізнаності, візуального упорядкування інформаційних потоків, а також впровадження засобів штучного інтелекту для підтримки прийняття рішень. ЛМІ, що розробляється, є інтегрованою частиною командного центру управління групою БПЛА, об'єднуючи дані з картографічних сервісів, сенсорних систем, відеоаналітики та телеметрії у єдину інформаційно-аналітичну панель.

Наголошується, що поєднання автономної навігації БПЛА та інтуїтивного керування через ЛМІ створює передумови для швидкої локалізації постраждалих осіб та організації ефективного рятування. Актуальність дослідження зумовлена не лише потребами цивільного захисту, а й зростанням обсягів внутрішнього та міжнародного туризму, що супроводжується зростанням кількості нештатних ситуацій у віддалених регіонах.

Ключові слова: БПЛА, пошуково-рятувальні операції, людино-машинна взаємодія, інтерфейс, когнітивна навігація, сенсорна інтеграція, ситуаційна обізнаність, штучний інтелект, автономні системи.

У сучасному світі спостерігається тенденція до зростання кількості туристичних подорожей у важкодоступні природні середовища: гори, ліси, пустелі тощо. Разом із цим підвищується ризик зникнення туристів, що зумовлює потребу в інноваційних засобах пошуку. Традиційні методи, які включають людські пошукові бригади та кінологічні підрозділи, мають обмеження за часом, фізичними ресурсами та доступністю.

Використання БПЛА у таких операціях демонструє значні переваги — можливість швидкого охоплення великої території, оперативне отримання відеозображень та телеметричної інформації, безпека для персоналу. Однак ефективність такої системи значною мірою залежить від способу керування і представлення даних оператору, який приймає рішення в режимі реального часу.

Під людино-машинним інтерфейсом у цьому контексті розуміється програмно-апаратний засіб, що забезпечує візуалізацію, керування та зворотний зв'язок між оператором і групою БПЛА. Основні функції ЛМІ включають:

- відображення геопросторової інформації на основі карти місцевості з накладенням зон пошуку, маршрутів польоту та зон виявлення;
- моніторинг стану кожного БПЛА: рівень заряду акумулятора, висота, швидкість, напрямок руху, канал зв'язку;
- передача відео з камер у режимі реального часу та збереження фрагментів для аналізу;
- автоматичне повідомлення про виявлення об'єктів, що відповідають ознакам людини;
- інтерактивне керування з можливістю видачі команд групі або окремому дрону.

Розроблений інтерфейс (рис. 1) включає поділ екранного простору на логічні блоки: контрольний (стан системи), ситуаційний (карта місцевості), аналітичний (відео та AI-підказки) та блок дій (управління та введення команд).

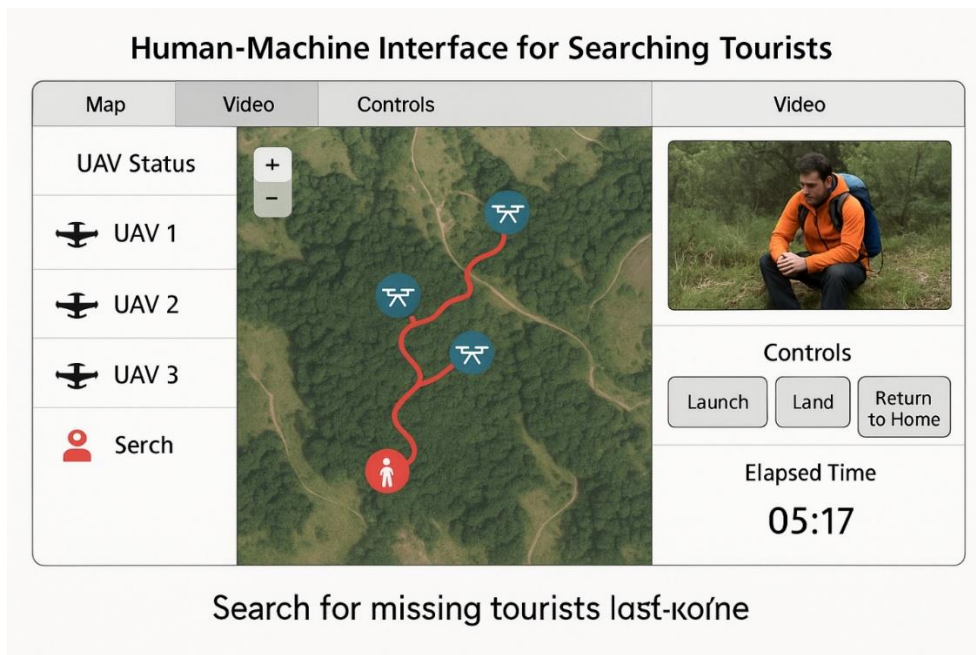


Рис. 1. Людино-машинний інтерфейс для пошуку зниклих туристів групою БПЛА

Людино-машинний інтерфейс розглядається як центральний компонент програмно-апаратного комплексу, що забезпечує ефективну взаємодію між оператором і групою БПЛА в реальному часі. Інтерфейс інтегрується в єдину архітектуру, яка включає в себе взаємодіючі модулі збору, обробки, візуалізації й аналізу даних, а також засоби прийняття рішень у критичних умовах.

До складу комплексу входять такі основні компоненти:

- наземна станція управління (НСУ);
- комунікаційний шлюз;
- група автономних БПЛА з сенсорними модулями;
- хмарна аналітична платформа.

Наземна станція управління (НСУ) — це ноутбук, планшет або захищений мобільний термінал, обладнаний спеціальним програмним забезпеченням, що реалізує функції ЛМІ. Вона є точкою централізованого управління та координації, забезпечує двосторонній зв'язок із БПЛА, прийом телеметрії, трансляцію відеопотоків та надає оператору змогу змінювати параметри місії у реальному часі.

Комунікаційний шлюз забезпечує стійкий та безперервний канал передачі даних між повітряними платформами та наземною станцією. Залежно від умов місцевості використовуються такі технології зв'язку:

- LTE/4G/5G мережі, якщо є покриття в районі операції;
- супутниковий зв'язок (Starlink, Iridium) — у віддалених регіонах;
- динамічні Mesh-сітки, що самовідновлюються, — при групових місіях із великою кількістю апаратів.

Група БПЛА складається з декількох автономних апаратів, які діють як єдина система. Кожен апарат обладнано бортовим комп'ютером, сенсорним комплексом та навігаційною системою. У типовій конфігурації присутні:

- оптична камера високої чіткості;
- тепловізор або інфрачервона камера;
- лазерний далекомір або лідар для об'ємного сканування;
- GPS/ГЛОНАСС навігаційний модуль з резервним інерціальним блоком.

Хмарна аналітична платформа виконує такі функції:

- зберігання історичних даних про польоти та виявлення об'єктів;
- обробка відео з використанням алгоритмів розпізнавання людини;
- машинне навчання на основі раніше виконаних місій;
- автоматичне формування звітів для рятувальників.

Завдяки хмарній обробці інформації значна частина обчислень знімається з оператора та апаратури на борту БПЛА, що дозволяє зменшити затримки в реакції системи.

Інтерфейс спроектовано за принципом багаторівневої взаємодії, що дозволяє одночасно відображати загальну ситуацію, деталі по окремих дронах та інструменти управління.

Основні інформаційні блоки інтерфейсу включають:

- головне картографічне поле із зоною пошуку, маршрутами, позначками тривоги;
- панель моніторингу БПЛА зі статусами: заряд акумулятора; швидкість; висота; координати; стан зв'язку;
- відеопотоки в реальному часі з можливістю перемикання між джерелами;
- блок керування польотами з основними командами: старт; пауза; повернення на базу; зміна маршруту;
- спливаючі повідомлення з автоматичними попередженнями про інциденти: втрату сигналу; перетин меж зони; виявлення підозрілих об'єктів.

Інтерфейс забезпечує адаптивне представлення інформації: залежно від ступеня загрози або важливості подій, змінюється колірна схема, з'являються візуальні попередження, звукові сигнали та вібровідгук (на мобільних пристроях).

Особливістю запропонованого інтерфейсу є наявність вбудованого модуля штучного інтелекту, який виконує роль цифрового помічника. Його функції:

- автоматичне розпізнавання людей на відео з точністю понад 90%;
- аналіз траєкторій руху БПЛА для оптимізації маршрутів покриття;
- рекомендації щодо зміни плану місії у разі виявлення аномалій або вичерпання ресурсу апаратів;
- прогнозування потенційного місця перебування зниклої особи на основі даних про рельєф, погоду, напрямки руху.

Модуль машинного навчання постійно оновлюється під час експлуатації системи, що дозволяє адаптувати поведінку інтерфейсу до специфіки кожної місцевості або категорії пошукових місій.

Для верифікації працездатності інтерфейсу було проведено симуляційні тести за допомогою платформи MATLAB Simulink із моделями польоту групи дронів. Основними метриками стали:

- час до виявлення об'єкта у заданій зоні;
- середній рівень навантаження на оператора (оцінюється за кількістю кліків та часу реакції);
- кількість помилкових спрацювань або пропущених цілей;
- загальна ефективність з точки зору охоплення зони та витрат енергії.

За результатами імітацій та обмежених польових випробувань встановлено, що новий ЛМІ дозволяє скоротити час виявлення зниклого на 25–35% у порівнянні з класичними способами керування.

У ході розробки ЛМІ особлива увага приділялася когнітивним характеристикам оператора, який, як правило, працює у стресових умовах, має обмежений час на ухвалення рішень та часто змушений взаємодіяти з великим обсягом інформації одночасно. Неправильно спроектований інтерфейс у такій ситуації може призвести до критичних помилок, що ставлять під загрозу ефективність усієї пошукової операції.

До ключових принципів когнітивного дизайну, що були реалізовані в інтерфейсі, належать:

- мінімізація надлишкового інформаційного потоку;
- виділення ключових візуальних маркерів через колір;
- пріоритетність подій — автоматичне фокусування на найважливіших ділянках екрану;
- використання знайомих оператору схем і піктограм;
- відсутність «схованих» функцій — кожна дія має бути очевидною.

Дослідження в галузі ергономіки показують, що час реакції людини на візуальні подразники в умовах навантаження може зростати на 40–60%. Саме тому в інтерфейсі реалізовані елементи пасивного контролю:

- автоматичне масштабування карти при виявленні підозрілої активності;
- спливаючі вікна із підказками AI-алгоритмів;
- сигналізація зміною кольору об'єктів (наприклад, жовтий — втрата сигналу; червоний — низький заряд).

Таким чином, ЛМІ не просто інформує, а активно підтримує прийняття рішень, зменшуючи когнітивне навантаження на людину.

Оскільки умови пошуку зниклих людей є унікальними для кожного випадку, критично важливою є здатність інтерфейсу до адаптації. Система проектувалася як масштабована, з можливістю:

- додавання нових БПЛА у групу без перезапуску місії;
- інтеграції нових сенсорів (наприклад, CO₂-датчиків для пошуку людей у завалах);
- зменшення/розширення функціональності залежно від рівня кваліфікації оператора.

Також підтримується декілька режимів роботи:

- «Режим експерта» — всі функції, доступ до алгоритмів, ручне управління маршрутом;
- «Стандартний режим» — тільки ключові дії та автоматичне наведення;
- «Режим навчання» — інтерфейс подає поетапні підказки для новачків.

Завдяки використанню адаптивних шаблонів інтерфейс коригується в реальному часі під потреби користувача, забезпечуючи максимальну ефективність у конкретному контексті.

Запропонований людино-машинний інтерфейс відповідає сучасним вимогам до систем управління БПЛА у пошуково-рятувальних місіях. Його головні переваги:

- когнітивна простота та зручність навігації;
- модульність та гнучкість налаштувань;
- інтеграція штучного інтелекту для аналітики даних;
- адаптація під різні сценарії використання.

На основі проведених симуляцій та обмежених польових тестів можна стверджувати, що ЛМІ підвищує ефективність і безпечність пошукових операцій, скорочуючи час реагування та оптимізуючи ресурси.

Подальші дослідження мають бути спрямовані на:

- вдосконалення алгоритмів машинного навчання для більш точного розпізнавання;
- розширення підтримки сенсорних та теплових модулів;
- підвищення кібербезпеки системи зв'язку;
- розробку мобільної версії ЛМІ для роботи на планшетах і смартфонах у польових умовах.

References

1. Haider SK, Nauman A, Jamshed MA, Jiang A, Batool S, Kim SW(2022) Internet of drones: routing algorithms. Tech Chall Math10(9):1488.
2. Sabino, Hullysses, Rodrigo V. S. Almeida, Lucas Baptista de Moraes, Walber Paschoal da Silva, Raphael Guerra, Carlos Malcher, Diego Passos, and Fernanda G. O. Passos. "A systematic literature review on the main factors for public acceptance of drones." Technology in Society (2022): 102097.
3. Chen, Long, Jianghao Wu, and Bo Cheng. "Leading-edge vortex formation and transient lift generation on a revolving wing at low Reynolds number." Aerospace Science and Technology 97 (2020): 105589.

DOI 10.34132/mspc2025.01.14.17

Oleksandr Skakodub

DEVELOPMENT OF A HUMAN-MACHINE INTERFACE FOR SEARCHING MISSING TOURISTS BY A UAV GROUP

The theses present a theoretical and methodological analysis of the development of a human-machine interface (HMI) for searching missing tourists using a group of unmanned aerial vehicles (UAVs). Human-machine interaction is viewed as a key component for improving the efficiency and safety of search and rescue operations under challenging geographic, weather, and temporal conditions.

At the core of this study lies a modern approach to designing an HMI focused on the cognitive capabilities of the operator, who must process a large volume of data within a limited timeframe. This approach incorporates principles of adaptive design, situational awareness, structured visualization of information flows, and the implementation of artificial intelligence tools to support decision-making. The developed HMI is an integrated part of a UAV group control command center, combining data from mapping services, sensor systems, video analytics, and telemetry into a unified informational-analytical dashboard.

It is emphasized that the combination of autonomous UAV navigation and intuitive control via HMI creates the foundation for rapid localization of affected individuals and the organization of effective rescue efforts. The relevance of the study stems not only from the needs of civil protection but also from the increasing volume of domestic and international tourism, which is accompanied by a growing number of emergency situations in remote regions.

Key words: UAV, search and rescue operations, human-machine interaction, interface, cognitive navigation, sensor integration, situational awareness, artificial intelligence, autonomous systems.

Відомості про автора / Information about the Author

Олександр Скакодуб, викладач кафедри інтелектуальних інформаційних систем Чорноморського національного університету імені Петра Могили, Миколаїв, Україна. E-mail: skakodub@chmnu.edu.ua, orcid.org/0000-0002-5643-775X

Oleksandr Skakodub, Lecturer at the Department of Intellectual Information Systems, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: skakodub@chmnu.edu.ua, orcid.org/0000-0002-5643-775X

© О. Скакодуб

Стаття отримана редакцією 19.05.2025.

DOI 10.34132/mspc2025.01.14.18

Микита Смоленський,
Євген Сіденко

ФОРМАЛІЗОВАНІ ПІДХОДИ ДО ІНТЕГРАЦІЇ RFID-ТЕХНОЛОГІЙ У ЦИФРОВУ ІНФРАСТРУКТУРУ МЕДИЧНИХ ЗАКЛАДІВ

У роботі представлено формалізований підхід до впровадження RFID-систем у сфері охорони здоров'я, заснований на математичному моделюванні критичних параметрів і вимог до точності, часу відгуку та захищеності даних. Підкреслюється важливість стандартизації технічних характеристик RFID-засобів відповідно до нормативів ISO/IEC, HL7 та IEC 60601. Наведено формули для оцінки ефективності зчитування та часу реакції системи, а також логічну схему покриття ключових зон медичного процесу. Обґрунтовано необхідність інтеграції RFID з існуючими інформаційними платформами на основі HL7 FHIR. Зазначено, що запропоновані підходи дозволяють попередньо оцінити RFID-інфраструктуру перед реальним впровадженням і можуть слугувати основою для подальшої автоматизації клінічних аудитів.

Ключові слова: RFID, охорона здоров'я, ідентифікація пацієнтів, формалізація, стандарти ISO, HL7, FHIR, час зчитування, точність, безпека даних.

Технології радіочастотної ідентифікації (RFID) мають потенціал значно підвищити ефективність і безпеку клінічних процесів, однак на практиці їх впровадження часто супроводжується технічними невідповідностями та недостатнім урахуванням вимог середовища. Актуальним завданням є розробка формалізованої системи технічних вимог, яка дозволить не лише обрати відповідний тип RFID-рішення, а й оцінити його відповідність до впровадження на етапі моделювання.

У роботі запропоновано набір метрик та аналітичних критеріїв, які можуть бути використані для кількісного аналізу RFID-систем у медичних умовах. Однією з базових величин є інтегральна ймовірність точного зчитування тегів у конкретному середовищі:

$$P_{read} = \frac{N_{valid}}{N_{total}}, P_{read} \geq 0.99, \quad (1)$$

де N_{valid} — кількість правильно зчитаних міток, N_{total} — загальна кількість спроб зчитування. Окрім цього, вводиться коефіцієнт латентності відповіді:

$$T_{resp} = t_{recv} - t_{scan}, T_{resp} \leq 500 \text{ мс} \quad (2)$$

Запропонована модель враховує вплив частотного діапазону, типу мітки, середовища (наявність рідин, металів), а також просторової конфігурації зчитувачів. Граничні умови встановлюються відповідно до стандартів ISO/IEC 18000-6C, ISO 14443 та IEC 60601 для медичних електронних систем [1].

Зокрема, для UHF-міток (860–960 МГц) було враховано, що ефективна дальність зчитування знижується до 60–70% у присутності вологи або металевих конструкцій. Це накладає обмеження на їх застосування поблизу пацієнтів і вимагає точного розміщення антен. У свою чергу, HF-мітки (13,56 МГц), хоч і мають меншу дальність, демонструють стабільнішу роботу в таких умовах і можуть бути рекомендовані для ідентифікації пацієнтів у ліжковому фонді.

На Рисунок 1 наведено модель типової RFID-інтеграції в лікарні, яка охоплює вузли ідентифікації пацієнта, контроль медикаментів, прив'язку до процедур та ведення історії хвороби. Кожен блок має бути пов'язаний із цільовими процедурами – сканування браслета, відстеження ліків, реєстрація дій персоналу – і відповідати інтерфейсам HL7/FHIR [2, 3].

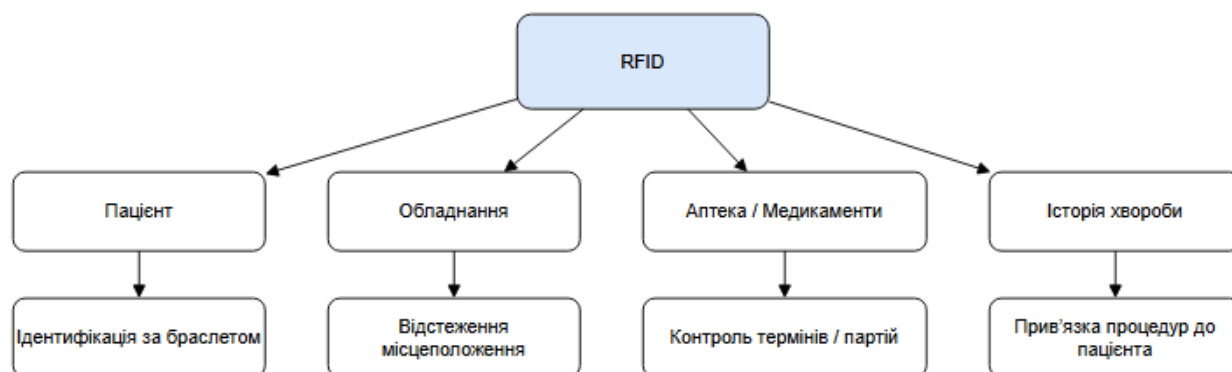


Рис. 1. Діаграма інтеграції RFID
 Джерело: сформовано автором

Пропонується структурувати архітектуру таким чином, щоб критичні події (видача препарату, початок процедури, зміна розташування пацієнта) супроводжувалися RFID-подією, що потрапляє у центральну інформаційну систему.

Розроблена модель формалізації також враховує можливість автоматизованого аудиту RFID-інфраструктури. Пропонується ввести набір KPI, зокрема:

– коефіцієнт критичного часу ідентифікації:

$$K_{crit} = \frac{N_{<300}}{N_{total}}, \quad (3)$$

де $N_{<300}$ — кількість зчитувань з відгуком до 300 мс;

– коефіцієнт інтероперабельності: частка подій, які успішно інтегрувались у HL7-повідомлення, розраховується через аудит логів системи.

Ці метрики можуть бути використані для прийняття рішень на рівні IT-відділів лікарень, а також у майбутніх державних закупівлях як критерії ефективності.

Задля перевірки відповідності RFID-системи до реальних умов, було змодельовано тестовий сценарій «Реєстрація введення препарату». У тесті з використанням HF-міток середній час зчитування складав $T_{resp}=240\pm35$ мс при 98,7% успішності. При цьому встановлено, що рівень хибних спрацювань (зчитування стороннього тега) не перевищував 0,3% за умови екранування антени зони сканування.

Таблиця 1.

Результати тестування RFID-сценаріїв у клінічних умовах

№	Тип мітки	Умови середовища	N_{valid}	N_{total}	P_{read}	Середній T_{resp} , мс	K_{crit}
1	HF (13,56 МГц)	Палата без перешкод	295 / 300	300	0.983	240	0.963
2	HF (13,56 МГц)	Металеве ліжко поруч	287 / 300	300	0.957	285	0.893
3	UHF (868 МГц)	Відкрите приміщення	299 / 300	300	0.997	170	0.993
4	UHF (868 МГц)	Металева шафа на фоні	274 / 300	300	0.913	230	0.870
5	HF + екран	Зона сканування локалізована	299 / 300	300	0.997	215	0.897

Джерело: сформовано автором

У таблиці 1 наведено результати тестування RFID-зчитування для кількох типових сценаріїв. Як видно з розрахунків, HF-мітки забезпечують стабільні результати в умовах, близьких до пацієнта, однак їх чутливість до металевих предметів дещо знижує точність. У той же час UHF-мітки демонструють вищу швидкість та точність

у відкритих зонах, але мають гірші результати в умовах відбиття сигналу. Найкращі результати показав сценарій з HF-міткою та екрануванням зони зчитування, де точність зчитування досягла 99,7%, а час реакції системи залишився в межах 215 мс, що відповідає критичним межах (2).

Додатково для кожного сценарію було обчислено коефіцієнт критичного часу ідентифікації K_{crit} (3), який показує частку зчитувань, виконаних із часом відповіді до 300 мс. Як видно з таблиці 1, сценарії з HF + екраном та UHF у відкритому приміщенні демонструють найвищу ефективність, з коефіцієнтом понад 0.98. У складніших умовах (наприклад, UHF при металевому фоні) цей показник знижується до 0.87, що є граничним рівнем, нижче якого система може бути визнана критично чутливою до затримок.

Окрім технічних вимірів, у роботі проведено порівняльний аналіз нормативної бази, прийнятої у США (FDA UDI Rule, HIPAA), ЄС (EU MDR + GDPR), та України (ISO/IEC адаптації через накази МОЗ та ДСТУ). Особливу увагу приділено стандарту ISO/IEC 29167, який встановлює механізми шифрування на рівні радіоінтерфейсу. Його реалізація дозволяє уникати несанкціонованого зчитування навіть у публічному середовищі. В європейському контексті з 2021 року вимагається використання унікального Device Identifier (UDI), який часто реалізується у вигляді RFID-мітки з підтримкою GTIN. Це забезпечує простежуваність обігу медичних виробів і гарантує сумісність між різними медичними установами.

Наукова новизна роботи полягає в математичному уточненні критичних параметрів RFID-систем, застосовуваних у лікарнях, та введенні формалізованого ядра вимог у вигляді граничних формул. Результати можуть бути використані як основа для проектування інтелектуальних RFID-систем контролю, а також для створення інструментів автоматизованого аудиту відповідності RFID-інфраструктури вимогам регуляторних стандартів.

У роботі представлено модель формалізації RFID-систем медичного призначення з урахуванням технічних, середовищних і нормативних параметрів. Запропоновано аналітичні формули, сценарії тестування та логічну структуру застосування RFID у закладах охорони здоров'я. Результати можуть бути адаптовані до будь-якого типу медичного закладу, який прагне інтегрувати цифрові рішення без шкоди для клінічної безпеки. Практична значущість полягає в можливості об'єктивної попередньої оцінки RFID-інфраструктури, що дає змогу скоротити ризики впровадження та зменшити залежність від конкретного виробника рішень.

References

1. Wamba, S.F., Anand, A., & Carter, L. (2013). A literature review of RFID-enabled healthcare applications and issues. *International Journal of Information Management*, 33(5), (pp. 875–891) [in English].
2. Gazzarata, R., Almeida, J., Lindsköld, L., Cangioli, G., Gaeta, E., Fico, G., & Chronaki, C.E. (2024). HL7 Fast Healthcare Interoperability Resources (HL7 FHIR) in digital healthcare ecosystems for chronic disease management: Scoping review. *International Journal of Medical Informatics*, 189, (Article 105507) [in English].
3. Dridi, A., Sassi, S., Chbeir, R., & Faiz, S. (2020). A flexible semantic integration framework for fully-integrated EHR based on FHIR standard. *Proceedings of the 12th International Conference on Agents and Artificial Intelligence (ICAART 2020)*, (pp. 684–691) [in English].

DOI 10.34132/mspc2025.01.14.18

Mykyia Smolenskyi,
Ievgen Sidenko,

FORMALIZED APPROACHES TO INTEGRATING RFID TECHNOLOGIES INTO THE DIGITAL INFRASTRUCTURE OF HEALTHCARE INSTITUTIONS

This paper introduces a formalized approach to implementing RFID systems in healthcare, based on mathematical modeling of critical performance parameters such as reading accuracy, response time, and data protection. The importance of standardizing RFID technical specifications in accordance with ISO/IEC, HL7, and IEC 60601 is emphasized. The paper includes formulas for calculating system efficiency and latency, along with a logical diagram covering core areas of clinical workflow. The integration of RFID systems with existing medical information platforms using HL7 FHIR is justified. The proposed methodology enables pre-deployment evaluation of RFID infrastructure and may serve as a foundation for future clinical audit automation.

Key words: RFID, healthcare, patient identification, formalization, ISO standards, HL7, FHIR, reading time, accuracy, data security.

Відомості про автора / Information about the Author

Микита Смоленський, аспірант кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: nikitsmolenskyi@gmail.com, orcid: <https://orcid.org/0009-0008-3071-4350>.

Mykyta Smolenskyi, PhD student of the Intelligent Information Systems Department of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: nikitsmolenskyi@gmail.com, orcid: <https://orcid.org/0009-0008-3071-4350>.

Євген Сіденко, канд. техн. наук, доцент кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Ievgen Sidenko, PhD, Associate Professor of the Intelligent Information Systems Department of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

DOI 10.34132/mspc2025.01.14.19

*Іван Сова,
Олексій Козлов*

АДАПТИВНІ МЕТОДИ МАШИННОГО НАВЧАННЯ ДЛЯ РОЗПІЗНАВАННЯ БПЛА В УМОВАХ ЗМІННОГО РАДІОЧАСТОТНОГО СЕРЕДОВИЩА

Тези присвячено аналізу проблеми розпізнавання безпілотних літальних апаратів (БПЛА) за радіочастотними ознаками в умовах змінного середовища. Наголошується, що традиційні моделі машинного навчання, навчені у статичних умовах, демонструють зниження продуктивності при зміні параметрів каналу, положення антен, впливі шумів і мультипасових ефектів. Обґрунтовується доцільність застосування адаптивних методів, зокрема адаптації до середовища (domain adaptation), які дозволяють зберігати точність без повного перенавчання. Запропоновано архітектуру системи з модулем виділення ознак, дискримінатором домену та симулятором змін середовища. Такий підхід відкриває перспективи для створення стійких до змін систем моніторингу БПЛА у військовій та цивільній сферах.

Ключові слова: БПЛА, розпізнавання, радіочастотний аналіз, адаптація до домену, машинне навчання, змінне середовище, RF fingerprinting.

Зі стрімким поширенням безпілотних літальних апаратів (БПЛА) у повсякденному житті, від рекреаційних польотів до розвідувальних і бойових місій, постає необхідність у побудові надійних систем їх виявлення та розпізнавання. Особливої актуальності ця задача набуває в умовах обмеженого доступу до візуальної або акустичної інформації, що притаманно як міському середовищу, так і театру бойових дій. Одним із напрямів, який розвивається в цьому контексті, є радіочастотне (RF) розпізнавання, що базується на аналізі випромінювання БПЛА або його зв'язку з оператором.

Однак, реалізація ефективних RF-систем стикається з кількома суттєвими труднощами. Найпоширенішою серед них є зміна середовища, тобто варіації каналу, положення передавачів і приймачів, присутність перешкод та інші фактори, які впливають на вигляд отриманого сигналу. Моделі машинного навчання, що використовуються для розпізнавання на основі таких даних, зазвичай демонструють зниження продуктивності, щойно умови тестування відрізняються від умов, у яких вони були навчені. Це явище, зміна середовища (domain shift), є загальною проблемою для ML-систем, але особливо гостро воно постає в задачах RF-розпізнавання.

Розпізнавання БПЛА на основі RF-сигналів за останнє десятиліття стало окремим підполем у задачах бездротової безпеки. Найбільш поширеним підходом є побудова моделей класифікації на основі IQ-семплів або спектрограм сигналу [1]. Для цієї мети використовують згорткові нейронні мережі (CNN), які добре зарекомендували себе у задачах розпізнавання образів [2]. У ряді робіт ці моделі доповнюються рекурентними компонентами, зокрема LSTM, що дозволяє враховувати часову динаміку передачі [3].

Іншим активно досліджуваним напрямом є так зване *RF fingerprinting*, або створення «відбитку» пристрою. Йдеться про використання дрібних особливостей у формуванні сигналу, які виникають через неточності електроніки: варіації у підсилювачах, генераторах, модуляторах. Навіть два однакові за моделлю дрони, у теорії, можуть бути розпізнані за своїм унікальним електромагнітним «почерком». Ці методи часто комбінують з CNN або автоенкодерами [4].

Однак, усі згадані рішення значною мірою залежать від умов, у яких здійснюється запис. Наприклад, зміна антен або розміщення приймача навіть на кілька метрів здатна повністю змінити вигляд сигналу у часовій або частотній області. Як наслідок, статично навчені моделі погано узагальнюються на нові сценарії.

Для подолання цього ефекту у суміжних галузях активно застосовуються методи *адаптації до середовища* (domain adaptation) – підходи, що дозволяють навчити модель бути інваріантною до ознак середовища. Найвідомішим є підхід DANN (Domain-Adversarial Neural Network), де мережа одночасно навчається виконувати класифікацію та «забувати» про домен через зворотну оптимізацію дискримінатора [5]. Візуально це виглядає як гра в «перетягування канату» між класифікатором і доменним розпізнавачем. Проте, у контексті RF-аналізу БПЛА цей клас методів поки що мало представлений у літературі, попри очевидну доцільність.

Таблиця 1 надає порівняльну характеристику наведених підходів для розпізнавання БПЛА за допомогою радіочастотної інформації.

Таблиця 1.

Порівняльна характеристика підходів до RF-розпізнавання БПЛА

Підхід	Тип вхідних даних	Стійкість до змін середовища	Необхідність перенавчання	Коментар
CNN на спектрограмах	Частотні зображення	Низька	Висока	Швидко деградує при зміні каналу
LSTM на IQ-сигналах	Часові ряди	Низька–середня	Висока	Чутливе до флуктуацій у фазі
RF fingerprinting	IQ/амплітудні профілі	Середня	Середня	Працює краще для ідентифікації
Domain adaptation (DANN)	Ознаки з CNN	Висока	Низька	Потребує невідмічених даних з нового середовища
Data augmentation (канали)	Будь-які	Середня–висока	Залежить від варіанту	Може застосовуватись до інших моделей

Джерело: сформоване автором на основі [2 - 5].

У задачах розпізнавання за RF-даними стабільність середовища не може бути гарантована навіть протягом короткого проміжку часу. Змінність може виникати як у фізичному просторі (зміна розміщення антен, перешкоди в каналі, відбиття від навколишніх об'єктів), так і у частотному або часовому вимірах.

До основних типів змін належать:

- *Часові варіації*: затримки сигналу, зсуви частоти, пов'язані зі зсувами генератора або рухом.
- *Багатопроменевість (multipath)*: сигнал доходить до приймача різними шляхами з різними затримками і фазами.
- *Погодні фактори*: дощ, вологість, туман можуть змінювати поглинання і відбиття сигналу.
- *Апаратні особливості*: калібрування антен, налаштування підсилювачів, фільтри.

У практичних експериментах відзначено, що навіть незначне переміщення приймальної антени (на 30–50 см) у середовищі зі значною кількістю відбиттів може істотно змінити характер спектру сигналу. Це особливо критично для моделей, натренованих лише на «чистих» даних, отриманих у контрольованих умовах. Без урахування цього різновиду змін створення стійких до помилок систем стає практично неможливим.

Враховуючи зазначені обмеження, пропонується архітектура, яка поєднує класичний модуль виділення ознак (наприклад, згорткова нейронна мережа) з модулем адаптації до домену. Основна ідея полягає у побудові системи, здатної розрізняти корисні ознаки (які відповідають класу БПЛА) і ігнорувати ті, що несуть лише інформацію про середовище. Блок-схема архітектури наведена на рисунку 1.

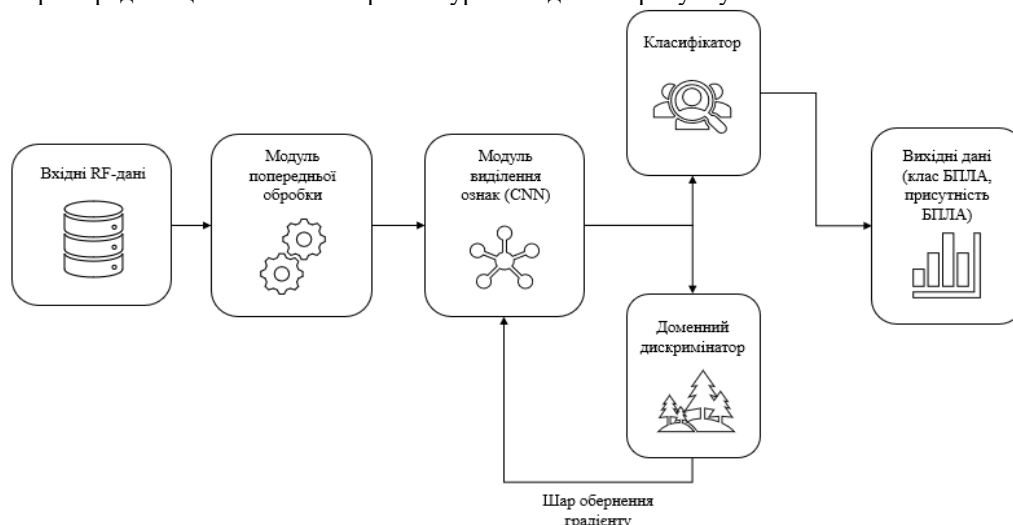


Рис. 1. Архітектура адаптивної системи RF-розпізнавання БПЛА

Джерело: сформоване автором на основі [5].

Архітектура включає такі компоненти:

1. *Модуль виділення ознак (CNN)* — обробляє спектрограми або IQ-сигнали, виділяючи ознаки.
2. *Класифікатор* — класифікує об'єкти (тип БПЛА або факт наявності).
3. *Доменний дискримінатор* — вчиться визначати, з якого середовища дані (середовища тренування або нового).
4. *Шар обернення градієнту* — зв'язує ознаки з дискримінатором і реалізує змагальне навчання.

Під час тренування, модель намагається бути «хорошим» класифікатором і водночас — «поганим» предиктором середовища. Такий механізм дозволяє їй навчитися витягати тільки релевантні ознаки.

Запропонований підхід відкриває шлях до створення систем розпізнавання БПЛА, здатних ефективно працювати у широкому спектрі середовищ без потреби у повному перенавчанні на кожен новий сценарій. У реальних умовах, де збір нових даних є дорогим, тривалим або взагалі неможливим, така властивість є критично важливою.

Сфери застосування охоплюють:

- Тактичні системи радіомоніторингу в зонах бойових дій;
- Системи охорони критичної інфраструктури;
- Безпекові рішення в урбанізованих просторах;
- Автономні сенсорні вузли з обмеженою обчислювальною потужністю.

У майбутньому передбачається реалізація програмно-апаратного прототипу, тестування системи у варіативних реальних умовах, а також дослідження можливостей її розширення для вирішення задач мультикласової класифікації, зокрема розпізнавання типу дрона, моделі, або навіть індивідуального пристрою за RF-відбитком.

References

1. Jafar N., Paeiz A., Farzaneh A. Automatic modulation classification using modulation fingerprint extraction. Journal of Systems Engineering and Electronics. 2021. Vol. 32, no. 4. P. 799–810.
2. Optimized Radio Frequency Footprint Identification Based on UAV Telemetry Radios / Y. Tian et al. Sensors. 2024. Vol. 24, no. 16. P. 5099.
3. Zhan, Y., et al. (2021). "UAV Detection Based on RF Fingerprints with CNN-LSTM Hybrid Model." Sensors.
4. Shi, Y., et al. (2020). "Deep Learning for RF Fingerprinting: A ResNet-Based Approach." IEEE Transactions on Vehicular Technology.
5. Tzeng, E., et al. (2017). "Adversarial Discriminative Domain Adaptation." CVPR.

DOI 10.34132/mspc2025.01.14.19

Ivan Sova,
Oleksiy Kozlov,

ADAPTIVE MACHINE LEARNING METHODS FOR UAV RECOGNITION IN A VARIABLE RADIO FREQUENCY ENVIRONMENT

The paper addresses the challenge of recognizing unmanned aerial vehicles (UAVs) based on radio frequency signals in variable operating conditions. It highlights that conventional machine learning models, trained in static environments, often suffer performance degradation when exposed to changes in channel characteristics, antenna positioning, noise, and multipath effects. The need for adaptive approaches is emphasized, particularly through domain adaptation techniques, which help preserve recognition accuracy without requiring full retraining. An architecture is proposed that incorporates a feature extraction module, a domain discriminator, and a simulated environment generator. This approach offers promising prospects for the development of robust UAV monitoring systems in both military and civilian applications.

Key words: UAV, recognition, radio frequency analysis, domain adaptation, machine learning, dynamic environment, RF fingerprinting.

Відомості про авторів / Information about the Authors

Іван Сова, аспірант факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: owlvano@gmail.com, orcid: <https://orcid.org/0000-0002-8972-4069>.

Ivan Sova, Ph. D. student of the Computer Science Faculty of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: owlvano@gmail.com, orcid: <https://orcid.org/0000-0002-8972-4069>.

Олексій Козлов, доктор технічних наук, професор, професор кафедри інтелектуальних інформаційних систем Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: kozlov_ov@ukr.net, orcid: <https://orcid.org/0000-0003-2069-5578>.

Oleksiy Kozlov, Doctor of Technical Sciences, Professor, Professor of the Department of Intellectual Information Systems of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: kozlov_ov@ukr.net, orcid: <https://orcid.org/0000-0003-2069-5578>.

DOI 10.34132/mspc2025.01.14.20

Богдан Сомряков
Світлана Боровльова

КЛАСИФІКАЦІЯ ВІЙСЬКОВОЇ ТЕХНІКИ ЗА ДОПОМОГОЮ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ (CNN)

У тезі представлено модель згорткової нейронної мережі (CNN) для автоматичної класифікації військової техніки на основі зображень. Дослідження базується на датасеті з понад 14 000 зображень, поділених на тренувальну та тестову вибірки, що охоплюють 10 класів, зокрема танки, артилерію, БМП та зенітні комплекси.

Результати підтверджують ефективність застосування глибинного навчання для розпізнавання військових об'єктів. Запропонований підхід має потенціал для використання у військовій аналітиці, дронівих системах спостереження, автоматичному виявленні цілей та в дослідженнях оборонного спрямування.

Ключові слова: Згорткова нейронна мережа (CNN), Класифікація військової техніки, Розпізнавання зображень, Глибинне навчання, Обробка зображень, Машинне навчання, Комп'ютерний зір.

У сучасних умовах ведення бойових дій автоматизовані системи розпізнавання військової техніки відіграють важливу роль у забезпеченні безпеки та ефективного прийняття рішень. Застосування згорткових нейронних мереж (CNN) у задачах класифікації зображень дозволяє досягти високої точності і швидкості обробки даних, що робить цю технологію перспективною для використання у військовій сфері.

У представленому дослідженні використано датасет [1], який складається із 10414 зображень, що належать до 10 різних класів військової техніки для навчальної вибірки (train), та 3719 зображень для тестової вибірки (test) (1). Кожне зображення представляє один із класів бойових машин, зокрема: танки, зенітні установки (Anti-aircraft), бойові машини піхоти (Infantry Fighting Vehicles), артилерія та інші категорії.

$$D = D_{train} \cup D_{test}, \quad |D_{train}| = \alpha |D|, \quad |D_{test}| = (1 - \alpha) |D| \quad (1)$$

Перед обробкою всі зображення були перетворені у відтінки сірого (grayscale) (2) для зменшення розмірності вхідних даних, після чого мітки класів були закодовані у вигляді бінарних векторів (one-hot encoding) (3).

$$Gray = [0.2989 \ 0.5870 \ 0.1140] * \begin{bmatrix} R \\ G \\ B \end{bmatrix} \quad (2)$$

$$v_i = \delta(y, c_i) = \begin{cases} 1, & \text{if } y = c_i \\ 0, & \text{else} \end{cases} \text{ for } i = 1, 2, 3 \dots, K \quad (3)$$

Для наочності на рисунку 1 наведено приклади зображень із навчальної вибірки, які ілюструють різні класи військової техніки з досліджуваного датасету, де елементи датасету представлені категоріальними мітками у вигляді векторів та зображення переведені у відтінки сірого.



зображення у відтінки сірого

нейронів.

model	=	Sequential([
Conv2D(32, (3, 3), activation='relu', input_shape=(256, 256, 3)),		
BatchNormalization(),		
MaxPooling2D(2, 2),		
Dropout(0.2),		
Conv2D(64, (3, 3), activation='relu'),		
BatchNormalization(),		
MaxPooling2D(2, 2),		
Dropout(0.2),		
Conv2D(128, (3, 3), activation='relu'),		
BatchNormalization(),		
MaxPooling2D(2, 2),		
Dropout(0.2),		
Conv2D(128, (3, 3), activation='relu'),		
BatchNormalization(),		

```
MaxPooling2D(2, 2),
Dropout(0.2),
Flatten(),
Dense(512, activation='relu'),
Dropout(0.5),
Dense(train_generator.num_classes, activation='softmax')
])
```

Для багатокласової класифікації на виході моделі застосовано функцію активації Softmax (5). Функція перетворює вихідні значення мережі на ймовірності для кожного класу, причому сума всіх ймовірностей дорівнює 1. Це дозволяє моделі ефективно класифікувати зображення в один з декількох можливих класів, на відміну від сигмоїдної (4) функції активації, яка зазвичай використовується для двокласових задач і повертає ймовірність для одного класу, що не є оптимальним для багатокласових сценаріїв.

$$\sigma(x) = \frac{1}{1+e^{-x}} \quad (4)$$

$$P(y_i) = \sum_{j=1}^N \frac{e^{z_i}}{e^{z_j}} \quad (5)$$

У даному дослідженні завдання класифікації полягає у розподілі об'єктів на 10 класів, що вказує на застосування багатокласового підходу замість бінарного. Це означає, що замість бінарної крос-ентропії (binary crossentropy) (6), яка застосовується при наявності лише двох класів, використовується категоріальна крос-ентропія (categorical crossentropy) (7). Такий підхід застосовується, коли задача полягає в класифікації з більше ніж двома класами, такими як танки, артилерія, ППО тощо.

$$-\frac{1}{n} \sum_{i=1}^n y_i * \log(\hat{y}_i) + (1 - y_i) * \log(1 - \hat{y}_i) \quad (6)$$

$$-\sum_{i=1}^C y_i * \log(p_i) \quad (7)$$

Рисунок 2 ілюструє розподіл ймовірностей для зображення танка, яке було завантажено в модель. Нейронна мережа успішно класифікувала зображення як танк, надавши найвищу ймовірність [2] для цього класу.

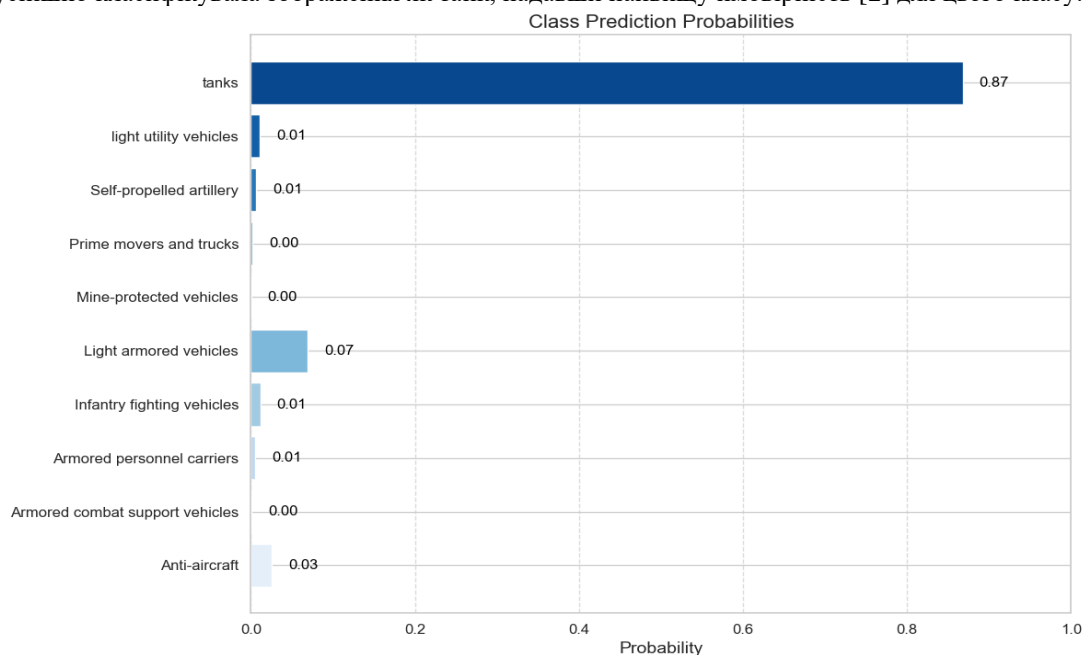


Рис. 2. – Розподіл ймовірностей для зображення танка

Результати тренування моделі показують, що на тренувальних даних досягнуто точності 0.9993 з втратою 0.0192, що вказує на надзвичайно високі результати на цих даних. Однак на тестових даних точність склала лише 0.5329, а втрата – 2.2691. Хоча результат 0.53 в тестуванні є в 5 разів кращим за випадкове вгадування, це також вказує на те, що модель може бути перевчена (overfitted) [3], оскільки вона демонструє значно кращі результати на тренувальних даних, ніж на тестових.

Розроблену модель можна застосовувати для автоматичної ідентифікації типів військової техніки на зображеннях, що може бути корисним у військовій аналітиці, моніторингу з дронів, системах безпеки або для дослідницьких цілей в оборонній сфері.

References

1. Military vehicles. URL: <https://www.kaggle.com/datasets/amanrajbose/military-vechiles>. (дата звернення 19.05.2025).
2. Probability. URL: <https://www.cuemath.com/data/probability/> (дата звернення 19.05.2025).
3. What is Overfitting? URL: <https://aws.amazon.com/what-is/overfitting/#:~:text=Overfitting%20is%20an%20undesirable%20machine,on%20a%20known%20data%20set>. (дата звернення 19.05.2025).
4. Military Applications of Machine Learning: A Bibliometric Perspective. URL: https://www.researchgate.net/publication/360162218_Military_Applications_of_Machine_Learning_A_Bibliometric_Perspective. (дата звернення 19.05.2025).

DOI 10.34132/mspc2025.01.14.20

Bohdan Somriakov
Svitlana Borovlova

CLASSIFICATION OF MILITARY VEHICLES USING CONVOLUTIONAL NEURAL NETWORKS (CNNs)

The thesis presents a convolutional neural network (CNN) model for the automatic classification of military vehicles based on images. The study is based on a dataset of over 14,000 images, divided into training and testing sets, covering 10 categories including tanks, artillery, infantry fighting vehicles, and anti-aircraft systems.

The results confirm the effectiveness of deep learning in recognizing military objects. The proposed approach has potential applications in military analytics, drone-based surveillance systems, automated target detection, and defense-related research.

Keywords: Convolutional Neural Network (CNN), Military Vehicles Classification, Image Recognition, Deep Learning, Image Processing, Machine Learning, Computer Vision.

Відомості про авторів / Information about the Authors

Богдан Сомряков, студент факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, Миколаїв, Україна. E-mail: pontius.bos@gmail.com, orcid: <https://orcid.org/0009-0003-7045-8586>

Bohdan Somriakov, student of the Faculty of Computer Sciences, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine E-mail: pontius.bos@gmail.com, ORCID: <https://orcid.org/0009-0003-7045-8586>

Світлана Боровльова, ст. викладачка кафедри інженерії програмного забезпечення Чорноморського національного університету імені Петра Могили, Миколаїв, Україна. E-mail: svetlana.borovlyova@chmnu.edu.ua, orcid: <https://orcid.org/0000-0003-1994-0556>.

Svitlana Borovlova, Senior Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: svetlana.borovlyova@chmnu.edu.ua, ORCID: <https://orcid.org/0000-0003-1994-0556>

DOI 10.34132/mspc2025.01.14.21

Євгеній Стоєв
Віктор Раленко

АРХІТЕКТУРНІ ПРИНЦИПИ ПОБУДОВИ ВІДМОВОСТІЙКИХ І ВИСОКОПРОДУКТИВНИХ ОБЧИСЛЮВАЛЬНИХ СИСТЕМ РЕАЛЬНОГО ЧАСУ: ВПЛИВ ПАТЕРНІВ ПРОЄКТУВАННЯ НА ФУНКЦІОНАЛЬНІ ХАРАКТЕРИСТИКИ

У дослідженні розглядаються архітектурні принципи побудови відмовостійких та високопродуктивних обчислювальних систем реального часу. Основна увага приділяється використанню патернів проєктування, таких як *Circuit Breaker*, *CQRS* та *Saga Pattern*, що сприяють підвищенню стабільності, продуктивності та узгодженості даних у розподілених системах. Оптимізація продуктивності досягається шляхом асинхронної обробки даних, подієво-орієнтованих підходів та використання платформ, таких як *Apache Kafka* і *Spark Streaming*. Додатково розглядається важливість автоматичного масштабування та ефективного кешування для забезпечення стабільної роботи високонавантажених систем.

Ключові слова: Відмовостійкі системи, високопродуктивні обчислювальні системи, патерни проєктування, *Circuit Breaker*, *CQRS*, *Saga Pattern*, стабільність, продуктивність, узгодженість даних, асинхронна обробка даних, подієво-орієнтовані підходи, автоматичне масштабування, балансування навантаження, горизонтальне масштабування, доступність.

У сучасному світі технологічний розвиток зумовлює необхідність створення обчислювальних систем, здатних працювати в режимі реального часу, гарантуючи при цьому високу продуктивність та відмовостійкість. Такі системи критично важливі для фінансових операцій, телекомунікацій, автоматизованого виробництва, транспорту та інших галузей, де швидкість обробки даних безпосередньо впливає на ефективність бізнес-процесів.

Архітектурні принципи проєктування відіграють ключову роль у формуванні надійних та масштабованих інфраструктур для обробки потоків даних реального часу. Одним із найважливіших аспектів є використання патернів проєктування, які допомагають структурувати систему таким чином, щоб вона могла ефективно обробляти запити, управляти транзакціями та гарантувати узгодженість даних.

Патерни проєктування в системах реального часу виконують низку функцій:

- Забезпечують відмовостійкість шляхом використання механізмів автоматичного відновлення та балансування навантаження (*Circuit Breaker*, *Failover Strategies*);
- підвищують продуктивність за допомогою структурованого розділення операцій читання та запису (*CQRS*), а також подієво-орієнтованого підходу (*Event-Driven Architecture*);
- оптимізують управління транзакціями в розподілених системах за допомогою *Saga Pattern*, що дозволяє гарантувати коректне виконання складних бізнес-операцій;
- спрощують обробку поточкових даних через використання *Data Streaming Patterns* та спеціалізованих рішень на базі *Kafka*, *Flink* або *Spark Streaming*.

Таким чином, архітектурні патерни не лише підвищують стабільність роботи обчислювальних систем реального часу, але й дозволяють ефективно масштабувати їх, забезпечуючи максимальну продуктивність навіть при значних навантаженнях.

Обчислювальні системи реального часу мають низку критичних вимог, що визначають їхню ефективність та стабільність у високонавантажених середовищах. Однією з найважливіших характеристик є низька латентність, оскільки такі системи повинні оперативно реагувати на події та обробляти великі обсяги даних без затримок. У багатьох випадках необхідно дотримуватися жорстких часових меж виконання операцій, що особливо актуально для фінансових платформ, промислових автоматизованих процесів або медичних систем моніторингу.

Висока доступність є ще одним ключовим фактором, оскільки критично важливі сервіси повинні працювати безперервно, навіть за умов часткових збоїв або відмови вузлів. Для цього застосовуються механізми резервного копіювання, автоматичного переключення на резервні сервери та шаблони проектування, такі як *Circuit Breaker*, які запобігають каскадним відмовам.

Масштабованість також відіграє важливу роль, оскільки системи реального часу повинні адаптуватися до динамічних змін навантаження. Горизонтальне масштабування дозволяє ефективно розподіляти ресурси та підтримувати обробку потоків даних на великих обсягах, що особливо важливо для телекомунікаційних платформ або фінансових ринків.

Узгодженість даних у розподілених обчислювальних середовищах потребує вибору правильного балансу між *Strong Consistency* та *Eventual Consistency*. Для складних бізнес-операцій використовується *Saga Pattern*, який забезпечує коректне управління довготривалими транзакціями у розподілених системах.

Крім того, захист і безпека інформації є важливими вимогами, оскільки системи реального часу обробляють конфіденційні та стратегічно важливі дані. Використання шифрування, механізмів автентифікації та інструментів для моніторингу активності допомагає запобігати загрозам та гарантує цілісність інформації.

Забезпечення всіх цих аспектів дозволяє системам реального часу ефективно працювати в умовах підвищеного навантаження, мінімізуючи ризики збоїв і забезпечуючи стабільну роботу в критично важливих галузях.

Продуктивність обчислювальних систем реального часу визначається їхньою здатністю обробляти великі обсяги даних з мінімальною затримкою. В умовах високонавантажених середовищ критично важливо досягти оптимального балансу між швидкістю виконання запитів, використанням ресурсів та узгодженістю даних.

Одним із ключових методів оптимізації є асинхронна обробка даних, яка дозволяє мінімізувати блокуючі виклики та збільшити загальну швидкість системи. Застосування подієво-орієнтованої архітектури (*Event-Driven Architecture, EDA*) дозволяє реагувати на зміни у даних без необхідності чекати завершення всіх процесів, що значно пришвидшує роботу у критичних сценаріях.

Розділення операцій читання та запису за допомогою *CQRS (Command Query Responsibility Segregation)* також сприяє оптимізації продуктивності. Цей підхід дає змогу розділити бізнес-логіку на незалежні компоненти, що дозволяє масштабувати запити та підвищити ефективність обробки даних. *CQRS* широко використовується у фінансових системах та інших середовищах, де важливо миттєво отримувати результати запитів без затримок у транзакціях.

Для підвищення продуктивності також застосовується обробка поточкових даних (*Data Streaming Patterns*). Використання платформ на кшталт *Apache Kafka*, *Flink* або *Spark Streaming* дозволяє працювати з постійними потоками інформації, що надходять від датчиків, користувачів або зовнішніх API, забезпечуючи швидку реакцію та адаптивність до змін.

Ще одним важливим фактором є автоматичне масштабування, яке дозволяє розподіляти навантаження між серверами залежно від поточної активності. Це досягається за допомогою оркестрації контейнерів у *Kubernetes*, який автоматично додає чи видаляє ресурси відповідно до запитів у режимі реального часу.

Крім того, оптимізація продуктивності включає ефективне кешування та застосування технологій попередньої обробки запитів. Використання *Redis* або *Memcached* допомагає значно знизити навантаження на основні бази даних та забезпечити швидкий доступ до критичних даних.

Побудова відмовостійких та високопродуктивних обчислювальних систем реального часу вимагає комплексного підходу, що включає сучасні архітектурні патерни, ефективне управління ресурсами та оптимізацію продуктивності. Таким чином, ефективна архітектура систем реального часу базується на принципах динамічного управління даними, відмовостійких механізмах та адаптивних підходах до продуктивності. Використання перевірених патернів проектування дає змогу побудувати стабільну, масштабовану та ефективну інфраструктуру для критично важливих бізнес-процесів.

References

1. Top 10 integration patterns for enterprise use cases. MuleSoft URL: <https://blogs.mulesoft.com/api-integration/patterns/top-10-integration-patterns/> (date accessed: 28.04.2025)
2. Understanding Enterprise Integration Patterns. HALSIMPLIFY. URL: <https://www.halsimplify.com/knowledge-center/enterprise-integration-patterns-guide> (date accessed: 28.04.2025)
3. A comprehensive guide to integration patterns in modern business systems. LONTI. URL: <https://www.lonti.com/blog/a-comprehensive-guide-to-integration-patterns-in-modern-business-systems> (date accessed: 30.04.2025)

ARCHITECTURAL PRINCIPLES FOR BUILDING FAULT-TOLERANT AND HIGH-PERFORMANCE REAL-TIME COMPUTING SYSTEMS: THE IMPACT OF DESIGN PATTERNS ON FUNCTIONAL CHARACTERISTICS

This study examines the architectural principles for building fault-tolerant and high-performance real-time computing systems. The focus is on the use of design patterns such as Circuit Breaker, CQRS, and Saga Pattern, which contribute to the improvement of stability, performance, and data consistency in distributed systems. Performance optimization is achieved through asynchronous data processing, event-driven approaches, and the use of platforms such as Apache Kafka and Spark Streaming. Additionally, the importance of auto-scaling and effective caching is considered to ensure the stable operation of high-load systems.

Keywords: Fault-tolerant systems, high-performance computing systems, design patterns, Circuit Breaker, CQRS, Saga Pattern, stability, performance, data consistency, asynchronous data processing, event-driven approaches, auto-scaling, load balancing, horizontal scaling, availability.

Відомості про авторів / Information about the Authors

Євгеній Стоєв, викладач кафедри інженерії програмного забезпечення, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: stoiev.yevhenii@chmnu.edu.ua, orcid: <https://orcid.org/0009-0008-8999-2267>.

Yevhenii Stoiev, Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: stoiev.yevhenii@chmnu.edu.ua, orcid: <https://orcid.org/0009-0008-8999-2267>.

Віктор Раленко, викладач кафедри інженерії програмного забезпечення, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

Viktor Ralenko, Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

DOI 10.34132/mspc2025.01.14.22

Микола Фісун,
Ігор Кандиба

СТВОРЕННЯ СЛОВНИКА ПРЕДМЕТНОЇ ГАЛУЗІ ЗАСОБАМИ PYTHON

Тези представляють аналіз можливості побудови словника предметної галузі засобами мови програмування Python. Представлено різні структури даних для зберігання сформованих словників предметної галузі та засоби Python для роботи з цими структурами. Описано варіант застосування моделі з підтримкою SPARQL-запитів. Визначено програмні засоби для аналізу тексту та виділення знайдених термінів. Коротко наведено опис запропонованих інструментів для пошуку термінів. Наведено опис процесу пошуку термінів з використанням засобів обробки природної мови. Запропоновано варіант застосування пошуку шаблонних конструкцій для визначення значень знайдених термінів. Описано можливість застосування регулярних виразів для пошуку та обробки шаблонних конструкцій. Розглянуто варіант використання великих мовних моделей для створення словнику предметної галузі.

Ключові слова: словник предметної галузі, RDF, NLTK, Python.

Словниками предметної галузі називають структуровані зібрання термінів, понять і визначень, що відображають специфічну лексику певної професійної або наукової галузі. А в сучасних інформаційних системах такі словники використовуються навіть у якості основи для побудови онтологій та семантичних моделей предметної галузі [1].

Словники, зазвичай, розробляються фахівцями предметної галузі, при цьому рівень автоматизації таких процесів залишається ще незадовільним. Тому питання аналізу, вибору та подальшого розвитку інструментів обробки текстової інформації, що подається природною мовою, залишаються актуальними.

У сучасних інформаційних технологіях для зберігання та обробки словників предметної галузі можуть використовуватися різні формати даних. Одними з найпоширеніших є формат електронних таблиць – Comma-Separated Values (CSV) [2]. Завдяки своїй простоті, цей формат зручний для представлення словників предметної галузі у вигляді таблиці з колонками "Термін", "Визначення", "Категорія", "Синоніми" тощо, що дозволяє легко обробляти такі дані за допомогою мов програмування загального призначення: C#, Java, Python, тощо.

Варто згадати про можливість використання формату JavaScript Object Notation (JSON) для представлення словника предметної галузі. Цей формат має складнішу структуру ніж CSV, але є більш гнучким і дозволяє створювати словники з додатковими властивостями. JSON дозволяє зберігати терміни разом з їхніми визначеннями, синонімами, категоріями та іншими атрибутами в зручній, для подальшої обробки структурі даних. Згаданий формат часто використовується при розробці вебзастосунків та при виконанні семантичного аналізу.

Найчастіше для зберігання та обробки словників предметної галузі використовується Resource Description Framework (RDF). Цей формат використовується для створення онтологій – формалізованих словників з логічними зв'язками. На відміну від описаних вище формат RDF призначений для моделювання даних, розроблений W3C для опису інформації у вигляді трійок "суб'єкт–предикат–об'єкт", що дозволяє формалізовано представляти знання та зв'язки між об'єктами. У контексті словників предметної галузі RDF використовується для опису термінів [2], їхніх визначень, синонімів, ієрархічних або семантичних зв'язків у вигляді графів. Варто зазначити, що на основі RDF частіше створюють онтології. Онтології є формалізованими структурами, які описують знання про предметну область у вигляді понять (класів), зв'язків між ними (відношень) та властивостей. Ці сутності дозволяють не лише визначити певні терміни предметної галузі, а й відобразити їх зв'язок.

Мова програмування Python підтримує можливість роботи з усіма перерахованими типами даних. Проте, найчастіше для роботи з словниками предметних галузей використовують RDF. Цей факт обумовлений використанням RDF у семантичному веб та в ПЗ, побудованому на використанні онтологіями. Інтеграція RDF-

словники з Python виконується за допомогою бібліотек, що підтримують обробку запитів до даних в форматі RDF, найрозповсюдженіша з яких є RDFlib. Ця бібліотека дозволяє створювати RDF-графи, додавати трійки, зчитувати RDF-файли в різних форматах (Turtle, RDF/XML, N-Triples, тощо) і виконувати SPARQL-запити. Для створення словника предметної галузі необхідно створення відповідного простору імен:

```
from rdflib import Graph, Literal, RDF, URIRef, Namespace
```

```
g = Graph()
```

```
EX = Namespace("https://chmnu.edu/")
```

```
g.add((EX["Python"], RDF.type, EX["Programming_language"]))
```

```
g.add((EX["Python"], EX["Supports_paradigm"], Literal("Object-oriented_programming ")))
```

```
g.add((EX["Python"], EX["Supports_paradigm"], Literal("structured")))
```

```
g.add((EX["Python"], EX["Supports_paradigm"], Literal("structured")))
```

```
g.add((EX["Python"], EX["Has_a_translator_type"], EX["interpreter"]))
```

```
g.serialize("python_example.rdf", format="xml")
```

```
for subj, pred, obj in g:
```

```
    print(f"{subj} -- {pred} --> {obj}")
```

Підтримка запитів мовою SPARQL є невід'ємною частиною ПЗ для створення словника предметної галу на основі RDF. SPARQL (SPARQL Protocol and RDF Query Language) являє собою мову запитів для вибірки та обробки даних, збережених у форматі RDF. Ця мова фактично є аналогом SQL у реляційних базах даних, SPARQL дозволяє звертатися до RDF-графів, шукати трійки, фільтрувати дані, обирати змінні, виконувати агрегацію та комбінувати результати:

```
qres = g.query("""
```

```
    SELECT ?Programming_language ?Supports_the_paradigm
```

```
    WHERE {
```

```
        ?Programming_language <https://chmnu.edu/Supports_paradigm> ?Supports_the_paradigm .
```

```
    }
```

```
""")
```

```
for row in qres:
```

```
    print(f"Programming          language:          {row.Programming_language},          Supports_the_paradigm: {row.Supports_the_paradigm}")
```

Варто зауважити, що повністю автоматизувати процес створення словника предметної галузі на основі аналізу тексту практично не можливо, але виділити основні терміни можливо. Засобами Python можлива часткова автоматизація процесу створення словника предметної галузі. Для цього необхідно використати додаткові бібліотеки для обробки природної мови: NLTK, Rake, Spacy тощо. Обробка природної мови (Natural Language Processing, NLP) засобами Python є вагомою галуззю штучного інтелекту, що полягає в реалізації аналізу, інтерпретації та генерувати тексти природною мовою за допомогою програмного забезпечення.

Одним з інструментів для автоматизованого створення словнику предметної галузі є бібліотека rake, що є реалізацією однойменного алгоритму. RAKE (Rapid Automatic Keyword Extraction) є алгоритмом для автоматичного виявлення ключових слів і фраз із тексту без потреби в попередньому навчанні чи складному лінгвістичному аналізі. Rake працює на основі статистики, розділяючи текст на фрази за «стоп-словами», оцінюючи значущість кожного слова, а потім обчислюючи вагу фраз як суми ваг її складових.

Хоча використання бібліотеки rake дасть змогу виділи основні ключові слова, але для формування словника предметної галузі необхідно також виділити їх визначення. Для реалізації цього процесу необхідно використання інструментарію NLTK. NLTK (Natural Language Toolkit) є однією з найпопулярніших бібліотек Python для обробки природної мови, що реалізує широкий функціонал для: токенизації, стемінгу, лематизації, визначення частини мови, синтаксичного розбору тощо [4].

Процес створення словника предметної галузі засобами NLTK необхідно виконати виділення термів. Для цього потрібно виконати кілька дій: видалити сполучники та розділові знаки методом stopwords.words(), провести токенизацію word_tokenize(), розпізнати частину мови pos_tag() та виділити лише іменники startswith('NN'), де «NN» – позначення іменників в однині. Такий підхід обмежений можливістю вибору лише термів іменників, але може бути модифікований шляхом додавання нового функціоналу і включати у якості термів інші частини мови. В Python код для такої обробки тексту з файлу «text.txt» виглядатиме наступним чином:

```
nltk.download('punkt')
nltk.download('averaged_perceptron_tagger') # Для англійської POS-розмітки
nltk.download('averaged_perceptron_tagger_eng')
nltk.download('stopwords')

with open("text.txt", "r", encoding="utf-8") as file:
    text = file.read()

tokens = word_tokenize(text.lower())
filtered = [w for w in tokens if w.isalpha() and w not in stopwords.words('english')]
tagged = pos_tag(filtered)
terms = [word for word, pos in tagged if pos.startswith('NN')]
```

Проте після виділення термів неможливо створити словник предметної галузі. Для цього потрібно виділити також значення самих термів, що можливо зробити, виконавши пошук шаблонних конструкцій з визначенням цих термів: «X – це ...», «X представляє собою...», «X є...» «Під X розуміють ...», «X визначається як ...», де X – це терм, визначений на попередньому кроці. Цей процес можна автоматизувати додатковими засобами: spaCy, Stanza, AllenNLP тощо, але ці засоби не підтримують можливість роботи з українською мовою. Можна здійснити пошук цих конструкцій за допомогою регулярних виразів використавши модуль RE та функцію re.findall(), що знаходить всі відповідності регулярним виразам.

Інший варіант формування словника предметної галузі – це застосування засобів ШІ, наприклад GPT-моделі [5]. Для цього можливо використати бібліотеку OpenAI, що дозволяє використання різних моделей. Перевагою цього методу є те, що модель може виконувати пошук по наявному тексту або знайти визначення певних термів самостійно використовуючи інформацію з мережі Інтернет.

Використання мовної моделі-GPT дозволяє автоматизувати завдання обробки природної мови. За допомогою бібліотек, таких як OpenAI, розробники можуть інтегрувати доступ до моделей GPT у свої застосунки, надсилаючи текстові запити та отримуючи текстову відповідь. Це відкриває широкі можливості для створення інтелектуальних пошукових систем, асистентів для написання коду чи текстів, а також інструментів для аналізу великих обсягів текстової інформації. GPT-моделі можна адаптувати під конкретні потреби, використовуючи додаткові параметри запиту або навчання на прикладах.

References

1. Veselovskaya, G. V., & Lebed, O. S. (2018). Models of the subject field in computer multimedia projection technologies. *Applied Questions of Mathematical Modeling*, 2, 203–211. <https://doi.org/10.32782/2618-0340-2018-2-203-211>
2. Nayak, A., Božić, B., & Longo, L. (2022). Data Quality Assessment of Comma Separated Values Using Linked Data Approach. *Y Business Information Systems Workshops* (p. 240–250). Springer International Publishing
3. Dongo, I., Cardinale, Y., & Chbeir, R. (2018). RDF-F: RDF Datatype inFerring Framework. *Data Science and Engineering*, 3(2), 115–135.
4. Nelli, F. (2018). Textual Data Analysis with NLTK. *Y Python Data Analytics* (p. 487–506). Apress.
5. Cao, Z., & Ren, Q. (2024). Towards python program repair with generative pre-trained transformer (GPT-3.5). *Theoretical and Natural Science*, 43(1), 102–112.

DOI 10.34132/mspc2025.01.14.22

**Mykola Fisun,
Ihor Kandyba**

PYTHON-BASED CREATION OF A DOMAIN DICTIONARY

This abstract presents an analysis of the feasibility of creating a domain-specific dictionary using the Python programming language. Various data structures for storing the generated domain-specific dictionaries and Python tools for working with these structures are presented. A use case for a model with SPARQL query support is described. Software tools for text analysis and the extraction of identified terms are defined. A brief description of the proposed tools for term

discovery is provided. The process of term discovery using natural language processing tools is described. An approach involving the search for pattern constructions (template matching) to determine the meanings of the identified terms is proposed. The applicability of regular expressions for searching and processing these pattern constructions is described. The option of using large language models for creating the domain-specific dictionary is considered.

Key words: domain vocabulary, RDF, NLTK, Python.

Відомості про автора / Information about the Author

Микола Фісун, д-р техн. наук, професор кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ntfis@chmnu.edu.ua, orcid: <https://orcid.org/0000-0003-1297-6230>.

Mykola Fisun, Doctor of Technical Sciences, Professor of the Department of Software Engineering Faculty of Computer Science the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ntfis@chmnu.edu.ua, orcid: <https://orcid.org/0000-0003-1297-6230>.

Ігор Кандиба PhD, старший викладач кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: igor.kandyba@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-8589-4028>

Ihor Kandyba, PhD, Senior Lecturer of Software Engineering Faculty of Computer Science the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ntfis@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-8589-4028>

DOI 10.34132/mspc2025.01.14.23

Альона Швед,
Євген Давиденко

ЗАСТОСУВАННЯ CBR-ПІДХОДУ В СИСТЕМАХ СИТУАЦІЙНОГО УПРАВЛІННЯ

У роботі досліджено проблематику ситуаційного управління та моделювання, розглянуто сучасні методи та інструменти реалізації ситуаційного менеджменту. Досліджено можливість та результативність застосування CBR-підходу при пошуку рішень нетипових, унікальних задач особливо в умовах обмеженості наявних ресурсів та невизначеності оточуючого середовища. Розглянуто ряд сучасних моделей подання і виведення знань в системах ситуаційного управління. Проаналізовано методи подання та вилучення знань на основі прецедентів.

Ключові слова: системи ситуаційного управління, прецедент, база знань, онтологія, метод міркувань на основі прецедентів

У сучасних умовах управління процесами, що протікають у сучасних соціальних, економічних, організаційних, технічних та інших системах стає все більш складним завданням, на вирішення якого безпосередньо впливає неминуче зростання складності таких систем, збільшення обсягів оброблюваної інформації, складна передбачуваність та посилення динамічності факторів зовнішнього середовища під впливом яких вони функціонують. В таких умовах застосування класичних моделей та принципів управління не сприяє підвищенню ефективності функціонування такої системи, що висуває задачу пошуку більш гнучких та ефективних механізмів та методів управління. Одним із таких підходів є ситуаційне управління, яке являє собою доповнення до стратегічного та оперативного менеджменту та полягає в пошуку шляхів вирішення поточної проблемної ситуації в реальному часі, спираючись на суб'єктивні і евристичні знання фахівців предметної області.

Основним завданням ситуаційного менеджменту є перенесення коригувальних дій у напрямі розвитку підприємства та розподілу ресурсів, які забезпечують реалізацію стратегічних цілей [2]. Важливо відзначити, що для вирішення проблемної ситуації створюються команди різних спеціалістів, які за результатами проведеного аналізу поточної проблемної ситуації шляхом командної роботи мають визначити найбільш прийнятний (ефективний, цінний) для підприємства з токи зору наявних ресурсів, часу реалізації та очікуваної вигоди варіант її вирішення.

Розглянемо формалізоване подання ситуаційного управління [1]. Під поточною ситуацією S розуміється сукупність поточного стану об'єктів (вектор станів X) та його зовнішньої сфери (вектор збурень F). Тоді поточна ситуація може бути описана кортежем виду $C = \langle X, F \rangle$.

Як «інструмент» реалізації принципів та підходів ситуаційного управління можуть бути використані інтелектуальні системи підтримки прийняття рішень реального часу (ІСППР РЧ), складовою частиною яких є бази знань (БЗ). В основу створення БЗ покладені експертні знання про предметні області ситуацій, за якими потрібно оперативно приймати рішення, та їх подання у вигляді оптимальної моделі. Для формування БЗ та організації швидкого доступу до них останнім часом широкого розповсюдження набув метод міркувань на основі прецедентів. Міркування на основі прецедентів ґрунтуються на накопиченні досвіду і подальшій адаптації рішення вже відомої задачі до вирішення нової. В наукових публікаціях *прецедент* (від лат. *praecedentis* – попередній) визначається як випадок (ситуація), що мав місце раніше і який є прикладом для наступних ситуацій подібного роду [5, 8–10].

Метод міркувань на основі прецедентів (*CBR – Case-Based Reasoning*) успішно використовується в різних областях людської діяльності: в системах експертного діагностування, системах підтримки прийняття рішень, системах машинного навчання, при вирішенні задач прогнозування, узагальнення накопиченого досвіду, пошуку рішень в маловивчених предметних областях [3, 7, 11–13].

Виведення на основі прецедентів є підходом, що дозволяє вирішувати нову, невідому задачу, використовуючи або адаптуючи вже накопичений досвід вирішення подібних задач. Цей метод з'явився внаслідок того, що велика кількість практичних задач є погано формалізованими, причому невизначеність має неймовірностний характер.

При пошуку рішень подібних задач необхідно застосовувати методи правдоподібного виведення, що дозволяють знайти прийнятне (яке може і не бути оптимальним) рішення. Одним з підходів до цього є факт того, що людині (експерту; аналітику; особі, що приймає рішення) притаманно на першому етапі пошуку рішення нової (невідомої) задачі намагатися використовувати рішення, які приймалися раніше в подібних ситуаціях, і при необхідності адаптувати їх до поточної проблеми (поточної проблемної ситуації).

Зазвичай, процес виведення на основі прецедентів складається із чотири послідовних етапів, які утворюють так званий цикл міркувань на основі прецедентів або CBR-цикл [5, 8–10].

Основними етапами CBR-циклу є:

- вилучення (*retrieve*) найбільш подібного до поточної ситуації прецеденту із бібліотеки прецедентів (БПр);
- повторне використання (*reuse*) вилученого прецеденту (за можливістю застосування існуючого в БПр рішення поточної проблеми);
- перегляд та адаптація (*revise*) – у випадку неможливості повторного використання прецеденту, відбувається вироблення нового рішення, шляхом внесення змін до вже існуючого рішення для обраного прецеденту для вирішення нової задачі (поточної проблемної ситуації);
- збереження (запам'ятовування) (*retain*) опису нової задачі та прийнятого рішення як частини нового прецеденту.

У загальному вигляді модель подання прецеденту містить опис ситуації і рішення для даної ситуації [5, 8–10]:

$$CASE = (Sit, Sol, Res), \quad (1)$$

де *Sit* – ситуація, що описує прецедент (параметри вихідної задачі); *Sol* – рішення що пропонується (рекомендації особи, що приймає рішення (ОПР); *Res* – результат застосування рішення, який може включати список виконаних дій, додаткові коментарі та посилання на інші прецеденти, а також в деяких випадках може приводитися обґрунтування вибору даного рішення і можливі альтернативи.

Здебільшого для опису прецедентів достатньо параметричного подання у вигляді набору параметрів з конкретними значеннями і рішеннями [5, 8-10]:

$$CASE = (x_1, x_2, \dots, x_n, Sol), \quad (2)$$

де $x_1, x_2, \dots, x_n$ – параметри ситуації, яку описує прецедент ($x_1 \in X_1, x_2 \in X_2, \dots, x_n \in X_n$), n – кількість параметрів, *Sol* – рішення і рекомендації ОПР, $X_1, X_2, \dots, X_n$ – області допустимих значень відповідних параметрів.

Проте, в ряді випадків при параметричному представленні важко враховувати залежність між параметрами прецедента (наприклад, причинно-наслідкові залежності). Одним із способів вирішення такої проблеми є подання прецедентів на основі методології онтологій предметних областей.

Для формалізованого подання онтологій широкого розповсюдження набули фрейми, семантичні мережі, продукційні моделі, та ін. Такий підхід дозволяє будувати гібридні моделі подання прецедентів.

Параметрична модель подання прецедентів виду (2) може бути розширена елементами продукційної моделі через введення продукційних правил виду («ЯКЩО» *умова*, «ТО» *дія*). Продукційні правила складаються із двох частин (антецедент та консеквент) та дозволяють встановлювати існуючі залежності між параметрами прецедентів і проблемної ситуації для конкретної предметної області (ПрО), а також отримати висновки щодо невідомих фактів (наприклад, встановити відсутні значення будь-яких параметрів при описі поточної ситуації тощо) [6].

Для подання знань про ПрО в наглядній і структурованій формі модель (2) може бути вдосконалена за рахунок використання семантичних мереж (СМ). В СМ структура знань про ПрО формалізується у вигляді орієнтованого графа з позначеними вершинами і дугами. Вершини позначають параметри прецеденту; дуги – дозволяють описати різні відношення між параметрами [6].

В даний час запропонована ціла низка методів вилучення прецедентів із БПр [5, 8-10], кожен з яких має свої переваги та недоліки. До найбільш поширених можна віднести метод вилучення на основі дерев рішень, метод вилучення з урахуванням їх застосовності, метод найближчого сусіда, метод вилучення прецедентів на основі знань.

В методі вилучення прецедентів на основі дерева рішень кожна вершина дерева зазначає, в якій її гілці слід продовжувати пошук рішень. Вибір гілки здійснюється на основі інформації щодо поточної проблемної ситуації. Задача полягає в знаходженні шляху до кінцевої вершини, яка відповідає одному або декільком прецедентам.

Метод вилучення прецедентів на основі знань дозволяє враховувати знання експерта (ОПР) в межах конкретної ПрО (коефіцієнти важливості параметрів, виявлення залежності тощо) при вилученні прецедентів. Метод реалізує підхід, в основі якого покладена індексація прецедентів, враховуючи важливість їх параметрів та іншої додаткової інформації про ПрО.

В роботі [4] запропонована модифікація процедури вилучення прецедентів яка використовує експертні знання в двоетапній процедурі звуження вихідної множини прецедентів, яка передбачає попередню фільтрацію

прецедентів, значення параметрів яких належать заданим околицям відповідних параметрів цільової ситуації на першому етапі, та додаткове звуження отриманої підмножини прецедентів методами теорії грубих множин на другому етапі. Запропонований алгоритм фільтрації прецедентів дозволяє відсікати частину БПр, яка не відповідає заданим граничним межах параметрів цільового прецеденту.

Процедура пошуку рішень в методі вилучення прецедентів з урахуванням їх застосовності базується не тільки на аналізі їх схожості із поточною проблемною ситуацією, але й на оцінці того наскільки якісну для бажаного результату модель вони собою являють. Метод припускає, що найбільш схожі з поточною проблемною ситуацією прецеденти є і найбільш застосовані в цій ситуації. Таким чином, на вибір вилучених прецедентів впливає можливість їх успішного застосування (адаптація) в конкретній ситуації. У ряді випадків ця проблема вирішується шляхом збереження прецедентів разом з коментарями щодо їх застосування.

В основу методу найближчого сусіда (*NN – Nearest Neighbor*) покладена процедура оцінки ступеня подібності поточної проблемної ситуації до прецедентів з БПр. Для цього на множині параметрів, що використовуються для опису прецедентів і поточної проблемної ситуації, вводиться деяка метрика. Далі відповідно до обраної метрики визначається відстань від цільової точки, що відповідає значенню параметрів поточної проблемної ситуації, до точок, що представляють опис параметрів прецедентів із БПр. Обирається точка, яка є найближчою до цільової.

Висновки. Проаналізовано основні концепти ситуаційного моделювання та управління. Досліджено питання подання та виведення знань в системах ситуаційного управління на основі методу міркувань за прецедентами та його модифікаціями. Особлива увага приділена аналізу методів вилучення прецедентів із бібліотеки прецедентів, оскільки від тривалості процедури пошуку в значній мірі залежить час відгуку системи. Розглянуті підходи, незважаючи на свої переваги, мають ряд недоліків, серед яких можна виділити наступні: збір та попередня підготовка додаткової інформації; відсутність обґрунтованого уніфікованого підходу до вибору міри близькості; суб'єктивність вибору вагових коефіцієнтів важливості параметрів прецеденту, не стійкість до аномальних даних (викидів) та інші. В цих умовах виникає необхідність пошуку шляхів вирішення актуальної проблеми спрямованої на оптимізацію та підвищення ефективності процедури пошуку та вилучення прецедентів.

References

1. Kovalenko, I. I., Shved, A. V., & Antipova, K. O. (2018) Modeli podannia ta vyvedennia znan u systemakh sytuatsiinoho upravlinnia [Models of knowledge representation and extraction in situational management systems]. Mykolaiv: Ilion. 91 p. [in Ukrainian]
2. Zadorozhna, L. M. (2017) Sytuatsiinyi menedzhment v systemi stratehichnoho i taktychnoho upravlinnia diialnistiu pidpriemstva [Situational management in the system of strategic and tactical management of the enterprise activity]. *Biznes-navihator – Business Navigator*, 3(42), (pp. 35–38) [in Ukrainian]
3. Shved, A. V. (2022) Osoblyvosti rozrobky system sytuatsiinoho upravlinnia realnoho chasu [Features of real-time situational control systems developing]. *Olviyskyi forum-2022 - Olvia Forum - 2022: Strategies of the Black Sea Region Countries in the Geopolitical Space: Proceedings of the XVI International Scientific and Practical Conference*. (Mykolaiv, Ukraine, 23–26 June. 2022) [in Ukrainian]
4. Shved, A. V., Davydenko, Ye. O., & Horban, H. V. (2024) Intellectual support of the processes of searching and extraction of precedents in Case-based reasoning approach. *Radio Electronics, Computer Science, Control*, (3), 107, (pp. 107–118). doi: 10.15588/1607-3274-2024-3-10 [in Ukrainian]
5. Aamodt, A., & Plaza, E. (1994) Case-based reasoning: fundamental issues, methodological variations and system approaches. *AI Communications*, 7(1), (pp. 39–59).
6. Hoekstra, R. J. (2009) Ontology Representation: design patterns and ontologies that make sense. IOS Press, Incorporated. 248 p. doi: 10.3233/978-1-60750-013-1-i
7. Kowalski, Z., Meler-Kapcia, M., Zieliński, S., & Drewka, M. (2005) CBR methodology application in an expert system for aided design ship's engine room automation. *Expert Systems with Applications*, 29(2), (pp. 256–263). doi: 10.1016/j.eswa.2005.03.002
8. Leake, D., Ye, X., & Crandall, D. (2021) Supporting case-based reasoning with neural networks: an illustration for case adaptation. *Combining Machine Learning and Knowledge Engineering: AAAI Spring Symposium AAAI-MAKE 2021, Palo Alto, California, USA, 22–24 March 2021: proceedings*. CEUR Workshop proceedings, 2846. Retrieved from: <https://ceur-ws.org/Vol-2846/paper1.pdf>
9. Pal, S. K., & Shiu, S. C. K. (2004) Foundation of soft case-based reasoning. New Jersey: John Wiley & Sons. 300 p.
10. Perner, P. (2019) Case-based reasoning – methods, techniques, and applications. In: Nyström I., Hernández Heredia Y., Milián Núñez V. (eds.) *Progress in Pattern Recognition, Image Analysis, Computer Vision, and Applications. CIARP 2019. Lecture Notes in Computer Science*, 11896, (pp. 16–30). doi: 10.1007/978-3-030-33904-3_2

11. Ramadhani, M., Sihombing, V., & Yanris, G. (2021) Application of the case based learning (CBR) method to diagnose conjunctivitis. *Sinkron*, 6, (pp. 176–182). doi: 10.33395/sinkron.v6i1.10908
12. Ramos-Quintana, F., Tovar-Sánchez, E., Saldarriaga-Noreña, H., Sotelo-Nava, H., Sánchez-Hernández, J.P., & Castrejón-Godínez, M-L. (2019) A CBR–AHP hybrid method to support the decision-making process in the selection of environmental management actions. *Sustainability*, 11(20): 5649. doi: 10.3390/su11205649
13. Zhong, Q.-Y., Zhang, Y.-J., Qu, X.-F., Ye, X., & Qu, Y. (2008) Research on method of CBR and its application in emergency commanding and decision-making. *Wireless Communications, Networking and Mobile Computing: 4th International Conference*, Dalian, China, 12–14 October 2008: proceedings. Institute of Electrical and Electronics Engineers (IEEE). P. 1–4. doi: 10.1109/WiCom.2008.2744

DOI 10.34132/mspc2025.01.14.23

Alona Shved
Yevhen Davydenko

APPLICATION OF THE CBR APPROACH IN SITUATION MANAGEMENT SYSTEMS

The paper explores the tasks of situational management and modeling, considers modern situational management methods and tools. The possibility and effectiveness of using the CBR approach for finding solutions to atypical, unique problems, especially in conditions of limited available resources and uncertainty of the environment, are investigated. A set of modern models of knowledge representation and extraction in situational management systems are considered. Methods of knowledge representation and extraction based on precedents have been analyzed.

Key words: situational management systems, precedent, knowledge base, ontology, CBR approach

Відомості про авторів / Information about the Authors

Альона Швед, д-рка техн. наук, професорка, професорка кафедри інженерії програмного забезпечення Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: avshved@chmnu.edu.ua, orcid: <https://orcid.org/0000-0003-4372-7472>.

Alona Shved, Doctor of Technical Sciences, Professor, Professor of the Department of Software Engineering of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: avshved@chmnu.edu.ua, orcid: <https://orcid.org/0000-0003-4372-7472>.

Євген Давиденко, канд. техн. наук, доцент, завідувач кафедри інженерії програмного забезпечення Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: davydenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-0547-3689>.

Yevhen Davydenko, PhD, Associate Professor, Head of the Department of Software Engineering of the Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: davydenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0002-0547-3689>.

DOI 10.34132/mspc2025.01.14.24

Олександр Яновський
Марина Фаленкова
Юрій Решетнік

ЕЛЕКТРОННА КОМЕРЦІЯ: СУЧАСНІ ТЕНДЕНЦІЇ, ВИКЛИКИ ТА ПРОГНОЗИ

Здійснено комплексний аналіз сучасного стану та перспектив електронної комерції (e-commerce) як драйвера цифрової економіки. Розглядається зростання e-commerce в Україні, підкреслюючи зростаючу частку онлайн-торгівлі в загальному обсязі роздрібних продажів та нові тенденції в цифрових бізнес-моделях. Особливу увагу приділено різним моделям електронної комерції, таким як B2C, B2B, C2C та C2B, та їх впливу на ринкові сегменти та поведінку споживачів.

У роботі також аналізується роль регуляторних frameworks як в Україні, так і в Європейському Союзі, у формуванні ринку електронної комерції. Окремо розглядаються впливи таких регуляцій, як PSD2 та GDPR в ЄС і їх аналоги в українському законодавстві. Крім того, досліджуються технологічні інновації, зокрема штучний інтелект (AI), доповнена реальність (AR) та сервіси BNPL (Buy Now Pay Later), що сприяють покращенню користувацького досвіду, підвищенню конверсії та зниженню ризиків у ланцюгах постачання. Здійснено прогноз подальшого розвитку електронної комерції, зокрема значним зростанням ринку та інтеграцією передових технологій у найближчі роки. Наголошено на необхідності адаптації бізнесу до змін в регуляторному, технологічному та споживчому середовищі для збереження конкурентоспроможності.

Ключові слова: електронна комерція, цифрова економіка, моделі e-commerce, регуляторні рамки, штучний інтелект, доповнена реальність, BNPL.

Електронна комерція (e-commerce) стала ключовим драйвером цифрової економіки, формуючи нові канали збуту, підходи до обслуговування клієнтів і бізнес-моделі. Гнучкість онлайн-торгівлі та зростаюча проникність Інтернету сприяють експоненційному збільшенню обсягів продажів у всіх регіонах світу. В Україні частка e-commerce в роздрібному товарообігу перевищила 12 % у 2024 р., що підтверджує високий попит на цифрові послуги [1].

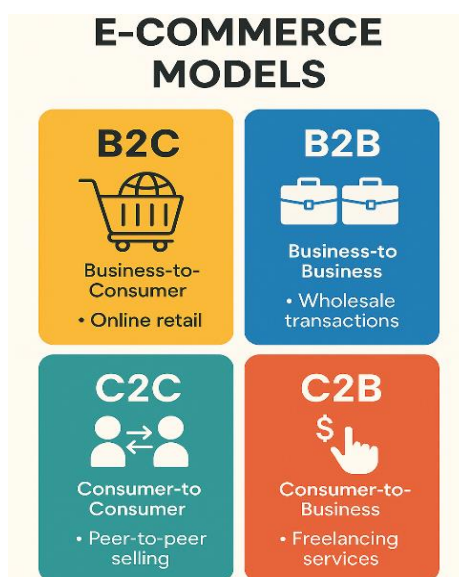


Рис. 1. Моделі комерції

Моделі електронної комерції поділяють на B2C, B2B, C2C та C2B (рис. 1). Кожна обслуговує різні сегменти ринку та передбачає відмінні механізми взаємодії між учасниками. Порівняння ключових характеристик наведено у табл. 1.

Таблиця 1.

Порівняння моделей електронної комерції

Модель	Сегмент	Ключові риси	Приклад
B2C	Бізнес → Споживач	Роздрібний продаж товарів/послуг онлайн	Rozetka, Amazon
B2B	Бізнес → Бізнес	Оптові закупівлі, корпоративні платформи	Prom.ua B2B, Alibaba
C2C	Споживач → Споживач	Peer-to-peer торгівля	eBay, OLX
C2B	Споживач → Бізнес	Фриланс-біржі, краудсорсинг	Upwork, Fiverr

За даними Української асоціації електронної комерції, упродовж 2021-2024 рр. обсяг онлайн-продажів зростає у середньому на 18 % щороку (табл. 2). Розвиток логістичної інфраструктури, зокрема PUDO-пунктів і мереж поштоматів, сприяє зниженню витрат на «остання милю» та підвищує задоволеність споживачів.

Таблиця 2.

Динаміка українського ринку e-commerce (2021–2024 рр.)

Рік	Обсяг ринку, млрд ₴	Δ, %	Онлайн-покупці, млн
2021	118	–	12.3
2022	135	14.4	13.7
2023	158	17.0	15.2
2024	186	17.7	17.6

Регуляторне середовище значно впливає на розвиток e-commerce. У ЄС запровадження PSD2 стимулює конкуренцію на ринку платежів, а GDPR визначає жорсткі вимоги до захисту персональних даних. В Україні діють Закон «Про електронну комерцію» № 675-VIII та оновлені правила дистанційної торгівлі, які гармонізують місцеве законодавство з директивами ЄС.

Вплив ключових технологій на e-commerce узагальнено в табл. 3 та на рис. 2.

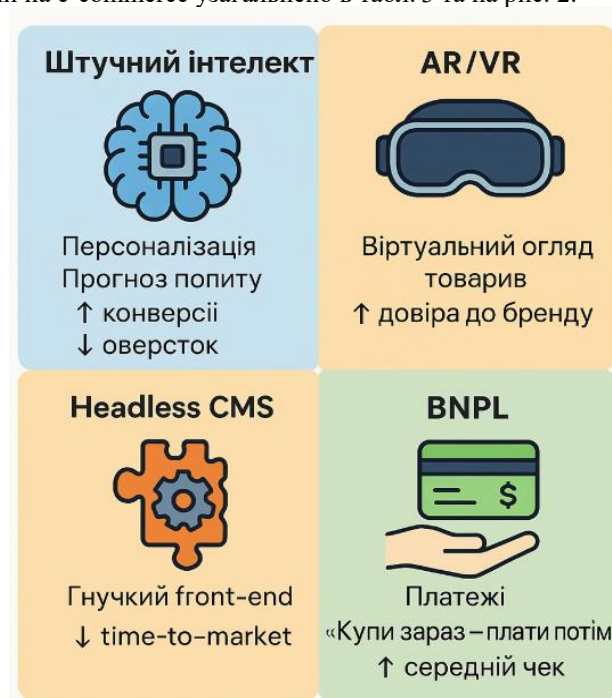


Рис. 2. Вплив провідних технологій на електронну комерцію

Основні виклики галузі включають кіберризик, забезпечення надійності ланцюгів постачання та необхідність персоналізованого сервісу без втрати приватності користувачів. Компанії впроваджують Zero-Trust Security та DDP-моделі, аби мінімізувати збитки від шахрайства й покращити клієнтський досвід.

Таблиця 3.

Вплив провідних технологій на електронну комерцію [2]

Технологія	Сфера застосування	Очікуваний ефект
Штучний інтелект	Персоналізація, прогноз попиту	↑ конверсії, ↓ оверсток
AR/VR	Віртуальний огляд товарів	↑ довіра до бренду
Headless CMS	Гнучкий front-end	↓ time-to-market
BNPL	Платежі	↑ середній чек

Провідні технологічні тренди 2025 р. (рис. 2):

- інтеграція III-рекомендацій, що підвищує коефіцієнт конверсії на 10–15 %;
- зростання частки платежів BNPL (Buy Now Pay Later) до 8 % у структурі онлайн-транзакцій;
- використання доповненої реальності (AR) для візуалізації товарів;
- впровадження headless-архітектур задля гнучкості front-end.

Згідно з прогнозами аналітиків eMarketer, глобальний обсяг ринку електронної комерції (e-commerce) до 2027 року перевищить 7 трильйонів доларів США, що свідчить про неймовірне зростання цієї галузі в наступні роки. Водночас, країни Центральної та Східної Європи, включаючи Україну, продовжать демонструвати значні темпи зростання: щорічний приріст у цих регіонах очікується на рівні 14–16% [3]. Це створює сприятливі умови для розвитку e-commerce у країнах, що ще не досягли максимальної насиченості ринку.

В Україні, в умовах стабільності макроекономічної ситуації та впровадження необхідних структурних реформ, сектор електронної комерції може досягти обороту в 280 мільярдів гривень до 2027 року. Такий показник стане результатом не тільки підвищення попиту на онлайн-послуги та товари, а й розвитку інфраструктури, зокрема, покращення логістичних процесів і адаптації місцевих компаній до змін в глобальних трендах.

Електронна комерція є одним з найбільш динамічних і перспективних сегментів цифрової економіки, оскільки вона не тільки забезпечує значні фінансові потоки, а й відкриває нові можливості для бізнесу, створюючи можливість залучення глобальних клієнтів без прив'язки до географічних меж. Проте для подальшого розвитку ринку e-commerce в Україні необхідно звернути особливу увагу на кілька ключових аспектів.

По-перше, адаптація бізнесу до постійно змінюваного регуляторного середовища стане важливим фактором. Регулювання електронної комерції вимагає постійного оновлення, оскільки зміни в законодавстві, таких як введення нових стандартів захисту персональних даних або фінансових транзакцій, можуть суттєво впливати на бізнес-процеси. По-друге, технологічні інновації, зокрема впровадження штучного інтелекту (AI), будуть визначати конкурентоспроможність на ринку. Інтелектуальні системи здатні покращити обслуговування клієнтів, підвищити персоналізацію та оптимізувати маркетингові стратегії.

Не менш важливим є забезпечення високого рівня безпеки даних користувачів. Захист персональних даних та фінансових транзакцій повинен залишатися в центрі уваги, оскільки зростання обсягів онлайн-торгівлі сприяє збільшенню кіберзагроз. Нарешті, підтримка ефективної логістики та омніканального підходу в обслуговуванні клієнтів є ще однією важливою умовою розвитку. В умовах зростаючої конкуренції на ринку e-commerce бізнеси мають бути готові до інтеграції різних каналів продажу, забезпечуючи зручність і швидкість обслуговування для своїх клієнтів, що підвищить їх лояльність та дозволить збільшити частку на ринку.

Таким чином, щоб максимально використати потенціал електронної комерції в Україні та на глобальному рівні, підприємства повинні не лише впроваджувати новітні технології, але й бути готовими до змін в регуляторному середовищі, активно працювати над підвищенням безпеки та ефективності своїх бізнес-процесів.

References

1. UNCTAD (2025). E-commerce and Digital Economy – Analytics and annual reports on the global development of electronic commerce. Retrieved from: <https://unctad.org/topic/ecommerce-and-digital-economy> (accessed: April 30, 2025).
2. Ecommerce Europe (2025). Reports & Insights – Industry reports and indicators of e-commerce in EU countries. Retrieved from: <https://ecommerce-europe.eu/research/> (accessed: April 30, 2025).
3. eMarketer / Insider Intelligence (2025). Statistics, forecasts, and trends in global retail online commerce. Retrieved from: <https://www.emarketer.com/> (accessed: April 30, 2025).

ELECTRONIC COMMERCE: CURRENT TRENDS, CHALLENGES AND FORECASTS

A comprehensive analysis of the current state and prospects of e-commerce as a driver of the digital economy has been carried out. The work examines the growth of e-commerce in Ukraine, highlighting the increasing share of online trade in the overall retail sales volume and new trends in digital business models. Special attention is given to various models of electronic commerce, such as B2C, B2B, C2C, and C2B, and their impact on market segments and consumer behavior.

The paper also analyzes the role of regulatory frameworks both in Ukraine and the European Union in shaping the e-commerce market. The influences of regulations such as PSD2 and GDPR in the EU, and their analogs in Ukrainian legislation, are discussed separately. In addition, technological innovations, particularly artificial intelligence (AI), augmented reality (AR), and Buy Now Pay Later (BNPL) services, are examined for their role in enhancing user experience, improving conversion rates, and reducing risks in supply chains.

The paper concludes with a forecast of the future development of e-commerce, including significant market growth and the integration of advanced technologies in the coming years. The necessity for businesses to adapt to changes in the regulatory, technological, and consumer environments to maintain competitiveness is emphasized.

Key words: *electronic commerce, digital economy, e-commerce models, regulatory frameworks, artificial intelligence, augmented reality, BNPL.*

Відомості про автора / Information about the Author

Яновський Олександр, здобувач спеціальності 121 «Інженерія програмного забезпечення», Чорноморський національний університет імені Петра Могили, м. Миколаїв, Україна. E-mail: ol.yanovskyi@gmail.com, orcid: <https://orcid.org/0009-0006-1105-4364>

Oleksandr Yanovskyi, 4th-year student of the specialty 121 “Software Engineering”, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ol.yanovskyi@gmail.com, orcid: <https://orcid.org/0009-0006-1105-4364>

Фаленкова Марина, старша викладачка кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили, м. Миколаїв, Україна. E-mail: marynakotova@gmail.com, orcid: <https://orcid.org/0000-0001-7797-0142>

Maryna Falenkova, Senior Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: marynakotova@gmail.com, orcid: <https://orcid.org/0000-0001-7797-0142>.

Решетнік Юрій, старший викладач кафедри інженерії програмного забезпечення, Чорноморський національний університет імені Петра Могили, м. Миколаїв, Україна. E-mail: jurij.reshetnik@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-7797-0142>

Jurij Reshetnik, Senior Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: jurij.reshetnik@chmnu.edu.ua, orcid: <https://orcid.org/0009-0004-7430-7467>.

DOI 10.34132/mspc2025.01.14.25

Roman Dinzhos
Nataliia Fialko

NON-ISOTHERMAL CRYSTALLIZATION OF POLYMER NANOCOMPOSITES BASED ON POLYCARBONATE FILLED WITH CARBON NANOTUBES

The article is devoted to theoretical and experimental studies of the dependence of the heat of crystallization of polymer nanocomposites on such factors as the mass fraction of the filler and the cooling rate of composites from the melt. Enthalpy relaxation in the cooling runs was monitored with the temperature-modulated DSC instrument (Perkin Elmer DSC-2 upgraded and supplied with signal processing software by the IFA GmbH, Ulm). Each sample was initially "overheated" by ~ 25 K above the apparent melting temperature of PP ($T_m \approx 490$ K), stored for 3 min. and cooled in the standard DSC mode to ~ 360 K at one of the six available constant cooling rates q (from 20 K/min down to 1 K/min). Studies of non-isothermal crystallization of polycarbonate and nanocomposites containing from 0 to 5 wt.% of carbon nanotubes have been conducted. It was established that the value of the heat of crystallization of nanocomposites significantly decreases with an increase in their cooling rate and an increase in the mass fraction of the filler. It is shown that the crystallization for polymer nanocomposites has the assumption of a superposition of intravenous initial (unimpeded) and secondary (limited) mechanisms of crystal formation, respectively.

Keywords: nanocomposite materials, carbon nanotubes, experimental studies, specific heat of crystallization.

Wide use of polymer nanocomposite materials requires a large amount of knowledge about their thermophysical properties [1-15]. One of the important thermophysical characteristics of polymer nanocomposites is their specific heat of crystallization. This determines the urgency of researching the patterns of change in the heat of crystallization of such materials.

This article is devoted to establishing the dependence of the specific heat of crystallization of polycarbonate-based composites filled with carbon nanotubes (CNTs) on a number of factors (melt cooling rate, mass fraction of filler, etc.).

The determination of the specific heat of crystallization q_c was based on the use of experimentally obtained exotherms of crystallization of the composite during its cooling from the melt at a given constant rate V_t (the method of construction of the indicated exotherms is given in [6, 8, 10]).

The value of q_c was determined by dependence

$$q_c = \frac{\int_{T_N}^{T_K} (Q_p - Q_p^{\min}) dT}{V_t}$$

where Q_p , $Q_p^{\min}$, is the current and minimum value of the specific heat flow; T – is the current temperature; T_N , T_K – temperature of the beginning and end of crystallization; V_t is the cooling rate [10, 12].

The value of the specified integral F was determined graphically as a component of the area under the crystallization exotherm.

Samples of materials for research were prepared by the method of hot pressing of the composition obtained as a result of mixing its components in a powder state in a magnetic stirrer. The carbon nanotubes used in the research were produced by the CVD (chemical vapor deposition) method.

The content of mineral impurities in them was ~ 0.1%. The specific surface area of CNTs determined by N_2 adsorption was 190 m²/g. The outer diameter of the CNT, found using the method of small-angle X-ray scattering, was 20 nm, the length - (1 ... 5) μm, the wall thickness ~ 5 nm. Manufacturer of carbon tubes - "Spetsmash" LLC.

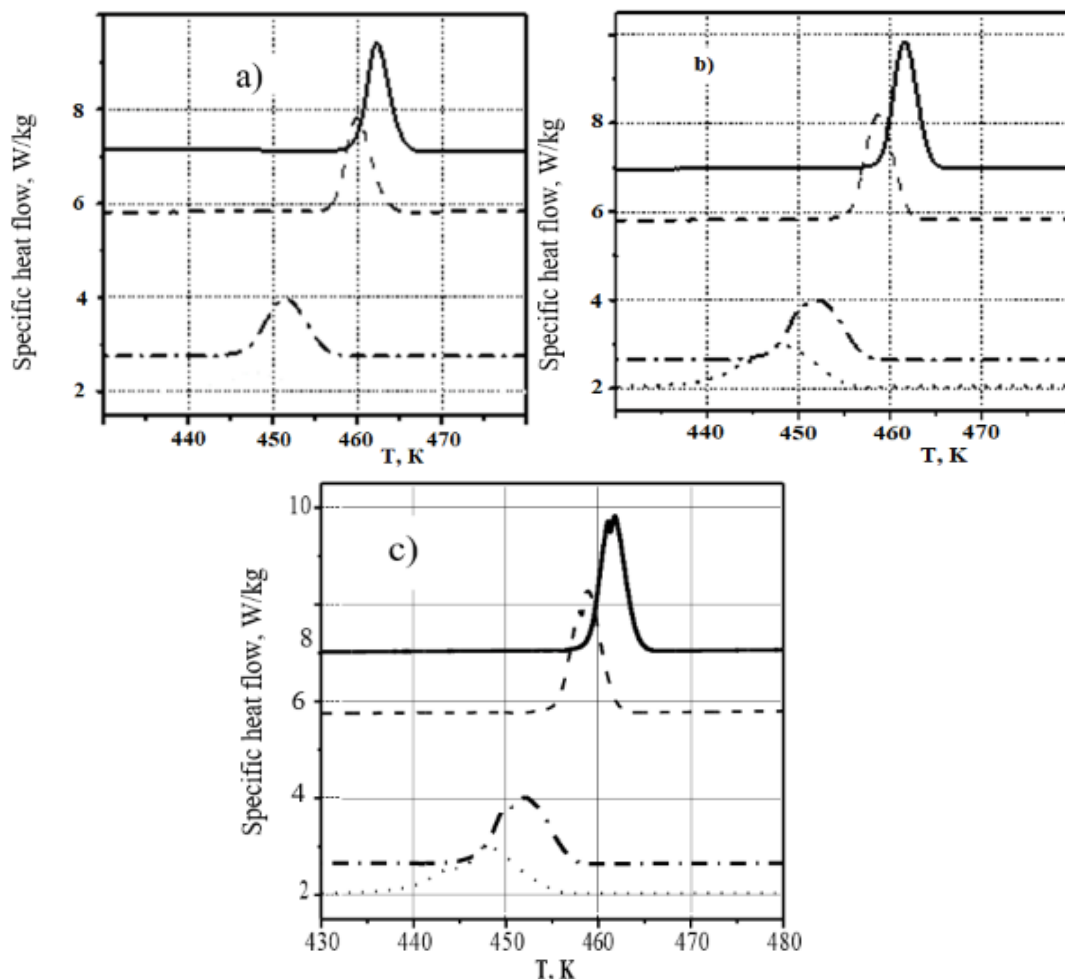


Fig. 1. Crystallization exotherms for polycarbonate-based composites with filler content $\omega = 0\%$ (a), 0.4% (b), 5% (c) for different cooling rates of the composites: 1 – $V_t = 0.0166$ K/c; 2 – $V_t = 0.0333$ K/c; 3 – $V_t = 0.1666$ K/c; 4 – $V_t = 0.3333$ K/c.

Figure 1 and table 1 present data relating to the determination of the specific heat of crystallization q_c for the considered composite materials. Figure 1 illustrates the crystallization exotherms obtained as a result of experimental studies, which were used to find the q_c values. In the table 1 shows the values of q_c , found according to dependence (1).

Let's move on to the consideration of experimentally obtained crystallization exotherms. Let us dwell briefly on the highlighted features of the influence of the cooling rate V_t and the mass fraction of the filler ω on the characteristics of the crystallization process for the investigated polymer nanocomposites (Fig. 1).

As for the cooling rate V_t , with an increase in V_t , there is a decrease in the maximum heat flow Q_p and its shift on the curve $Q_p = f(T)$ to the region of lower temperatures. In addition, with an increase in the cooling rate, there is a decrease in the temperatures of the beginning T_N and the end T_K of crystallization. For example, at $\omega = 4.0\%$ for $V_t = 0.0166$ K/c the temperature of the beginning and end of crystallization is 465.6 K and 458.0 K, and for $V_t = 0.3333$ K/c – 455.6 K and 436.7 K, respectively.

According to the obtained data, the mass fraction of CNTs significantly affects the nature of crystallization exotherms for polycarbonate-based composites. Namely, as ω increases, the unimodal peak on the curve $Q_p = f(T)$ transforms into a bimodal one.

Let us further analyze the dependence of the specific heat of crystallization q_c on various factors (table 1).

Table 1.

The value of the heat of crystallization q_c and the integral F for polymer nanocomposites based on polycarbonate filled with CNTs at different filler content ω and different cooling rates V_t of the composites

V_t , K/s	F , (W·K)/kg	q_c , J·kg
$\omega = 0\%$		
0,0166	8,78	529
0,0333	8,72	262
0,0833	8,69	104
0,1666	8,65	52
0,3333	8,06	24
$\omega = 0,4\%$		
0,0166	8,55	515
0,0333	8,49	256
0,0833	8,26	102
0,1666	8,08	51
0,3333	7,97	26
$\omega = 5,0\%$		
0,0166	7,46	449
0,0333	7,23	219
0,0833	7,03	86
0,1666	6,90	42
0,3333	3,17	11

The conducted studies showed that the heat of crystallization of composites decreases with the growth of the mass fraction of the filler. For example, at a cooling rate of $V_t = 0.0166$ K/s and $\omega = 4.0\%$, the value of q_c for a composite based on polycarbonate is less than q_c for pure polycarbonate by approximately 1.2 times.

It also follows from the obtained data that the value of q_c for the considered composites significantly decreases with increasing speed of their cooling. As can be seen from the table 1, for $\omega = 4.0\%$ with an increase in the speed of V_t from 0.0166 K/s (1 K/min) to 0.166 K/s (10 K/min), the value of q_c decreases from 448 J/kg to 41 J/kg about 10.9 times.

Thus, based on the results of experimental studies, the value of the specific heat of crystallization of nanocomposites based on polycarbonate when used as a filler for carbon nanotubes was determined. It is shown that the heat of crystallization of these composite materials depends significantly on the mass fraction of the filler and the cooling rate of the composite.

References

1. Privalko, V. P., Kawai, T., Lipalov, Y. S.: Crystallization of filled nylon 6 III. Non-isothermal crystallization. Colloid and Polymer Science 257 (10), 1042–1048 (1979).
2. Wunderlich, B.: Thermodynamics and kinetics of crystallization of flexible molecules. Journal of Polymer Science Part B: Polymer Physics 46 (24), 2647–2659 (2008).
3. Gao, Q., Jian, Z., Xu, J., Zhu, M., Chang, F., Han, A.: Crystallization kinetics of the $\text{Cu}_{50}\text{Zr}_{50}$ metallic glass under isothermal conditions. Journal of Solid State Chemistry 244, 116–119 (2016).
4. Rasana, N., Jayanarayanan, K., Pegoretti, A.: Non-isothermal crystallization kinetics of polypropylene/short glass fibre/multiwalled carbon nanotube composites. RSC Advances 8 (68), 39127–39139 (2018).
5. Roushan, A. H., Omrani, A.: Non-isothermal crystallization of NaX nanocrystals/poly (vinyl alcohol) nanocomposite. Journal of Thermal Analysis and Calorimetry 144(6) (2020).
6. De Melo, C.C.N., Beatrice, C.A.G., Pessan, L.A., Oliveira, A.D., Machado, F.M.: Analysis of nonisothermal crystallization kinetics of graphene oxide - reinforced polyamide 6 nanocomposites. Thermochimica Acta 667, 111–121 (2018).
7. Montanheiro, T. L., Menezes, B.R.C., Montagna, L.S., Beatrice, C. A. G., Marini, J., Lemes, A.P., Thim, G.P.: Non- Isothermal Crystallization Kinetics of Injection Grade PHBV and PHBV/Carbon Nanotubes Nanocomposites Using Isoconversional Method. Journal of Composites Science 4 (2), 52 (2020).

8. Menezes, B., Campos, T., Montanheiro, T., Ribas, R., Cividanes, L., Thim, G.: Non-Isothermal Crystallization Kinetic of Polyethylene/Carbon Nanotubes Nanocomposites Using an Isoconversional Method. *Journal of Composites Science* 3 (1), 21 (2019).
9. Mata-Padilla, J.M., Ávila-Orta, C.A., Almendárez-Camarillo, A., Martínez-Colunga, J.G., Hernández-Hernández, E., Cruz-Delgado, V.J., et. al.: Non-isothermal crystallization behavior of isotactic polypropylene/copper nanocomposites. *Journal of Thermal Analysis and Calorimetry* 143(57), (2020).
10. Privalko, V.P., Dinzhos, R.V., Privalko, E.G.: Enthalpy relaxation in the cooling/heating cycles of polypropylene/organosilica nanocomposites I: Non-isothermal crystallization *Thermochimica Acta* 432(1), 76-82 (2005).
11. Privalko, V.P., Dinzhos, R.V., Privalko, E.G.: Enthalpy relaxation in the cooling/heating cycles of polyamide 6/organoclay nanocomposites. II. Melting behavior. *Journal of Macromolecular Science - Physics* 44 (4), 431-443 (2005).
12. Dinzhos, R.V., Privalko, E.G., Privalko, V.P.: Enthalpy relaxation in the cooling/heating cycles of polyamide 6/organoclay nanocomposites. I. Nonisothermal crystallization. *Journal of Macromolecular Science - Physics* 44 (4), 421-430 (2005).
13. Privalko, V.P., Dinzhos, R.V., Privalko, E.G.: Enthalpy relaxation in the cooling/heating cycles of polypropylene/organosilica nanocomposites II. Melting behavior. *Thermochimica Acta* 428(1-2), 31-39 (2005).
14. Privalko, V.P., Dinzhos, R.V., Privalko, E.G.: Melting behavior of the nonisothermally crystallized polypropylene/organosilica nanocomposite. *Journal of Macromolecular Science - Physics* 43 (5), 979-988 (2004).
15. Privalko, V.P., Dinzhos, R.V., Rekheta, N.A., Calleja, F.J.B.: Structure-diamagnetic susceptibility correlations in regular alternating terpolymers of ethene and propene with carbon monoxide. *Journal of Macromolecular Science - Physics* 42 (5), 929-938 (2003).
16. Alekseev, O.M., Alekseev, S.O., Bulavin, L.A., Lazarenko, M.M., Maiko, O.M.: Phase transitions in chain molecular polycrystals of 1-octadecene. *Ukr. J. Phys.* 53, 882 (2008).

DOI 10.34132/mspc2025.01.14.25

Roman Dinzhos
Nataliia Fialko

НЕІЗОТЕРМІЧНА КРИСТАЛІЗАЦІЯ ПОЛІМЕРНИХ НАНОКОМПЗИТІВ НА ОСНОВІ ПОЛІКАРБОНАТУ, НАПОВНЕНОГО ВУГЛЕЦЕВИМИ НАНОТРУБКАМИ

Стаття присвячена теоретичним та експериментальним дослідженням залежності теплоти кристалізації полімерних нанокомпозитів від таких факторів, як масова частка наповнювача та швидкість охолодження композитів з розплаву. Релаксацію ентальпії в циклах охолодження відстежували за допомогою термомодульованого ДСК (Perkin Elmer DSC-2, модернізований і оснащений програмним забезпеченням для обробки сигналів фірми IFA GmbH, Ульм). Кожен зразок спочатку «перегрівали» на ~ 25 К вище видимої температури плавлення ПП ($T_m \approx 490$ К), витримували протягом 3 хв. і охолоджували в стандартному режимі ДСК до ~ 360 К при одній з шести доступних постійних швидкостей охолодження q^- (від 20 К/хв до 1 К/хв). Проведено дослідження неізотермічної кристалізації полікарбонату та нанокомпозитів, що містять від 0 до 5 мас.ч. вуглецевих нанотрубок. Встановлено, що значення теплоти кристалізації нанокомпозитів суттєво зменшується зі збільшенням швидкості їх охолодження та збільшенням масової частки наповнювача. Показано, що кристалізація для полімерних нанокомпозитів має припущення про суперпозицію внутрішньовенного початкового (безперешкодного) і вторинного (обмеженого) механізмів кристалоутворення відповідно.

Відомості про автора / Information about the Author

Роман Дінжос, докт. тех. наук, професор кафедри фізики та математики факультету комп'ютерних наук, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: dinzhosrv@gmail.com, orcid: <https://orcid.org/0000-0003-1105-2642>.

Roman Dinzhos, Doctor of Technical Sciences, Professor of the Department of Physics and Mathematics, Faculty of Computer Science, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: alnievt@ukr.net, orcid: <https://orcid.org/0000-0003-1105-2642>.

Наталія Фіалко, член. кор. НАН України, докт. тех., наук, проф., відділ теплофізики енергоефективних теплотехнологій Інституту технічної теплофізики НАН України, Київ, Україна, E-mail: nmfialko@ukr.net, orcid: <https://orcid.org/0000-0003-0116-7673>.

Nataliia Fialko, Corresponding Member of the National Academy of Sciences of Ukraine, Doctor of Technical Sciences, Professor, Department of Thermophysics of Energy Efficient Heat Technologies, Institute of Technical Thermophysics, National Academy of Sciences of Ukraine, Kyiv, Ukraine, ORCID: 0000-0003-0116-7673

DOI 10.34132/mspc2025.01.14.26

Bohdan Somryakov,
Viktor Ralenko

SQL DATA ANALYSIS FOR BIG DATA OPTIMIZATION: BEST PRACTICES AND PITFALLS

The theses investigates SQL optimization techniques essential for enhancing query performance, reducing resource consumption, and ensuring efficient data retrieval in large-scale database systems. It outlines seven key strategies: avoiding functions on WHERE columns to leverage indexes, partitioning large tables for faster searches, filtering data early to minimize processing, utilizing indexes for frequent queries, applying TOP (LIMIT) for testing, optimizing ON clauses in joins, and avoiding unnecessary DISTINCT operations.

These methods enhance query execution speed, scalability, and system performance in big data environments. Practical examples and common pitfalls are highlighted, providing actionable insights for database administrators and developers to build responsive, resource-efficient applications while managing complex queries and large datasets effectively.

Key words: SQL optimization, big data, query performance, indexing, partitioning, database scalability, data retrieval, query efficiency, performance tuning.

SQL optimization is crucial for improving query performance, reducing resource consumption, and ensuring faster data retrieval [1]. By efficiently structuring queries, databases can handle large datasets more effectively, minimizing the time taken for processing and enhancing the overall system performance [2]. This helps in achieving more responsive applications, especially when dealing with complex operations or large volumes of data. Proper optimization ensures that queries execute quickly, making the system scalable and able to handle increasing workloads without sacrificing performance. The top seven SQL optimization techniques include:

1) Avoid Functions on WHERE Columns.

Using functions on columns in the WHERE clause can prevent Sybase IQ from utilizing indexes effectively, which forces the database to perform full table scans, leading to slower query performance.

```
SELECT *  
FROM sales  
WHERE YEAR(sale_date) = 2025;
```

YEAR(sale_date) applies a function to the sale_date column, and Sybase IQ might not use an index on sale_date. Avoid Functions on Indexed Columns. Instead, rewrite the query to perform a direct comparison.

```
SELECT *  
FROM sales  
WHERE sale_date BETWEEN '2025-01-01' AND '2025-12-31';
```

2) Partition Large Tables.

Partitioning improves performance by distributing the data across smaller, more manageable physical blocks, which makes searches faster and more efficient when querying specific ranges of data.

To partition a table by date, for example, you might partition it by month or year. This way, queries filtering on the sale_date column will only scan the relevant partition.

```
CREATE TABLE sales (  
    sale_id INT,  
    sale_date DATE,
```

```
amount DECIMAL(10, 2)
)
PARTITION BY RANGE (sale_date)
(PARTITION p2020 VALUES LESS THAN '2023-01-01',
PARTITION p2021 VALUES LESS THAN '2024-01-01',
PARTITION p2022 VALUES LESS THAN '2025-01-01');
```

Queries that filter by year can take advantage of partitioning by scanning only the relevant partition. This minimizes unnecessary I/O operations and leads to faster query execution and improved overall performance.

```
SELECT *
FROM sales
WHERE sale_date BETWEEN '2024-01-01' AND '2024-12-31';
```

3) Filter Early

Filtering early in the query reduces the size of the data being processed, which can lead to significant performance gains. Apply WHERE filters as early as possible, especially before joins or aggregations.

```
SELECT customer_name, SUM(amount)
FROM customers
JOIN orders ON customers.customer_id = orders.customer_id
WHERE orders.status = 'completed'
GROUP BY customer_name;
```

Filtering the orders table first reduces the amount of data involved in the join operation.

```
SELECT customer_name, SUM(amount)
FROM customers
JOIN (
    SELECT customer_id, amount
    FROM orders
    WHERE status = 'completed'
) AS filtered_orders ON customers.customer_id = filtered_orders.customer_id
GROUP BY customer_name;
```

4) Use Indexes

If you often query by email in the customers table, create an index on that column. Queries filtering by email will use the index, significantly improving query performance.

```
CREATE INDEX idx_email ON customers(email);
SELECT *
FROM customers
WHERE email = 'bob@example.com';
```

That being said, too many indexes can slow down insert and update operations because the indexes must also be updated. Choose only relevant columns to index, especially those involved in frequent filtering and joining operations.

5) TOP (LIMIT)

TOP is useful for testing and to limit the result set in exploratory queries. It helps prevent full table scans and large result sets during query development.

```
SELECT TOP 10 *
FROM employees;
```

6) Filter early in ON clause

.....

A LEFT JOIN ensures that all rows from the left table (e.g., employees) are returned, even if there is no matching row in the right table (e.g., departments). If there is no matching row in the right table, the columns from the right table are filled with NULL values. The WHERE clause is used to filter rows after the join has been completed. It is applied to the final result set, which means it comes into play after SQL performs the join operation and has already combined the rows from multiple tables.

```
SELECT e.name, d.department_name
FROM employees e
LEFT JOIN departments d
  ON e.department_id = d.department_id
WHERE d.department_name = 'Engineering';
```

The ON clause is used to specify the join condition itself. It dictates how the rows from the two tables are combined. The condition d.department_name = 'Engineering' is part of the join condition itself, which means only the rows where the department is 'Engineering' are considered for joining with the employees table.

```
SELECT e.name, d.department_name
FROM employees e
LEFT JOIN departments d
  ON e.department_id = d.department_id
  AND d.department_name = 'Engineering';
```

The ON clause can be used to filter data before the join operation happens, so that unnecessary rows from the right table are excluded from the join process. This can help improve performance because you limit the amount of data being processed during the join. It also ensures that the result of a LEFT JOIN includes rows from the left table even if there is no match in the right table (i.e., it can preserve rows from the left table with NULL values in the columns of the right table).

7) Avoid DISTINCT when unnecessary

The query selects unique combinations of department_id and the count of active employees in each department.

```
SELECT DISTINCT department_id, COUNT(*)
FROM employees
WHERE status = 'active'
GROUP BY department_id;
```

DISTINCT is unnecessary when you're already using GROUP BY. The GROUP BY will automatically ensure that you get one row per department_id, so there's no need to add DISTINCT.

```
SELECT department_id, COUNT(*)
FROM employees
WHERE status = 'active'
GROUP BY department_id;
```

Using DISTINCT is very resource-intensive because it requires the database to sort and remove duplicate rows, which can increase processing time and memory usage, especially on large datasets.

In conclusion, optimizing SQL queries is crucial for maintaining efficient and scalable database performance, especially as data volumes increase. By employing strategies such as avoiding functions on WHERE columns, filtering early, using indexes, partitioning large tables, avoiding DISTINCT when unnecessary, and limiting data for testing, you can significantly enhance query performance and reduce resource consumption. These techniques help ensure that your database can handle complex queries and large datasets swiftly, leading to faster response times, improved user experience, and overall better system performance.

References

1. "SQL Performance Explained" by Markus Winand, 2012.
2. "Joe Celko's SQL for Smarties: Advanced SQL Programming" (The Morgan Kaufmann Series in Data Management Systems) by Joe Celko, 2005.

DOI 10.34132/mspc2025.01.14.26

**Богдан Сомряков
Віктор Раленко**

**АНАЛІЗ ДАНИХ SQL ДЛЯ ОПТИМІЗАЦІЇ ВЕЛИКИХ ДАНИХ: НАЙКРАЩІ ПРАКТИКИ ТА ПІДВОДНІ
КАМЕНІ**

Дослідження розглядає техніки оптимізації SQL, важливі для підвищення продуктивності запитів, зменшення споживання ресурсів та забезпечення ефективного отримання даних у великих базах даних. Визначено сім ключових стратегій: уникнення функцій у стовпцях WHERE для використання індексів, партиціонування великих таблиць для швидшого пошуку, рання фільтрація даних для зменшення обробки, використання індексів для частих запитів, застосування TOP (LIMIT) для тестування, оптимізація умов ON у з'єднаннях та уникнення непотрібного DISTINCT.

Ці методи підвищують швидкість виконання запитів, масштабованість і продуктивність системи у середовищах великих даних. Наведено практичні приклади та типові помилки, що надають корисні рекомендації адміністраторам баз даних і розробникам для створення швидких, ефективних застосунків при роботі зі складними запитами та великими наборами даних..

Ключові слова: оптимізація SQL, великі дані, продуктивність запитів, індексування, розділення на частини, масштабованість баз даних, отримання даних, ефективність запитів, налаштування продуктивності.

Відомості про авторів / Information about the Authors

Богдан Сомряков, здобувач кафедри інтелектуальних інформаційних систем факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна.

Bohdan Somryakov, student of the Faculty of Computer Sciences at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine.

Віктор Раленко, викладач кафедри інженерії програмного забезпечення факультету комп'ютерних наук Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

Viktor Ralenko, Lecturer at the Department of Software Engineering, Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ralenko.v@chmnu.edu.ua, orcid: <https://orcid.org/0009-0009-4161-8468>.

© Bohdan Somryakov, Viktor Ralenko

Стаття отримана редакцією 19.05.2025.

DOI 10.34132/mspc2025.01.14.27

Bohdan Somriakov,
 Ievgen Sidenko,
 Yuriy Kondratenko

COMPARATIVE ANALYSIS OF THE APPLICATION RESULTS OF MAX-MIN CONVOLUTION BASED ON VARIOUS T-NORM AND S-NORM OPERATORS

This paper presents a comparative analysis of the results of applying max-min convolution based on different t-norm and s-norm operators. The addition of two fuzzy triangular numbers is considered as an example. The authors developed software to compare t-norms and s-norms in max-min convolution using discretization of fuzzy triangular numbers. The results highlight the flexibility of t-norms and s-norms in modeling uncertainty, which makes them valuable for applications in control systems, data analysis, and decision support systems.

Keywords: Fuzzy logic, min-max convolution, t-norm, s-norm, fuzzy triangular number.

Fuzzy logic is a mathematical approach designed to manage uncertainty by permitting values between 0 and 1, instead of relying solely on binary true/false logic [1, 2]. Formula 1 illustrates the addition of two triangular fuzzy numbers. Formula 2 defines a fuzzy set A as a set of pairs $(x, \mu_A(x))$, where $\mu_A(x)$ represents the membership degree of each element x. Formula 3 presents the max-min convolution used for adding two fuzzy numbers [3]. Figure 1 illustrates the manual addition of two fuzzy numbers, (4,7,9) and (7,9,12), using the max-min convolution method.

$$A (+) B = [a_1 + b_1, a_2 + b_2], \forall A, B \in \mathbb{R} \quad (1)$$

$$A = \{ (x, \mu_A(x)) \}, \forall x \in E \quad (2)$$

$$\mu_{A+B}(y) = \max \text{ over all } x_1 + x_2 = y \text{ of } \min(\mu_A(x_1), \mu_B(x_2)) \quad (3)$$

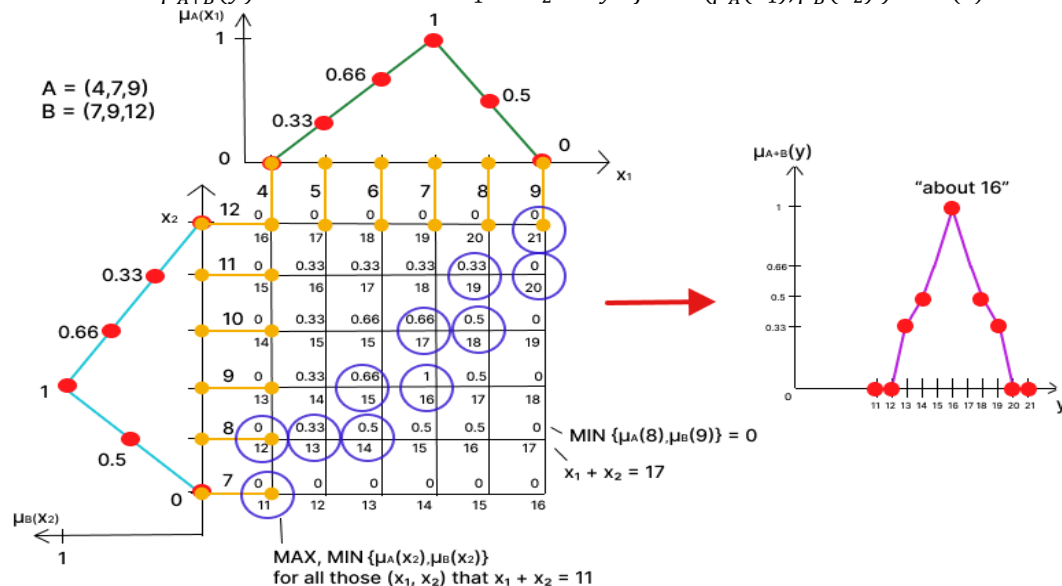


Fig. 1. Max-Min Convolution of Two Fuzzy Numbers Performed Manually
 Source: formed by the author based on [3]

Although minimum and maximum are the most commonly used operations in fuzzy logic convolutions, t-norms and s-norms provide more general ways to combine fuzzy values. T-norms model fuzzy intersection (AND) and s-norms model fuzzy union (OR), extending basic operations while satisfying important properties like associativity and commutativity. Different choices of t-norms and s-norms offer flexible modeling options, allowing for better alignment with the specific needs and characteristics of the problem at hand.

While the minimum (MIN) operation (4) is a commonly used t-norm in fuzzy logic, there are several other t-norms as well, including the product (PROD) (5), Hamacher product (6), Einstein product (7), drastic product (8), and bounded difference (9).

$$\mu_{A \cap B}(x) = \text{MIN}(\mu_A(x), \mu_B(x)) \quad (4)$$

$$\mu_{A \cap B}(x) = \mu_A(x) \cdot \mu_B(x) \quad (5)$$

$$\mu_{A \cap B}(x) = \frac{\mu_A(x) \cdot \mu_B(x)}{\mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x)} \quad (6)$$

$$\mu_{A \cap B}(x) = \mu_A(x) \cdot \frac{\mu_A(x) \cdot \mu_B(x)}{2 - [\mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x)]} \quad (7)$$

$$\mu_{A \cap B}(x) = \begin{cases} \text{MIN}(\mu_A(x), \mu_B(x)), & \text{for } \text{MAX}(\mu_A(x), \mu_B(x)) = 1 \\ 0, & \text{otherwise} \end{cases} \quad (8)$$

$$\mu_{A \cap B}(x) = \text{MAX}\{0, \mu_A(x) + \mu_B(x) - 1\} \quad (9)$$

The authors of this thesis developed software to compare different t-norms. Figure 2 illustrates the addition of the same two fuzzy numbers, (4,7,9) and (7,9,12), performed using the software through discretization into 20 points, with different t-norms applied for the calculation.

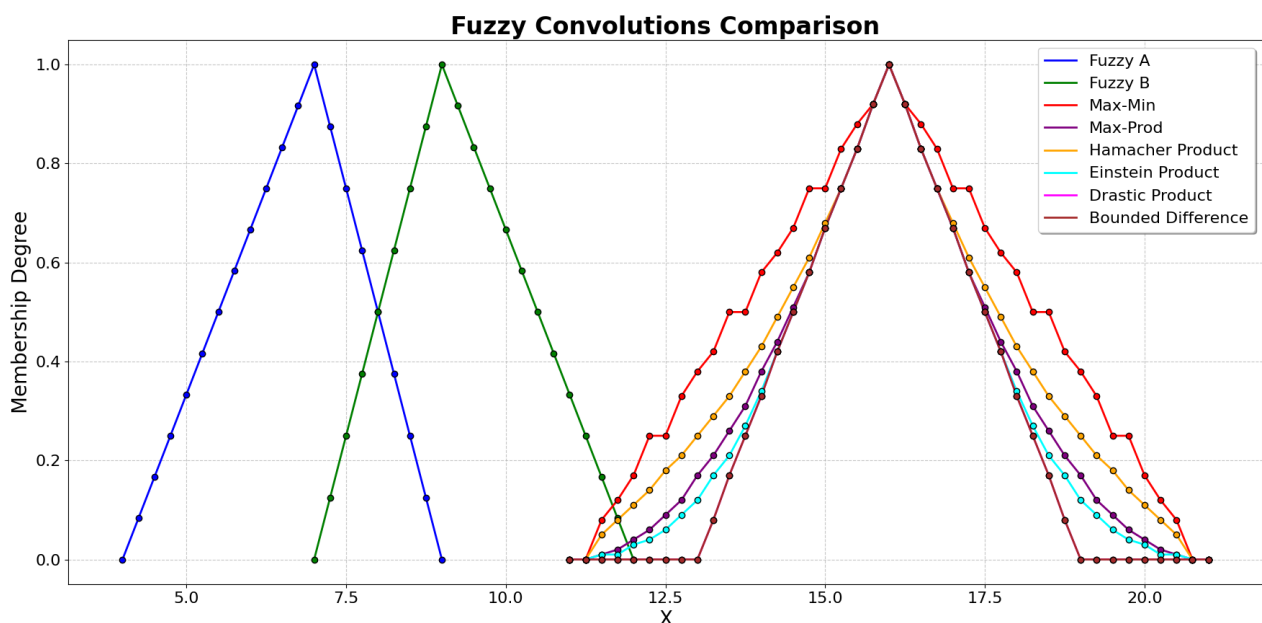


Fig. 2. Comparison of T-Norm Operators for Addition of Two Fuzzy Numbers (4,7,9) and (7,9,12) Performed via Software

Source: formed by the author

While the maximum (MAX) operation (10) is a commonly used s-norm in fuzzy logic, there are several other s-norms as well, including the algebraic sum (11), Gamakher's sum (12), Einstein's sum (13), reinforced sum (14), and bounded sum (15).

$$\mu_{A \cup B}(x) = \text{MAX}(\mu_A(x), \mu_B(x)) \quad (10)$$

$$\mu_{A \cup B}(x) = \mu_A(x) + \mu_B(x) - \mu_A(x) \cdot \mu_B(x) \quad (11)$$

$$\mu_{A \cup B}(x) = \frac{\mu_A(x) + \mu_B(x) - 2 \cdot \mu_A(x) \cdot \mu_B(x)}{1 - \mu_A(x) \cdot \mu_B(x)} \quad (12)$$

$$\mu_{A \cup B}(x) = \frac{\mu_A(x) + \mu_B(x)}{1 + \mu_A(x) \cdot \mu_B(x)} \quad (13)$$

$$\mu_{A \cup B}(x) = \begin{cases} \text{MAX}(\mu_A(x), \mu_B(x)), & \text{for } \text{MIN}(\mu_A(x), \mu_B(x)) = 0 \\ 1, & \text{otherwise} \end{cases} \quad (14)$$

$$\mu_{A \cup B}(x) = \text{MIN}\{1, \mu_A(x) + \mu_B(x)\} \quad (15)$$

The software also compares different s-norms. Figure 3 illustrates the addition of the same two fuzzy numbers, (4,7,9) and (7,9,12), performed using the software through discretization into 20 points, with different s-norms applied for the calculation.

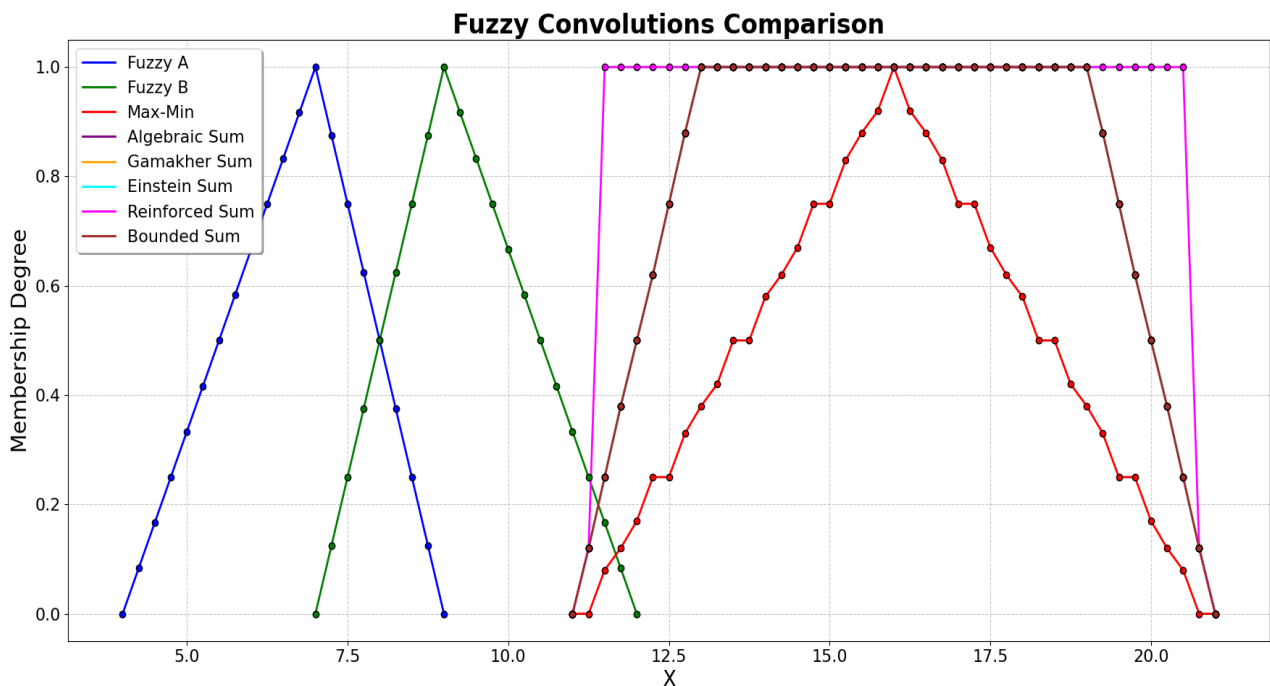


Fig. 3. Comparison of S-Norm Operators for Addition of Two Fuzzy Numbers (4,7,9) and (7,9,12) Performed via Software

Source: formed by the author

The use of different t-norms and s-norms in fuzzy logic provides a versatile framework for modeling uncertainty and handling complex decision-making scenarios. By offering flexibility in combining fuzzy values, these operators can be effectively applied in various fields such as control systems, data analysis, and decision support systems. Their ability to model different types of relationships between fuzzy sets makes them valuable tools for improving the accuracy and adaptability of systems dealing with imprecise or uncertain information.

References

1. Zadeh, L.A. (1965) Fuzzy sets. *Information and Control*, vol. 8, iss. 3, pp. 338-353.

2. Kondratenko, Y.P., Kondratenko, N.Y. (2018) Synthesis of Analytic Models for Subtraction of Fuzzy Numbers with Various Membership Function's Shapes. *Gil-Lafuente, A., Merigó, J., Dass, B., Verma, R. (eds) Applied Mathematics and Computational Intelligence. FIM 2015. Advances in Intelligent Systems and Computing*, vol. 730. Springer, Cham, pp. 87-100.
3. Piegat, A. (2001) Fuzzy Modeling and Control. Physica Heidelberg.

DOI 10.34132/mspc2025.01.14.27

**Богдан Сомряков,
Євген Сіденко
Юрій Кондратенко**

ПОРІВНЯЛЬНИЙ АНАЛІЗ РЕЗУЛЬТАТІВ ЗАСТОСУВАННЯ MAX-MIN ЗГОРТКИ НА ОСНОВІ РІЗНИХ ОПЕРАТОРІВ T-НОРМИ ТА S-НОРМИ

У даній роботі представлено порівняльний аналіз результатів застосування max-min згортки на основі різних операторів t-норми та s-норми. В якості прикладу розглянуто додавання двох нечітких трикутних чисел. Автори розробили програмне забезпечення для порівняння t-норм і s-норм в max-min згортці, застосовуючи дискретизацію нечітких трикутних чисел. Результати підкреслюють гнучкість t-норм і s-норм у моделюванні невизначеності, що робить їх цінними для застосувань у системах керування, аналізі даних і системах підтримки прийняття рішень.

Ключові слова: Нечітка логіка, max-min згортка, t-норми, s-норми, нечітке трикутне число.

Information about the Authors / Відомості про авторів

Bohdan Somriakov, student of the Computer Science Faculty, specialty 122 Computer Science, at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: pontius.bos@gmail.com, orcid: <https://orcid.org/0009-0003-7045-8586>.

Богдан Сомряков, студент факультету комп'ютерних наук, спеціальність 122 Комп'ютерні науки, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: pontius.bos@gmail.com, orcid: <https://orcid.org/0009-0003-7045-8586>.

Ievgen Sidenko, PhD, Associate Professor of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Євген Сіденко, канд. техн. наук, доцент кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: ievgen.sidenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-6496-2469>.

Yuriy Kondratenko, Doctor of Science, Professor, Head of the Intelligent Information Systems Department at Petro Mohyla Black Sea National University, Mykolaiv, Ukraine. E-mail: yuriy.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-7736-883X>.

Юрій Кондратенко, д-р техн. наук, професор, завідувач кафедри інтелектуальних інформаційних систем, Чорноморського національного університету імені Петра Могили, м. Миколаїв, Україна. E-mail: yuriy.kondratenko@chmnu.edu.ua, orcid: <https://orcid.org/0000-0001-7736-883X>.

Наукове видання

**МАТЕРІАЛИ НАУКОВО-ПРАКТИЧНИХ КОНФЕРЕНЦІЙ
ЧОРНОМОРСЬКОГО НАЦІОНАЛЬНОГО УНІВЕРСИТЕТУ
ІМЕНІ ПЕТРА МОГИЛИ**

**MATERIALS OF THE SCIENTIFIC AND PRACTICAL CONFERENCE
OF PETRO MOGYLA BLACK**

**Матеріали XXII Міжнародної наукової конференції
«ОЛЬВІЙСЬКИЙ ФОРУМ – 2025: стратегії країн
Причорноморського регіону в геополітичному просторі»**

№ 1 (1), 2025

СЕРІЯ: ТЕХНІЧНІ НАУКИ

Адреса оргкомітету:

м. Миколаїв, вул. 68 Десантників, 10, 54003
Тел.: 8 (0512) 50–03–32, 8 (0512) 76–55–81, 8 (0512) 76–55–99, факс: 50–00–69, 50–03–33,
E-mail: avi@chmnu.edu.ua, rector@chmnu.edu.ua

Редактор: *О. Михайлова.*
Комп'ютерна верстка, дизайн: *К. Гросу-Грабарчук*
Друк: *С. Волинець.*
Фальцювальні-палітурні роботи: *О. Мішалкіна.*

Підп. до друку 12.06.2025.
Формат 60х84^{1/8}. Папір офсет.
Гарнітура «Times New Roman». Друк ризограф.
Ум. друк. арк. 8. Обл.-вид. арк. 9,4. Зам. № 7056

Видавець і виготовлювач: ЧНУ ім. Петра Могили.
54003, м. Миколаїв, вул. 68 Десантників, 10.
Тел.: 8 (0512) 50–03–32, 8 (0512) 76–55–81, e-mail: rector@chmnu.edu.ua.
Свідоцтво суб'єкта видавничої справи ДК № 6124 від 05.04.2018.