

MINISTRY OF EDUCATION
AND SCIENCE OF UKRAINE

NATIONAL UNIVERSITY
OF LIFE AND ENVIRONMENTAL
SCIENCES OF UKRAINE

FACULTY OF INFORMATION
TECHNOLOGY

МІНІСТЕРСТВО ОСВІТИ
І НАУКИ УКРАЇНИ

НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ
БІОРЕСУРСІВ І
ПРИРОДОКОРИСТУВАННЯ УКРАЇНИ

ФАКУЛЬТЕТ ІНФОРМАЦІЙНИХ
ТЕХНОЛОГІЙ

PROCEEDINGS

XIII International scientific
and practical conference

**GLOBAL AND
REGIONAL PROBLEMS OF
INFORMATIZATION IN
SOCIETY AND
NATURE USING
'2025**

13-14 November 2025

Kyiv, NULES of Ukraine

Kyiv 2025

МАТЕРІАЛИ

XIII Міжнародної науково-
практичної конференції

**ГЛОБАЛЬНІ ТА
РЕГІОНАЛЬНІ ПРОБЛЕМИ
ІНФОРМАТИЗАЦІЇ В
СУСПІЛЬСТВІ І
ПРИРОДОКОРИСТУВАННІ
'2025**

13-14 листопада 2025 року

Київ, НУБіП України

Київ 2025

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ БІОРЕСУРСІВ
І ПРИРОДОКОРИСТУВАННЯ УКРАЇНИ
ФАКУЛЬТЕТ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ

МАТЕРІАЛИ

XIII Міжнародної науково-практичної конференції

ГЛОБАЛЬНІ ТА РЕГІОНАЛЬНІ ПРОБЛЕМИ ІНФОРМАТИЗАЦІЇ В СУСПІЛЬСТВІ І ПРИРОДОКОРИСТУВАННІ '2025

13-14 листопада 2025 року

Київ, НУБіП України

Київ 2025

УДК 004

Рекомендовано до друку вченою радою факультету інформаційних технологій Національного університету біоресурсів і природокористування України (протокол № 4 від 18.12.2025).

Укладач: д.т.н., доцент Шкарупило В.В.

Збірник матеріалів XIII Міжнародної науково-практичної конференції "Глобальні та регіональні проблеми інформатизації в суспільстві і природокористуванні '2025", 13–14 листопада 2025 року, НУБіП України, Київ. – К.: НУБіП України, 2025. – 206 с.

Відповідальність за зміст публікацій несуть автори.

© Національний університет біоресурсів
і природокористування України, 2025

Зміст

SECTION 1. MODELS, METHODS AND INFORMATION TECHNOLOGIES IN ECONOMICS / МОДЕЛІ, МЕТОДИ ТА ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ В ЕКОНОМІЦІ 5

ПЕРСПЕКТИВИ РОЗВИТКУ РИНКУ ТЕХНІЧНИХ КУЛЬТУР В УКРАЇНІ	5
<i>Тетяна Коваль</i>	5

ОЦІНКА СТАРТАПІВ ТА ЇЇ ВИЗНАЧАЛЬНІ ФАКТОРИ В УМОВАХ ГЕТЕРОГЕННОСТІ	8
<i>Христина Кашитанова, Марина Негрей</i>	8

ІНФЛЯЦІЙНІ ПРОЦЕСИ В УКРАЇНІ: МАКРОЕКОНОМІЧНИЙ АНАЛІЗ ТА ВПЛИВ НА ВИЩУ ОСВІТУ	11
<i>Інна Костенко, Валерія Яковлева</i>	11

АНАЛІЗ СПОЖИВЧОГО ПОПИТУ НА ТОВАРИ ТА ПОСЛУГИ В УМОВАХ ЦИФРОВОЇ ЕКОНОМІКИ УКРАЇНИ	15
<i>Інна Костенко, Юлія Терещенко</i>	15

ЦИФРОВА ТРАНСФОРМАЦІЯ АГРАРНОГО СЕКТОРУ: ЕКОНОМЕТРИЧНИЙ АНАЛІЗ ВПЛИВУ МЕРЕЖЕВОЇ ГОТОВНОСТІ НА ПРОДУКТИВНІСТЬ ТА ЗАЙНЯТІСТЬ	18
<i>Сергій Костенко</i>	18

ЦИФРОВА ТРАНСФОРМАЦІЯ АГРАРНОГО СЕКТОРУ: ЕКОНОМЕТРИЧНИЙ АНАЛІЗ ВПЛИВУ МЕРЕЖЕВОЇ ГОТОВНОСТІ НА ПРОДУКТИВНІСТЬ ТА ЗАЙНЯТІСТЬ	23
<i>Наталія Рогоза</i>	23

ANALYSIS OF THE VEGETABLE MARKET IN UKRAINE IN CONDITIONS OF WAR	26
<i>Катерина Наконечна</i>	26

АНАЛІЗ ТЕМ І НАСТРОЇВ У ЦИФРОВОМУ ДИСКУРСІ АГРАРНОГО СЕКТОРУ	30
<i>Роман Руденський, Володимир Кравченко</i>	30

ПРО ДЕЯКІ ПІДХОДИ ДО МОДЕЛЮВАННЯ СИСТЕМИ ПЕНСІЙНОГО СТРАХУВАННЯ	33
<i>Людмила Галаєва</i>	33

SECTION 2. COMPUTER SYSTEMS AND NETWORKS, CYBERSECURITY / КОМП'ЮТЕРНІ СИСТЕМИ І МЕРЕЖІ, КІБЕРБЕЗПЕКА 36

ВБУДОВАНА ІОТ СИСТЕМА КОНТРОЛЮ ДОСТУПУ З БІОМЕТРИЧНОЮ АВТЕНТИФІКАЦІЄЮ НА БАЗІ ESP32 ДЛЯ ІНФРАСТРУКТУРИ РОЗУМНИХ БУДІВЕЛЬ	36
<i>Максим Місюра, Сергій Мамченко</i>	36

ВИКОРИСТАННЯ MESH-СИСТЕМ У АГРАРНО-ПРОМИСЛОВОМУ КОМПЛЕКСІ УКРАЇНИ	40
<i>Сергій Мамченко, Максим Місюра, Микита Журавльов</i>	40

ОПТИМІЗАЦІЯ ТЕСТОВИХ КОНФІГУРАЦІЙ БАГАТОКОМПОНЕНТНИХ ІНФОРМАЦІЙНИХ СИСТЕМ ЗА ДОПОМОГОЮ PAIRWISE TESTING	43
<i>Кирило Кохан</i>	43
РОЗРОБКА СИСТЕМИ МОНІТОРИНГУ ТА ПРОГНОЗУВАННЯ РОЇННЯ БДЖІЛ З ВИКОРИСТАННЯМ ІОТ-ТЕХНОЛОГІЙ	46
<i>Ірина Бенцал, Леся Дмитроца</i>	46
ОГЛЯД ДОСЛІДЖЕНЬ ФЕНОМЕНУ ВТОМИ ВІД ЗАПИТІВ НА ДОСТУП	49
<i>Олексій Коваленко, Микола Богдюк</i>	49
SURVEY ON THE EVOLUTION AND KEY TECHNOLOGIES OF DIFFERENTIAL DISTINGUISHER DESIGN	52
<i>Цзян Сюе, Lakhno Valerii</i>	52
DEVELOPING A SYSTEM FOR ONTOLOGY-BASED DESIGN OF ITC ELECTRONIC COURSES	56
<i>Nazym Sabitova</i>	56
VIDEO GAME SECURITY IN 2025: HYBRID ANTI-CHEAT ARCHITECTURES, BEHAVIORAL ANALYTICS, AND LESSONS FOR REAL-TIME CYBER-PHYSICAL SYSTEMS	59
<i>Volodymyr Nazarenko</i>	59
ПІДХІД ДО ПЕРЕВІРКИ ТА КОРИГУВАННЯ ПОМИЛОК В ТЕКСТІ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ	62
<i>Микола Вертман, Степан Скрупський</i>	62
СИСТЕМНИЙ ПІДХІД ДО ВИБОРУ ТЕХНОЛОГІЙ ВЕБРОЗРОБКИ	67
<i>Олександр Рюмін, Степан Скрупський</i>	67
СИСТЕМА ЗБИРАННЯ ТА ВІДОБРАЖЕННЯ ЗОБРАЖЕНЬ НА ОСНОВІ STM32	72
<i>Олексій Коваленко, Лі Лі</i>	72
МОДЕЛЬНО-ПРЕДИКТИВНЕ КЕРУВАННЯ В ROS 2 ДЛЯ МОБІЛЬНИХ ПЛАТФОРМ ІЗ LIDAR-ЛОКАЛІЗАЦІЮ: АРХІТЕКТУРА ТА ЗВ'ЯЗОК	75
<i>Богдан Остроушко, Семен Волошин</i>	75
SMALL TARGET DETECTION METHOD BASED ON YOLOV8S UAV PERSPECTIVE	78
<i>Liu Shijun, Vadym Shkarupylo</i>	78
ПРИНЦИПОВА СХЕМА ІНТЕЛЕКТУАЛЬНОГО РОБОТИЗОВАНОГО КОМПЛЕКСУ МОНІТОРИНГУ ҐРУНТІВ УРАЖЕНИХ ВНАСЛІДОК ВІЙСЬКОВИХ ДІЙ	82
<i>Олексій Циганов, Ігор Болбот</i>	82
SECTION 3. DATA PROCESSING AND SOFTWARE SYSTEMS DEVELOPMENT/ ТЕХНОЛОГІЇ ОБРОБКИ ДАНИХ ТА СТВОРЕННЯ ПРОГРАМНИХ СИСТЕМ	85
DEFORMABLE ATTENTION U-NET NETWORK BASED ON FULLY CONVOLUTIONAL DISCRIMINATOR FOR SEMANTIC SEGMENTATION OF FETAL BRAIN ULTRASOUND IMAGES	85
<i>Yulong Li, Dmytro Nikolaienko</i>	85
ІНТЕРАКТИВНИЙ НАВЧАЛЬНИЙ ТРЕНАЖЕР ДЛЯ ВІЗУАЛІЗАЦІЇ ОСНОВНИХ ПОНЯТЬ ДИСКРЕТНОЇ МАТЕМАТИКИ	88
<i>Яніна Криворучко, Богдан Возний</i>	88

ПОРІВНЯЛЬНИЙ АНАЛІЗ ПРОДУКТИВНОСТІ ТА БЕЗПЕКИ ТИПІВ ORM SEQUELIZE ТА PRISMA У СУЧАСНИХ NODE.JS-СИСТЕМАХ	91
<i>Дарій Ратушний, Роман Руденський</i>	91
SEMANTIC SEARCH IN DOCUMENTATION TO ENHANCE THE EFFICIENCY OF UI COMPONENT GENERATION USING LANGUAGE MODELS	94
<i>Maksym Nedoshev, Viktor Kyrychenko</i>	94
ВЕКТОРНІ БАЗИ ДАНИХ У СУЧАСНИХ ШІ-ЧАТАХ	97
<i>Ярослав Гордій</i>	97
КОНТРОЛЬ ЯКОСТІ 3D-ДРУКОВАНИХ ВИРОБІВ ЗА ДОПОМОГОЮ КОМП'ЮТЕРНОГО ЗОРУ	102
<i>Матвій Горбач</i>	102
РОЗПІЗНАВАННЯ, КЛАСИФІКАЦІЯ ТА ВІДСТЕЖЕННЯ ОБ'ЄКТІВ В УМОВАХ ДИНАМІЧНОГО СПОСТЕРЕЖЕННЯ	105
<i>Денис Віннічук</i>	105
МЕТОДОЛОГІЧНИЙ АНАЛІЗ ПІДХОДІВ В АНАЛІТИЦІ BIG DATA	108
<i>Олег Шубалий, Леся Дмитроца</i>	108
ТЕОРЕТИЧНА МОДЕЛЬ МЕТАРІВНЕВОГО ВИБОРУ АЛГОРИТМІВ КЛАСТЕРИЗАЦІЇ ЗА ОПИСОВИМИ ХАРАКТЕРИСТИКАМИ НАБОРУ ДАНИХ ДЛЯ СИСТЕМ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ	111
<i>Ганна Вайганг, Юрій Науринський</i>	111
ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ДЛЯ СИСТЕМИ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ПРИ ВИРОЩУВАННІ ЗЕРНОВИХ КУЛЬТУР	115
<i>Марина Лендел, Белла Голуб</i>	115
ГІБРИДНА МОДЕЛЬ НЕЙРОННОГО ПОШУКУ: ПОЄДНАННЯ СЛОВНИКОВОГО ТА СЕМАНТИЧНОГО ПІДХОДІВ	118
<i>Ганна Вайганг, Іван Корнілов</i>	118
ОСОБЛИВОСТІ ВИКОРИСТАННЯ НЕЙРОМЕРЕЖ В СИСТЕМАХ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ УПРАВЛІННЯ ВРОЖАЙНІСТЮ	122
<i>Володимир Хиленко, Олексій Степанов, Bence Ladóczy</i>	122
ВІЗУАЛІЗАЦІЯ АЛГОРИТМУ ПОШУКУ ОПТИМАЛЬНОГО ШЛЯХУ МІЖ ДВОМА ТОЧКАМИ КЛІТИННОГО ЛАБІРИНТУ З ПЕРЕШКОДАМИ, ЩО ДИНАМІЧНО ЗМІНЮЮТЬСЯ	125
<i>Юрій Міловідов</i>	125
SECTION 4. INFORMATION SYSTEMS AND TECHNOLOGIES IN THE ECONOMY, TECHNOLOGY AND NATURAL USE / ІНФОРМАЦІЙНІ СИСТЕМИ І ТЕХНОЛОГІЇ В ЕКОНОМІЦІ, ТЕХНІЦІ ТА ПРИРОДОКОРИСТУВАННІ	128
ІНТЕГРАЦІЯ LMS MOODLE ТА CMS WORDPRESS ДЛЯ СТВОРЕННЯ ВІДКРИТОГО ОСВІТНЬОГО СЕРЕДОВИЩА ДЛЯ ДОДАТКОВОГО САМОСТІЙНОГО НАВЧАННЯ МАГІСТРІВ	128
<i>Максим Мокрієв</i>	128
АВТОМАТИЗАЦІЯ ЕКОЛОГІЧНОГО МОНІТОРИНГУ ЗАСОБАМИ СУЧАСНИХ ІТ	131
<i>Вікторія Усік</i>	131
ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ МОНІТОРИНГУ СТАНУ ДОВКІЛЛЯ В УМОВАХ ЗМІНИ КЛІМАТУ	134

<i>Анна Паламарчук</i>	134
МОДУЛЬ МАШИННОГО ЗОРУ ДЛЯ РОЗПІЗНАВАННЯ ОБ'ЄКТІВ	137
<i>Вікторія Смолій, Натан Смолій</i>	137
ГІБРИДНА БАЗА ДАНИХ ДЛЯ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ ОЦІНКИ ДЕГРАДОВАНИХ ҐРУНТІВ	140
<i>Олексій Коваль, Ігор Болбот</i>	140
МОДЕЛІ ВПРОВАДЖЕННЯ ЕЛЕКТРОННОГО ДОРАДНИЦТВА У СФЕРІ АГРОПРОМИСЛОВОГО КОМПЛЕКСУ УКРАЇНИ	143
<i>Христина Дрогобицька</i>	143
ВИЗНАЧЕННЯ ФУНКЦІОНАЛЬНИХ МОДУЛІВ ТА АРХІТЕКТУРНІ РІШЕННЯ ДЛЯ ІНТЕГРАЦІЇ МОБІЛЬНОГО ЗАСТОСУНКУ В ЕКОСИСТЕМУ УНІВЕРСИТЕТУ	146
<i>Ярослав Понзель</i>	146
ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ МАКРОЕКОНОМІЧНИХ ПОКАЗНИКІВ У ВИЗНАЧЕННІ ТЕНДЕНЦІЙ РОЗВИТКУ ЕКОНОМІК КРАЇН G7 ЗА ПЕРІОД 1999– 2023 РР.	148
<i>Владислав Шантала</i>	148
ОГЛЯД МОДЕЛЕЙ ПРОГНОЗУВАННЯ РЕЗУЛЬТАТІВ НАВЧАННЯ СТУДЕНТІВ У ЗАКЛАДАХ ВИЩОЇ ОСВІТИ НА ОСНОВІ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ	151
<i>Вадим Олійник, Тетяна Волошина</i>	151
ВИКОРИСТАННЯ НЕЙРОННИХ МЕРЕЖ В МОДЕЛЮВАННІ УПРАВЛІННЯ АГРАРНИМИ ПІДПРИЄМСТВАМИ	154
<i>Михайло Садко</i>	154
РОЗВИТОК ПЛАТФОРМИ ETHEREUM ТА ЇЇ ЗАСТОСУВАННЯ В СІЛЬСЬКОМУ ГОСПОДАРСТВІ	157
<i>Костянтин Рогоза</i>	157
ЦИФРОВЕ ОСВІТНЄ СЕРЕДОВИЩЕ: ВІД ІНСТРУМЕНТУ ДО СИСТЕМИ УПРАВЛІННЯ. ПРОБЛЕМИ АРХІТЕКТУРИ ТА ВПРОВАДЖЕННЯ	160
<i>Владислав Стеценко, Олена Глазунова, Петро Дрозд, Валентина Корольчук</i>	160
ТРАНСФОРМАЦІЯ СІЛЬСЬКОГОСПОДАРСЬКОГО ЕЛЕКТРОННОГО ДОРАДНИЦТВА ЗА ДОПОМОГОЮ ІНСТРУМЕНТІВ ШТУЧНОГО ІНТЕЛЕКТУ	163
<i>Михайло Швиденко, Сергій Саяпін</i>	163
SECTION 5. AUTOMATION, COMPUTER-INTEGRATED TECHNOLOGIES, ROBOTICS, ARTIFICIAL INTELLIGENCE/АВТОМАТИЗАЦІЯ, КОМП'ЮТЕРНО-ІНТЕГРОВАНІ ТЕХНОЛОГІЇ, РОБОТОТЕХНІКА, ШТУЧНИЙ ІНТЕЛЕКТ	167
ФУНКЦІЇ БАЖАНОСТІ ХАРІНГТОНА ДЛЯ ОЦІНКИ ЯКІСНИХ ПАРАМЕТРІВ БІОСИРОВИНИ ДЛЯ ЗБРОДЖУВАННЯ	167
<i>Сергій Павлов, Віталій Лисенко, Тарас Лендел, Катерина Наконечна</i>	167
КОНЦЕПТУАЛЬНИЙ ФРЕЙМВОРК УНІВЕРСАЛЬНОЇ ІШІ ДЛЯ ОПТИМІЗАЦІЇ РОБІТ БІОГАЗОВИХ УСТАНОВОК	171
<i>Станіслав Усенко, Максим Клименко, Євгеній Шаповалов, Сергій Жадан</i>	171

NOTEBOOKLM ЯК ПОТУЖНИЙ ШІ ІНСТРУМЕНТ ДЛЯ ОПРАЦЮВАННЯ ДЖЕРЕЛ З АКАДЕМІЧНИХ ДИСЦИПЛІН	177
<i>Валентина Марусяк</i>	177
ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ МОДЕЛЮВАННЯ ТА ОПТИМІЗАЦІЇ РОЗКЛАДУ ЗАНЯТЬ У ВИЩИХ НАВЧАЛЬНИХ ЗАКЛАДАХ	180
<i>Віталій Котляр</i>	180
ШТУЧНИЙ ІНТЕЛЕКТ ТА МУЛЬТИАГЕНТНІ СИСТЕМИ ДЛЯ ВИРІШЕННЯ ЗАДАЧ КООРДИНАЦІЇ АВТОНОМНИХ АГЕНТІВ	182
<i>Дмитро Кравченко</i>	182
ІДЕНТИФІКАЦІЯ СТАНО РОСЛИН У ЗАКРИТОМУ ҐРУНТІ НА ОСНОВІ СПЕКТРАЛЬНИХ ІНДЕКСІВ	185
<i>Іван Савченко, Ігор Болбот</i>	185
ШІ ДЛЯ РІШЕНЬ У БІЗНЕСІ ТА ПУБЛІЧНОМУ СЕКТОРІ БІБЛІОМЕТРИЧНИЙ ОГЛЯД	188
<i>Андрій Якушин, Марина Негрей</i>	188
ІНТЕЛЕКТУАЛЬНА СИСТЕМА АНАЛІТИКИ ДЛЯ ПІДТРИМКИ УПРАВЛІНСЬКИХ РІШЕНЬ НА ВИРОБНИЧОМУ ПІДПРИЄМСТВІ	191
<i>Владислав Бездетко, Анастасія Тройніна, Вікторія Рувінська</i>	191
ПРОГНОЗУВАННЯ ВРОЖАЙНОСТІ ТА ЯКОСТІ ПОСІВІВ КУКУРУДЗИ ДЛЯ СТЕПОВОЇ ЗОНИ УКРАЇНИ	195
<i>Віталій Лисенко, Микола Доля, Тарас Лендєл</i>	195
RESEARCH ON DEHAZING OF SURVEILLANCE IMAGES BASED ON DEEP LEARNING	198
<i>Gao Hui, Zhou Meng zhen</i>	198

SECTION 1. MODELS, METHODS AND INFORMATION TECHNOLOGIES IN ECONOMICS / МОДЕЛІ, МЕТОДИ ТА ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ В ЕКОНОМІЦІ

Тетяна Коваль

К. ф.-м. н., доцент кафедри економічної кібернетики

НУБіП України, факультет ІТ

ORCID ID 0000-0002-3981-5843

tkoval@nubip.edu.ua

ПЕРСПЕКТИВИ РОЗВИТКУ РИНКУ ТЕХНІЧНИХ КУЛЬТУР В УКРАЇНІ

Анотація. Розглянуто поточний стан ринку технічних культур в Україні та перспективи його подальшого розвитку. Визначено ключові тенденції, виклики та можливості зростання в галузі. Проаналізовано вплив процесів на міжнародному ринку, зміни в структурі виробництва та потенціал внутрішньої переробки, а також запропоновано напрямки підвищення ефективності ринку технічних культур в умовах міжнародної конкуренції.

Ключові слова: технічні культури, ринок, виробництво, експорт, рентабельність, біоенергетика.

ВСТУП

Постановка проблеми.

В Україні технічні культури є одним із найважливіших секторів сільськогосподарського виробництва, що забезпечує сировиною харчову, легку, хімічну, фармацевтичну та біоенергетичну промисловість. До них належать соняшник, ріпак, соя, цукровий буряк, льон, коноплі, гірчиця та сафлор. Ринок технічних культур розвивається завдяки зростаючому попиту на продукти переробки, такі як рослинні олії, цукор, біодизель, клітковина та білкові концентрати. Україна, маючи сприятливі кліматичні умови та родючий ґрунт, є одним із провідних експортерів промислових культур, зокрема соняшнику, ріпаку та цукрового буряка. Збільшення виробництва цих культур свідчить про високий ринковий потенціал, але необхідна подальша модернізація, диверсифікація структури культур та розвиток переробки на місці.

Аналіз останніх досліджень і публікацій.

Теоретичні та практичні аспекти розвитку ринку технічних культур розглядали у своїх працях І.В. Чехова [1], Т.А. Жадан [2], С.В. Тимошенко [3]. Дослідники проаналізували тенденції виробництва, експортний потенціал, динаміку цін та перспективи диверсифікації галузі.

Вчені зазначили, що ключовим до сталого розвитку ринку є зосередження на внутрішній переробці, впровадження інноваційних технологій та розширення асортименту нішевих культур.

Мета публікації. Узагальнити та дослідити наявні тенденції розвитку ринку технічних культур. Проаналізувати сучасний стан та визначити перспективи розвитку ринку технічних культур в Україні, а також розроблення рекомендацій щодо підвищення його ефективності й конкурентоспроможності.

РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Виробництво промислових культур в Україні стабільно зростає протягом останнього десятиліття. Основними напрямками виробництва є соняшник (понад 50% загальної площі промислових культур), ріпак (25-30%), соя та цукровий буряк. Виробництво льону та конопель відновлюється завдяки розвитку нішевих та органічних ринків. [1]

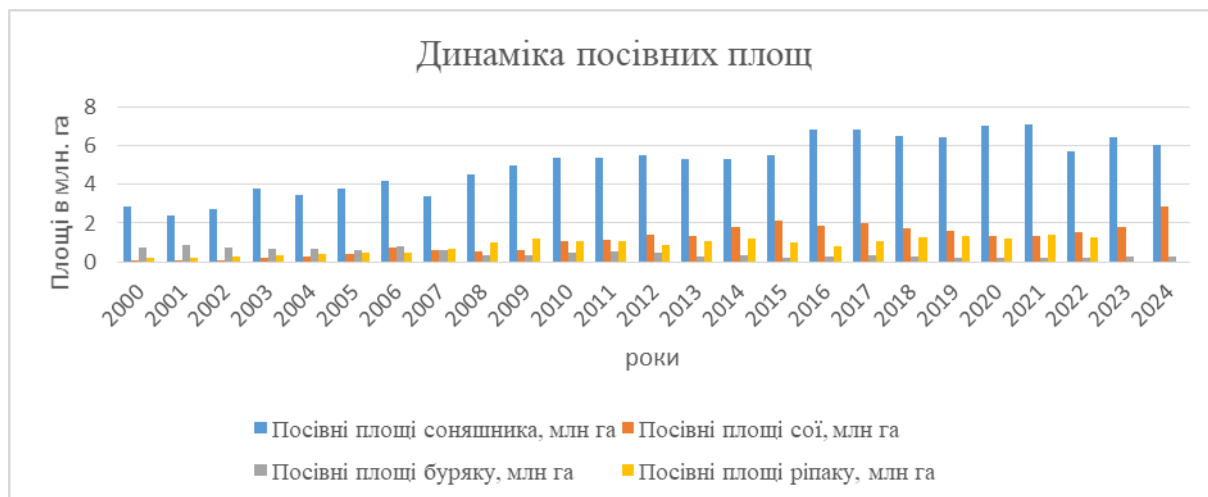


Рис. 1. Динаміка посівних площ. Складено на основі[4]

За даними Державної служби статистики України, загальний урожай промислових культур у 2023 році перевищив 20 мільйонів тонн. Висока частка експорту підкреслює експортно-орієнтований характер галузі. Приблизно 70% продукції промислових культур продається на зовнішніх ринках, переважно в ЄС, Індію, Китай та Туреччину. [2,4]

Світовий ринок технічних культур характеризується стабільним зростанням попиту на продукти переробки — рослинні олії, біопаливо та цукор. Основними виробниками технічних культур є США, Китай, Індія, Бразилія, Аргентина та країни ЄС.

Розвиток біоенергетичної галузі сприяє підвищенню цін на сировину, зокрема ріпак та цукровий буряк, що стимулює їх вирощування. Україна має потенціал посилити свою присутність на цьому ринку за рахунок збільшення виробництва ріпаку та сої для виробництва біодизелю. [3]

Розглянемо прогнозування такої важливої технічної культури ,як цукровий буряк.

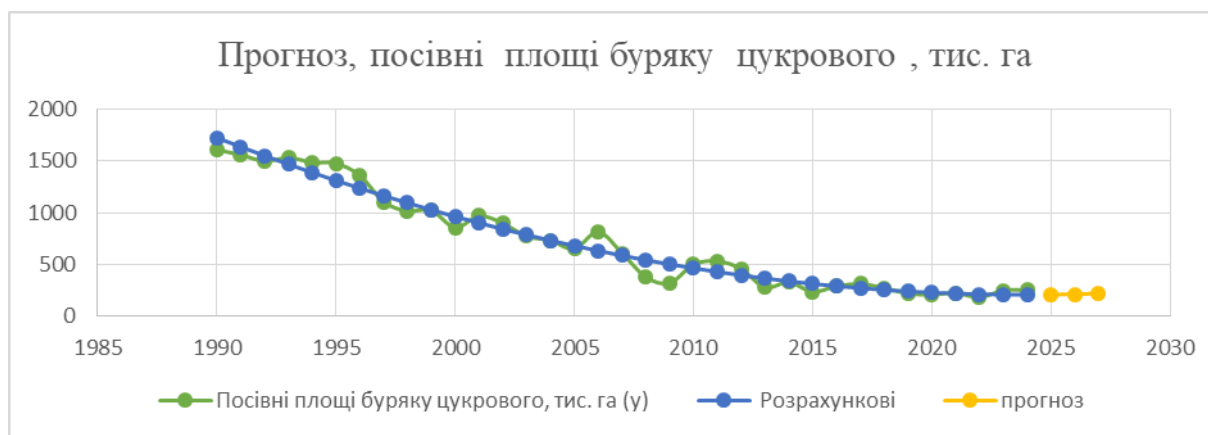


Рис. 2. Прогноз посівних площ буряку цукрового, тис. га. Складено на основі[4]

Так, аналіз статистичних даних 1995-2024 рр. Показав, що після значного скорочення посівних площ у 1990-х — на початку 2000-х років, вирощування цукрового буряка в Україні у 2020-х демонструє часткове відновлення. Втім, за

прогнозами на 2025 рік очікується нове зменшення площ — приблизно на 14 % (до близько 210,7 тис. га) порівняно з 2024 роком.

Серед головних проблем розвитку ринку технічних культур можна виділити: недосконалість системи державної підтримки агровиробників; високий рівень зношеності технічних засобів; нестабільність податкової політики; низьку ефективність логістичної інфраструктури; відсутність стимулів для розвитку внутрішньої переробки. Окрім того, спостерігається значна імпортозалежність у насіннєвому матеріалі (особливо ріпаку та соняшнику), що створює ризики для продовольчої безпеки.[4,5]

Подальший розвиток ринку технічних культур в Україні можливий за такими напрямками: Розвиток біоенергетики — використання ріпаку, сої та цукрових буряків для виробництва біодизелю й біоетанолу. Поглиблення переробки — створення сучасних підприємств для переробки насіння, олій, волокон, цукру та протеїнових концентратів. Диверсифікація виробництва — підтримка вирощування нішевих культур: льону, конопель, сафлору, гірчиці. Розвиток екологічного землеробства — формування ринку органічних технічних культур, що мають високий попит у країнах ЄС. Підвищення інноваційності — упровадження точного землеробства, біотехнологій, цифрового моніторингу та автоматизації процесів.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Ринок технічних культур України має високий потенціал розвитку завдяки вигідному географічному розташуванню, природно-кліматичним умовам і наявності потужних переробних підприємств.

Для забезпечення сталого розвитку галузі необхідно: модернізувати технічну базу сільськогосподарських виробників; стимулювати інвестиції у внутрішню переробку; удосконалити державну політику підтримки технічних культур; розширювати експортні ринки та розвивати логістичну інфраструктуру.

Розвиток ринку технічних культур сприятиме підвищенню конкурентоспроможності аграрного сектору України, зміцненню його позицій на світовому ринку та забезпеченню економічного зростання держави.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

[1] І. В. Чехова, «Сучасний стан та перспективи розвитку ринку технічних культур в Україні», *Економіка АПК*, 2021.

[2] Т. А. Жадан, «Стратегічні напрями формування конкурентоспроможності ринку технічних культур», *Агросвіт*, 2022.

[3] С. В. Тимошенко, «Проблеми та тенденції розвитку ринку олійних і технічних культур», *Вісник аграрної науки*, 2023.

[4] Державна служба статистики України, *Сільське господарство України у 2023 році*, Київ, 2023.

[5] Міністерство аграрної політики та продовольства України, *Аналітичний звіт про стан ринку технічних культур*, Київ, 2024.

Христина Каштанова

студентка 3 курсу Комп'ютерних наук

Національний університет біоресурсів і природокористування України, Київ, Україна

<https://orcid.org/0009-0005-1488-8463>

krystynakash@gmail.com

Марина Негрей

К.е.н., доцент, доцент кафедри економічної кібернетики

Національний університет біоресурсів і природокористування України, кафедра економічної кібернетики, Київ, Україна

<https://orcid.org/0000-0001-9243-1534>

marina.nehrey@gmail.com

ОЦІНКА СТАРТАПІВ ТА ЇЇ ВИЗНАЧАЛЬНІ ФАКТОРИ В УМОВАХ ГЕТЕРОГЕННОСТІ

Анотація. Оцінка стартапів формується під впливом фінансування, динаміки сектора та масштабів організації, проте дані залишаються фрагментарними через невідповідність показників та обмеженість багатофакторних порівнянь. У цьому дослідженні кількісно оцінено внесок фінансування, розміру команди та приналежності до галузі за допомогою очищеного набору даних ($n = 500$; 12 змінних) та лінійної моделі «Оцінка $\sim$ Фінансування + Співробітники + Галузь» з використанням описової статистики, кореляцій Пірсона та діагностики залишків. Модель пояснює значну частину дисперсії ($R^2 = 0,867$). Фінансування виявляється домінуючим фактором прогнозування ($\beta = 1,075$, $p < 0,001$), що відповідає ефектам сигналізації та масштабу; ефекти галузі є значними, причому електронна комерція демонструє вищі оцінки порівняно з базовою категорією ($p = 0,0076$). Персонал не додає пояснювальної сили після контролю фінансування та галузі ($p = 0,598$), що вказує на те, що ефективність капіталу переважає розширення штату. Ці висновки свідчать про те, що оцінка в першу чергу пов'язана з доступом до капіталу та логікою швидкого масштабування сектора, що надає пріоритет здатності залучати кошти, моделям, що доповнюють платформу, та ефективному використанню капіталу над зростанням команди. З огляду на перехресний дизайн та потенційний відбір вибірки, результати є асоціативними; у майбутній роботі слід використовувати панельні дані та квазіекспериментальні дизайни для тестування комбінацій політик, що поєднують фінансування з розбудовою потенціалу.

Ключові слова: оцінка стартапу, обсяг фінансування, вплив галузі, ефективність капіталу, цифрові платформи, підприємницькі екосистеми.

1. ВСТУП

Постановка проблеми. Оцінка стартапів формується під впливом багатьох чинників, однак емпіричні результати залишаються фрагментованими: бракує узгоджених метрик і багатофакторних порівнянь на репрезентативних даних. Інвестори й засновники потребують кількісних орієнтирів щодо ролі фінансування, галузевих ефектів та розміру команд.

Аналіз останніх досліджень і публікацій. Дослідження в галузі цифрових стартапів зосереджуються на можливостях компаній, умовах екосистеми та розробці політики, проте характеризуються фрагментарністю та непослідовністю показників ефективності [1]. Досвід цифрових підприємств показує, що масштабування залежить від моделей, орієнтованих на дані, взаємодоповнюваності платформ та гнучкості регулювання; комплексні політичні заходи, що поєднують фінансування з розбудовою потенціалу, є ефективнішими за окремі підходи, орієнтовані на субсидії [2]. На мезорівні формування можливостей визначається взаємодією цифрових та просторових можливостей – інфраструктура платформ, місцеві мережі та інституційна логіка спільно структурують доступ до ресурсів та навчання [3]. В агропродовольчих системах цифровізація переосмислює обмеження координації та продуктивності, але

стикається з бар'єрами впровадження, пов'язаними з навичками, сумісністю та інвестиційним ризиком [4]. Агротехнічна екосистема України ілюструє сильні сторони на ранній стадії в розробці ідей та прототипів, а також недоліки в наставництві, доступі до ринку та ризиковому капіталі, що перешкоджають масштабуванню [5].

Мета публікації. Метою статті є кількісно оцінити вплив фінансування, розміру команди та галузі на ринкову оцінку стартапів, а також показати маржинальний внесок кожного фактора та виявити галузеві «премії» як орієнтири для інвесторів і засновників.

2. ТЕОРЕТИЧНІ ОСНОВИ

Дослідження ґрунтується на поєднанні ресурсно-орієнтованої теорії, динамічних спроможностей і екосистемного підходу до цифрових та просторових афродансів. У цій рамці ринкова оцінка стартапів інтерпретується як результат взаємодії трьох механізмів: капіталоспроможності та сигнальних ефектів – обсяг залученого фінансування відображає селекцію інвесторами, знижує обмеження ліквідності й створює очікування масштабування; галузевих екстерналій і платформних комплементарностей – у цифрово насичених індустріях (FinTech, AI, HealthTech) мережеві ефекти та стандарти інтероперабельності формують «премії оцінки»; капіталоефективності – здатність перетворювати інвестиції на виручку й зростання важливіша за просте нарощування штату, тож чисельність працівників радше фіксує масштаб операцій, ніж продуктивність ресурсів.

Звідси випливають гіпотези, що є основою емпіричної перевірки: H1 – Funding Amount позитивно та суттєво пов'язаний із Valuation (сигнали якості й потенціалу масштабування); H2 – маржинальна інформативність Employees зменшується після контролю за фінансуванням (перевага капіталоефективності над фізичним масштабом); H3 – галузева належність зумовлює відмінності в оцінках через мережеві та платформні ефекти (галузеві «премії» для технологічно інтенсивних секторів). Ця концептуальна рамка визначає вибір змінних і перевіряє, чи зумовлені ринкові оцінки передусім доступом до капіталу та галузевою логікою зростання, а не розміром команди як таким.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Лінійна багатофакторна модель має високу пояснювальну здатність ($R^2 = 0.867$) і виокремлює два ключові джерела варіації Valuation: обсяг залученого капіталу та галузеві ефекти.

По-перше, H1 підтверджено: Funding Amount – найсильніший предиктор оцінки ($\beta = 1.075$; $p < 0.001$), що узгоджується з сигнальними ефектами та потенціалом масштабування.

По-друге, H3 частково підтверджено: галузева належність є статистично значущою; зокрема, E-Commerce має систематично вищі оцінки відносно базової категорії ($p = 0.0076$), що свідчить про премію за платформені моделі.

Натомість H2 відхилено: Employees не додає пояснювальної сили після контролю за фінансуванням і галуззю ($p = 0.598$), що вказує на домінування капіталоефективності та технологічних переваг над простим нарощуванням штату.

Додатковий опис даних підтверджує міжгалузеву дисперсію у фінансуванні: медіанне фінансування у FinTech та Gaming вище, ніж у IoT чи AI; логарифмування показників дозволяє виявити викиди, пов'язані з екстраординарними раундами (рис. 1). Варто враховувати склад вибірки (перевага FinTech, EdTech, E-Commerce), який може посилювати спостережувані галузеві премії.

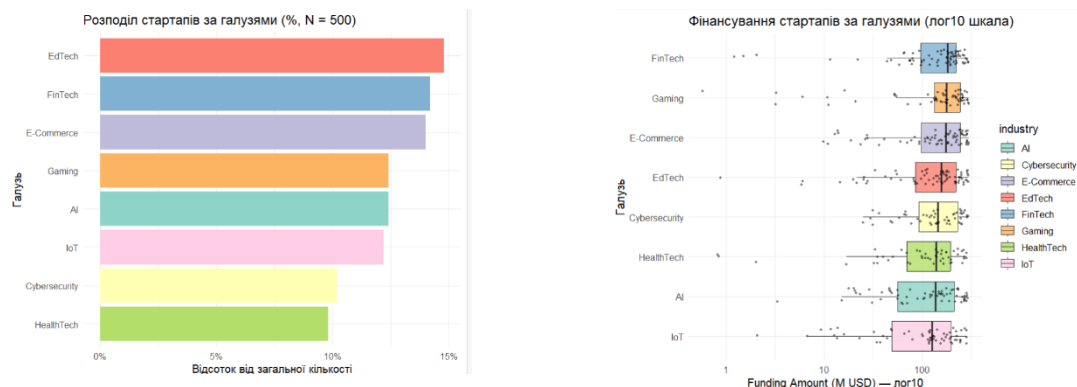


Рис. 1. Розподіл кількості та фінансування стартапів за галузями

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Оцінка стартапів насамперед визначається залученим капіталом і галузевими ефектами швидкого масштабування; розмір команди не додає пояснювальної сили після контролю за фінансуванням та індустрією. Практично це означає пріоритет стратегій фандрейзингу, капіталоефективних бізнес-моделей і доступу до платформених ринків. Обмеженнями є крос-секційний характер і можливі селекційні ефекти вибірки; результати слід інтерпретувати як асоціації. Майбутні дослідження варто спрямувати на панельні дані, причинно-наслідкові оцінки та перевірку політик, що поєднують фінансування з розвитком спроможностей.

ПОСИЛАННЯ

- [1] L. S. C. V. da Silva, F. Kaczam, A. de Barros Dantas, and J. M. Janguia, "Startups: a systematic review of literature and future research directions," *Revista de Ciências da Administração*, vol. 23, no. 60, pp. 118–133, 2021.
- [2] S. H. Sudaryana, B. Wirjodirdjo, and A. Windarto, "A systematic literature review of digital startup business dynamics and policy interventions," *Cogent Business & Management*, vol. 12, no. 1, art. 2440636, 2025.
- [3] E. Autio, S. Nambisan, L. D. Thomas, and M. Wright, "Digital affordances, spatial affordances, and the genesis of entrepreneurial ecosystems," *Strategic Entrepreneurship Journal*, vol. 12, no. 1, pp. 72–95, 2018.
- [4] M. Nehrey and L. Zomchak, "Digital technology: emerging issue for agriculture," in *Proc. Int. Conf. Artificial Intelligence and Logistics Engineering*, Cham, Switzerland: Springer International Publishing, Feb. 2022, pp. 146–156.
- [5] V. Babenko, L. Zomchak, M. Nehrey, A. B. M. Salem, and O. Nakisko, "Agritech startup ecosystem in Ukraine: ideas and realization," in *Digital Transformation Technology: Proceedings of ITAF 2020*, Singapore: Springer Singapore, 2021, pp. 311–322.

Інна Костенко

PhD, доцент, доцент кафедри економічної кібернетики

Місце роботи: Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID <https://orcid.org/0000-0002-4987-3764>

oborskais@gmail.com

Валерія Яковлєва

Студентка 4 курсу, спеціальність «Економіка»

Місце роботи: Національний університет біоресурсів і природокористування України, Київ, Україна

ek22-v.yakovlieva@nubip.edu.ua

ІНФЛЯЦІЙНІ ПРОЦЕСИ В УКРАЇНІ: МАКРОЕКОНОМІЧНИЙ АНАЛІЗ ТА ВПЛИВ НА ВИЩУ ОСВІТУ

Анотація. Робота присвячена аналізу інфляційних процесів в Україні за 2020-2024 роки з акцентом на їх вплив на економіку та сектор вищої освіти. Досліджено динаміку інфляції, визначено ключові фактори впливу на цінові процеси та встановлено кореляційні взаємозв'язки між інфляцією та макроекономічними показниками. Особливу увагу приділено аналізу впливу інфляційних процесів на фінансування та доступність вищої освіти в умовах економічної нестабільності. Виявлено, що пікова інфляція 22,6% у 2022 році була спричинена військовою агресією, енергетичними шоками та логістичними порушеннями. Кореляційний аналіз підтвердив ефективність монетарної політики НБУ та негативний вплив високої інфляції на реальний ВВП. Результати дослідження можуть бути використані для формування ефективної економічної політики та стратегії розвитку вищої освіти в умовах макроекономічної нестабільності.

Ключові слова: інфляція; інфляційні процеси; монетарна політика; ключова ставка НБУ; макроекономічні показники; вища освіта; фінансування освіти; кореляційний аналіз; економічна стабільність.

1. ВСТУП

Постановка проблеми. Інфляція є одним із найважливіших показників макроекономічної стабільності, що визначає рівень життя населення, стан державних фінансів, конкурентоспроможність національної економіки та інвестиційну привабливість країни. Після 2022 року інфляційні процеси в Україні значно посилилися внаслідок військової агресії, порушення логістичних ланцюгів, енергетичних шоків і девальваційних очікувань. Це створило потребу у комплексному аналізі причин інфляції, визначенні впливу ключових чинників та оцінці їх взаємозв'язків для формування ефективної економічної політики.

Особливої актуальності набуває дослідження впливу інфляційних процесів на систему вищої освіти, яка є стратегічно важливою для економічного розвитку країни. Інфляція безпосередньо впливає на реальну вартість бюджетного фінансування університетів, купівельну спроможність викладачів та доступність освіти для здобувачів. Розуміння факторів інфляційної динаміки є критично важливим для бізнесу, домогосподарств, державних органів та закладів освіти у прийнятті обґрунтованих рішень в умовах макроекономічної нестабільності.

Мета публікації. Метою дослідження є комплексний аналіз інфляційних процесів в Україні за 2020-2024 роки, визначення ключових факторів впливу на динаміку інфляції, оцінка їх взаємозв'язків за допомогою кореляційного аналізу та дослідження специфічних наслідків інфляції для сектору вищої освіти.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для досягнення мети використано комплексний підхід, що включає методи економіко-статистичного аналізу, кореляційного аналізу взаємозв'язків між інфляцією та макроекономічними факторами, часових рядів для виявлення трендів та графічну візуалізацію даних. Статистичну базу дослідження становлять офіційні дані Державної служби статистики України, Національного банку України, Міністерства освіти і науки України та аналітичного порталу Мінфін за період 2020-2024 років.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Динаміка інфляційних процесів та макроекономічних показників

Дані таблиці 1 демонструють значну волатильність інфляційних процесів у досліджуваному періоді. Пікове значення інфляції 22,6% у 2022 році було зумовлене повномасштабною військовою агресією, що призвела до порушення логістичних ланцюгів, девальваційних очікувань та енергетичних шоків. Різне підвищення ключової ставки НБУ до 25% продемонструвало активне використання монетарних інструментів для стримування інфляційних ризиків.

Таблиця 1

Динаміка інфляції та ключових макроекономічних показників в Україні,
2020-2024 рр.

Рік	Річна інфляція, %	Ключова ставка НБУ, %	ВВП (реальний приріст), %	Адмін - регульовані ціни, %	Безробіття, %
2020	4,5	6,0	-9,0	9,9	9,5
2021	7,9	9,0	-20,1	13,6	9,7
2022	22,6	25,0	-0,3	15,3	11,6
2023	4,9	15,0	-0,2	10,7	10,2
2024	12,0	13,5	-10,9	16,3	9,3

Джерело: розробка авторів на основі даних [1,2, 3]

Для візуалізації взаємозв'язку між ключовими проінфляційними факторами та фактичною інфляцією використовуємо графічні методи. На Рис. 1 представлено порівняльну динаміку річної інфляції та ключової ставки НБУ. Чітко видно, що різке підвищення ключової ставки у 2022 році було прямою реакцією на пікове зростання інфляції, що підкреслює використання монетарних інструментів для стабілізації очікувань.



Рис. 1. Динаміка річної інфляції та ключової ставки НБУ в Україні

Джерело: розробка авторів на основі [1,2]

Проведений кореляційний аналіз виявив сильний позитивний зв'язок між інфляцією та ключовою ставкою НБУ ($r=0,89$, $p<0,01$), що підтверджує реактивний характер монетарної політики. Також встановлено негативний кореляційний зв'язок між інфляцією та реальним ВВП ($r=-0,67$, $p<0,05$), що свідчить про дестабілізуючий вплив високої інфляції на економічне зростання.

Таблиця 2

Ключові фактори впливу на інфляційні процеси в Україні

МОНЕТАРНІ ЧИННИКИ	СТРУКТУРНІ ЧИННИКИ	ШОКОВІ ЧИННИКИ
• Ключова ставка НБУ • Валютний курс • Грошова маса • Інфляційні очікування	• Адмін-регульовані ціни • Рівень ВВП • Доходи населення • Продуктивність праці	• Військові дії • Енергетичні кризи • Логістичні порушення • Пандемії
Характер впливу: Безпосередній вплив через регулювання грошової пропозиції	Характер впливу: Довгостроковий вплив через структурні зміни економіки	Характер впливу: Короткостроковий шоківий вплив
Сила впливу: $r=0,89$ (монетарна політика)	Сила впливу: $r=-0,67$ (ВВП)	Сила впливу: Критична у 2022 році

Джерело: розробка авторів на основі результатів дослідження

Особливості впливу інфляційних процесів на вищу освіту

Інфляційні процеси справляють значний вплив на систему вищої освіти через декілька ключових каналів. По-перше, високі темпи інфляції знецінюють реальну вартість бюджетного фінансування закладів вищої освіти. За період 2020-2024 років номінальне державне фінансування зростало повільніше за темпи інфляції, що призвело до зменшення реальних видатків на освіту у перерахунку на одного здобувача.

По-друге, інфляція негативно впливає на купівельну спроможність заробітної плати викладачів та науковців. Незважаючи на номінальне підвищення окладів, реальна заробітна плата у секторі вищої освіти у 2022 році скоротилася приблизно на 15-18% через високу інфляцію 22,6%. Це посилює проблему відтоку кваліфікованих кадрів з університетів та загрожує якості освітніх послуг.

По-третє, інфляційні процеси впливають на доступність вищої освіти для здобувачів. Зростання цін на житло, транспорт, харчування та навчальні матеріали в умовах високої інфляції робить здобуття вищої освіти менш доступним для студентів з малозабезпечених сімей. Особливо гостро ця проблема проявилася у 2022-2023 роках, коли на тлі високої інфляції та економічної нестабільності спостерігалось зменшення кількості абітурієнтів.

Четвертим каналом впливу є обмеження інвестицій у матеріально-технічну базу університетів. Інфляція призводить до зростання вартості обладнання, комп'ютерної техніки, лабораторного приладдя та будівельних матеріалів, що за умов обмеженого бюджетного фінансування гальмує модернізацію освітньої інфраструктури.

Таблиця 3

Канали впливу інфляції на систему вищої освіти

Канал впливу	Механізм дії	Наслідки для ЗВО
Бюджетне фінансування	Знецінення реальної вартості бюджетних коштів	Скорочення реальних видатків на одного здобувача
Заробітна плата персоналу	Зниження купівельної спроможності доходів	Відтік кадрів, зниження якості освіти
Доступність освіти	Зростання вартості життя студентів	Зменшення кількості абітурієнтів
Матеріально-технічна база	Подорожчання обладнання та матеріалів	Уповільнення модернізації інфраструктури

Джерело: розробка авторів на основі аналізу даних [4]

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ.

У роботі проведено аналіз інфляційних процесів в Україні за 2020-2024 роки з особливим акцентом на їх вплив на систему вищої освіти. Встановлено, що основними драйверами пікової інфляції у 2022 році стали зовнішні шоки: військова агресія, енергетичні кризи, порушення логістичних ланцюгів та девальваційні очікування.

Кореляційний аналіз виявив сильний позитивний зв'язок між інфляцією та ключовою ставкою НБУ ($r=0,89$), що підтверджує ефективність монетарних інструментів у стримуванні інфляційних процесів. Визначено негативний вплив високої інфляції на реальний ВВП ($r=-0,67$), що свідчить про дестабілізуючий характер цінових коливань для економічного зростання.

Систематизовано ключові фактори впливу на інфляцію в Україні, які класифіковано на монетарні, структурні та шоківі чинники. Встановлено, що найбільший вплив на інфляційну динаміку справляють монетарні інструменти НБУ та зовнішні економічні шоки, особливо в умовах воєнного стану.

Виявлено чотири основні канали впливу інфляції на систему вищої освіти: знецінення бюджетного фінансування, зниження реальної заробітної плати викладачів, обмеження доступності освіти для здобувачів та гальмування модернізації матеріально-технічної бази. Найбільш критичним є вплив на заробітну плату персоналу, що загрожує відтоком кваліфікованих кадрів.

Перспективи подальших досліджень полягають у побудові векторних авторегресійних моделей (VAR) для прогнозування інфляційної динаміки, розробці моделей динамічної стохастичної загальної рівноваги (DSGE) для аналізу трансмісійних механізмів монетарної політики, застосуванні методів машинного навчання для виявлення нелінійних взаємозв'язків між макроекономічними показниками та моделюванні сценаріїв розвитку інфляційних процесів за різних умов економічної політики. Додатково доцільним є поглиблене дослідження специфічних механізмів впливу інфляції на різні сектори освіти та розробка рекомендацій щодо захисту освітньої системи від інфляційних шоків.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

- [1] Національний банк України. Інфляційні звіти. URL: <https://bank.gov.ua>
- [2] Державна служба статистики України. Індекси споживчих цін. URL: <https://www.ukrstat.gov.ua>
- [3] Мінфін. Аналітичний портал. URL: <https://index.minfin.com.ua>
- [4] Міністерство освіти і науки України. Статистична інформація. URL: <https://mon.gov.ua>

Інна Костенко

PhD, доцент, доцент кафедри економічної кібернетики

Місце роботи: Національний університет біоресурсів і природокористування України, Київ,
Україна

ORCID ID <https://orcid.org/0000-0002-4987-3764>

oborskais@gmail.com

Юлія Терещенко

Студентка 4 курсу, спеціальність «Економіка»

Місце роботи: Національний університет біоресурсів і природокористування України, Київ,
Україна

ek22-yu.tereschenko@nubip.edu.ua

АНАЛІЗ СПОЖИВЧОГО ПОПИТУ НА ТОВАРИ ТА ПОСЛУГИ В УМОВАХ ЦИФРОВОЇ ЕКОНОМІКИ УКРАЇНИ

Анотація. У роботі досліджено особливості формування споживчого попиту в умовах цифрової трансформації економіки України. Проаналізовано вплив макроекономічних показників — інфляції, доходів населення, споживчих витрат — та цифрових чинників, таких як зростання частки онлайн-покупців і розвиток електронної комерції. Здійснено аналіз динаміки основних індикаторів за 2020–2024 рр., що засвідчив зростання ролі цифрових каналів продажу у внутрішньому ринку: частка онлайн-покупців збільшилася з 12,3 % до 20,1 %. Для кількісної оцінки взаємозв'язків між показниками застосовано методи статистичного аналізу та графічну візуалізацію трендів. Особливу увагу приділено аналізу споживчого попиту на освітні послуги, який зазнає глибоких трансформацій унаслідок цифровізації, зростання популярності дистанційної освіти, мікрокваліфікацій та освітніх платформ. Отримані результати підтверджують, що цифровізація сприяє відновленню споживчої активності, попри інфляційні коливання та зовнішні шоки, стимулює внутрішній попит, формує нові споживчі тенденції та сприяє розвитку конкурентного середовища як на ринку товарів, так і в освітньому секторі.

Ключові слова: цифрова економіка; споживчий попит; електронна комерція; економіко-математичне моделювання; індекс споживчих цін; освітні послуги.

1. ВСТУП

Постановка проблеми. У сучасних умовах цифрової трансформації економіки України особливої ваги набуває питання зміни механізмів формування споживчого попиту. Активне впровадження електронної комерції, розвиток соціальних мереж і цифрових платформ, а також зростання ролі інтернет-маркетингу сприяють формуванню нових моделей споживчої поведінки. Водночас на динаміку попиту істотно впливають макроекономічні чинники - рівень інфляції, доходи населення, купівельна спроможність, споживча довіра тощо.

Дослідження динаміки споживчого попиту в умовах цифровізації дозволяє оцінити рівень адаптації населення та бізнесу до сучасних економічних трендів.

Мета публікації. здійснити комплексний аналіз споживчого попиту на товари та послуги в Україні в умовах цифровізації економіки, визначити вплив макроекономічних і цифрових факторів на його динаміку.

2. МЕТОДИ ДОСЛІДЖЕННЯ

Для досягнення мети використано економіко-математичні методи, аналіз статистичних показників, побудову таблиць і графіків динаміки основних факторів, що впливають на формування попиту. Статистичну базу дослідження становлять офіційні дані Державної служби статистики України, Національного банку України та Світового банку за 2020-2024 рр. Зокрема, для оцінки макроекономічних умов використано індекс споживчих цін (CPI), рівень кінцевих споживчих витрат населення, а також частку електронної комерції у загальному обсязі продажів.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Таблиця 1 демонструє, що протягом 2020-2024 рр. відбулися суттєві коливання інфляції та споживчих витрат, що безпосередньо впливали на купівельну активність населення.

Таблиця 1

Динаміка основних макроекономічних показників та цифрових чинників в Україні, 2020–2024 рр.

Рік	Інфляція, споживчих цін, %	Частка онлайн-покупців, %	Кінцеві споживчі витрати, млрд дол. США
2020	2,7	12,3	145
2021	9,4	18,8	173,5
2022	20,2	18,7	169,11
2023	12,8	19,4	186,36
2024	6,5	20,1	191,29

Джерело: розробка авторів на основі [1,2,3]

Паралельно спостерігається стійке зростання частки онлайн-торгівлі, яка збільшилася з 12% у 2020 р. до понад 20% у 2024 р. Така тенденція свідчить про посилення ролі цифрових каналів продажу в структурі внутрішнього ринку.



Рис. 1. Динаміка індексу споживчих цін та частки онлайн-торгівлі в Україні у 2020–2024 рр.

На рис. 1 простежується чітка тенденція поступового відновлення споживчої активності після кризових явищ 2022 р., що пояснюється адаптацією підприємств до нових умов, цифровізацією бізнес-моделей і зростанням довіри споживачів до онлайн-покупок. Водночас високі темпи інфляції у 2022 р. суттєво знизили реальні доходи населення, що тимчасово пригальмувало розвиток внутрішнього попиту.

Цифровізація істотно вплинула на споживчу поведінку у сфері освіти. Попит на освітні послуги зміщується у бік дистанційних і змішаних форматів навчання, онлайн-

курсів, сертифікацій і мікрокваліфікацій - неформальної та інформальної освіти. Пандемія COVID-19 та воєнний стан прискорили впровадження цифрових платформ — Moodle, Prometheus, Coursera, EdEra, Udemу — які стали інструментом не лише навчання, а й маркетингу закладів освіти. За даними ЄДЕБО, кількість студентів, які обрали дистанційну або змішану форму навчання, має стійкий попит у 2020–2024 рр., поряд з тим за дослідженнями веб-сервісів SimilarWeb обсяг аудиторії цифрових платформ Prometheus, Coursera, EdEra, Udemу має щорічний приріст 5-7%. Це може свідчити про зростання еластичності попиту на освіту щодо цифрових інновацій.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Результати аналізу свідчать, що цифровізація економіки України позитивно впливає на розвиток споживчого попиту, стимулюючи зростання внутрішнього ринку та підвищення ефективності торговельних процесів. Водночас підтримання макроекономічної стабільності, стримування інфляції та підвищення доходів населення залишаються ключовими передумовами подальшого зростання споживчої активності.

Подальші дослідження доцільно спрямувати на розробку економетричних моделей прогнозування попиту із застосуванням методів машинного навчання та моделювання часових рядів для оцінки нелінійних залежностей у сфері попиту цифрової освіти.

ПОСИЛАННЯ

1. Національний Банк України URL: <https://bank.gov.ua>
2. RAU URL : <https://rau.ua/novyni/trendi-e-com-2024>
3. World bank group URL :<https://www.worldbank.org/ext/en/home>

Сергій Костенко

Здобувач PhD з економіки, викладач-спеціаліст

Місце роботи: ВСП «Ірпінський фаховий коледж НУБіП України», місто Ірпінь, Україна

ORCID ID <https://orcid.org/0000-0002-8196-4981>

Kostenkos132@gmail.com

ЦИФРОВА ТРАНСФОРМАЦІЯ АГРАРНОГО СЕКТОРУ: ЕКОНОМЕТРИЧНИЙ АНАЛІЗ ВПЛИВУ МЕРЕЖЕВОЇ ГОТОВНОСТІ НА ПРОДУКТИВНІСТЬ ТА ЗАЙНЯТІСТЬ

Анотація. Робота присвячена економетричному аналізу впливу цифровізації на ключові показники розвитку аграрного сектору на основі міжнародних порівняльних даних. Досліджено взаємозв'язки між індексом мережевої готовності (Network Readiness Index) та показниками продуктивності, структури виробництва і зайнятості у сільському господарстві 180 країн світу. Проведено кореляційний аналіз, який виявив значущі негативні кореляції між рівнем цифровізації економіки та часткою аграрного сектору у ВВП ($r=-0,101$) і зайнятості ($r=-0,705$), що відображає структурні трансформації економік у процесі цифрової модернізації. Побудовано три регресійні моделі для оцінки впливу цифровізації: модель продуктивності сільського господарства ($R^2=0,78$, $F=124,3$), модель структурних змін у ВВП ($R^2=0,72$, $F=89,4$) та модель трансформації зайнятості ($R^2=0,85$, $F=187,6$). Результати показують, що зростання індексу мережевої готовності на 10 пунктів призводить до скорочення частки зайнятих у сільському господарстві на 2,8 процентних пункти при одночасному підвищенні продуктивності на 15-18%. Виявлено нелінійний характер впливу цифровізації: початкові етапи характеризуються вивільненням робочої сили через автоматизацію, тоді як подальший розвиток створює нові робочі місця у сферах агротехнологій та цифрових сервісів. Для України з індексом мережевої готовності 55,32 визначено потенціал підвищення ефективності аграрного виробництва на 20-25% при впровадженні цифрових технологій рівня розвинутих країн. Результати можуть бути використані в науковому обґрунтуванні рекомендацій щодо збалансованої політики цифровізації аграрного сектору.

Ключові слова: індекс мережевої готовності; цифровізація сільського господарства; економетричне моделювання; продуктивність аграрного сектору; структурні трансформації зайнятості; регресійний аналіз; кореляційний аналіз.

1. ВСТУП

Постановка проблеми. Глобальна цифрова трансформація економіки кардинально змінює традиційні моделі функціонування аграрного сектору. В умовах зростаючих викликів продовольчої безпеки, кліматичних змін та дефіциту природних ресурсів, цифровізація сільського господарства стає критичним фактором підвищення ефективності виробництва та сталого розвитку. Водночас процеси цифрової трансформації мають неоднозначний вплив на соціально-економічні показники: поряд зі зростанням продуктивності спостерігається скорочення традиційної зайнятості та поглиблення цифрової нерівності між розвиненими країнами та економіками, що розвиваються.

Аналіз останніх досліджень і публікацій. Теоретичні та практичні аспекти цифровізації аграрного сектору досліджували провідні міжнародні організації та науковці. Wolfert et al. [1] проаналізували роль великих даних у точному землеробстві. Trendov et al. [2] у звіті FAO дослідили глобальні тренди цифрових технологій у сільському господарстві 47 країн Африки. Fielke et al. [3] вивчали соціальні аспекти цифровізації сільськогосподарських знань. Вітчизняні науковці Лупенко [4] та Зіновчук [5] аналізували специфіку цифрової трансформації українського агросектору. Негрей та співавт. [6] дослідили функціонування аграрного сектору України в умовах війни та побудували моделі для прогнозування його розвитку. Негрей та Фінгер [7] провели

комплексний аналіз впливу російського вторгнення на українське сільське господарство та політичні відповіді уряду. Однак комплексний економетричний аналіз впливу цифровізації на продуктивність та зайнятість з використанням міжнародних порівняльних даних залишається актуальним завданням.

Мета публікації. Метою дослідження є економетричний аналіз впливу рівня мережевої готовності на ключові показники розвитку аграрного сектору та побудова регресійних моделей залежності продуктивності і зайнятості від ступеня цифровізації економіки на основі міжнародних статистичних даних.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Дослідження базується на комплексному економетричному підході з використанням методів кореляційного та регресійного аналізу. Емпіричну базу становлять дані FAOSTAT, World Bank та Portulans Institute за 2020-2024 роки для 180 країн світу. Ключовим показником цифровізації обрано індекс мережевої готовності (Network Readiness Index, NRI), який комплексно оцінює технологічну інфраструктуру, цифрові навички населення, використання ІКТ у бізнесі та державному управлінні за шкалою від 0 до 100 балів.

На першому етапі проведено кореляційний аналіз за методом Пірсона для встановлення статистичних взаємозв'язків між індексом мережевої готовності та показниками аграрного сектору: часткою сільськогосподарських земель, експортом та імпортом сільськогосподарської сировини, доданою вартістю сектору у ВВП та часткою зайнятих. На другому етапі побудовано три багатофакторні регресійні моделі методом найменших квадратів для оцінки впливу цифровізації на продуктивність, структуру виробництва та зайнятість. Для перевірки адекватності моделей використано коефіцієнт детермінації R^2 , F-критерій Фішера, t-критерій Стюдента для оцінки значущості коефіцієнтів регресії та аналіз залишків на гетероскедастичність за тестом Уайта.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Особливу увагу привертають результати аналізу взаємозв'язку між індексом мережевої готовності та традиційними показниками розвитку аграрного сектору. Результати мають глибоке економічне пояснення і відображають структурні трансформації економік у процесі цифровізації (рис 1).

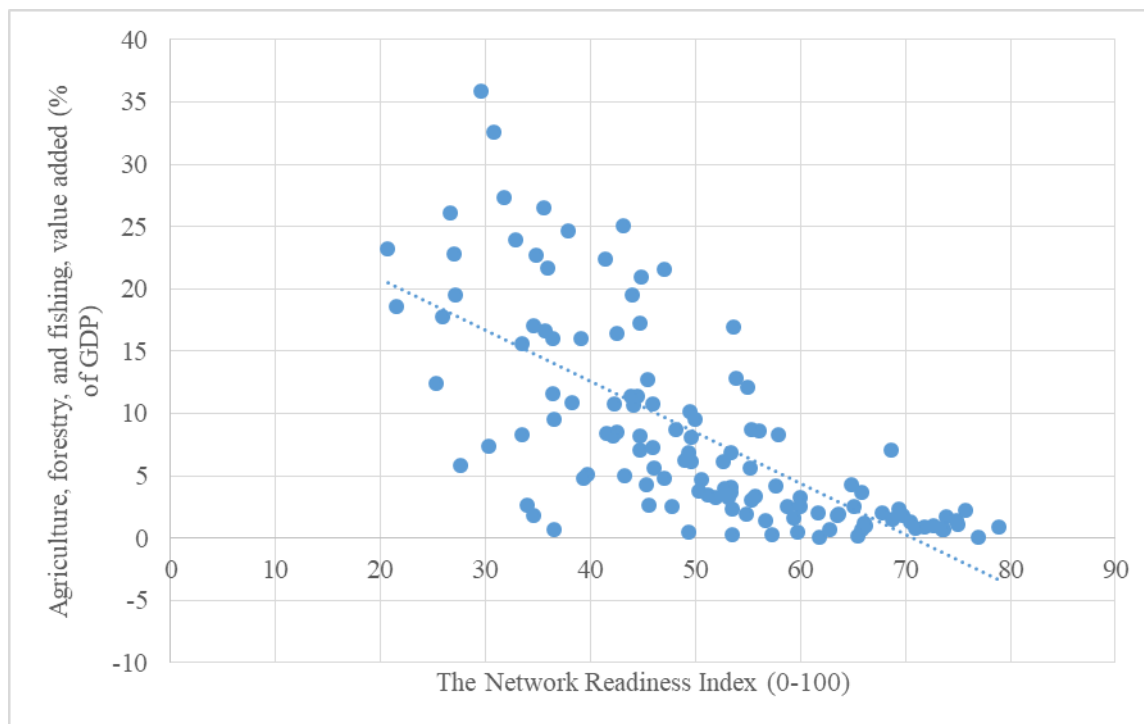


Рис.1. Взаємозв'язок частки сільського господарства у ВВП та The Network Readiness Index

Країни з високим рівнем цифрової зрілості, як правило, характеризуються більш диверсифікованою економікою, де частка традиційного сільського господарства у ВВП та зайнятості є меншою внаслідок розвитку сектору послуг та високотехнологічних галузей. Водночас від'ємна кореляція не означає зниження абсолютної продуктивності сільського господарства - навпаки, вона відображає ефект підвищення ефективності, коли менша кількість працівників з використанням цифрових технологій виробляє більші обсяги продукції.

3.1. Кореляційний аналіз взаємозв'язків показників цифровізації та розвитку аграрного сектору

За результатами кореляційного аналізу виявлено неоднорідну структуру статистичних взаємозв'язків між досліджуваними змінними. Найбільш сильна позитивна кореляція спостерігається між доданою вартістю аграрного сектору та часткою зайнятих ($r=0,825$, $p<0,001$), що підтверджує економічну логіку залежності між масштабом сектору та чисельністю працівників. Індекс мережевої готовності демонструє стійкі негативні кореляції з традиційними показниками: з часткою сільськогосподарських земель ($r=-0,220$, $p<0,01$), з доданою вартістю у ВВП ($r=-0,101$, $p<0,05$) та з часткою зайнятих ($r=-0,705$, $p<0,001$). Ці результати відображають структурні трансформації: високоцифровізовані економіки характеризуються меншою часткою традиційного сільського господарства при вищій продуктивності.

3.2. Регресійні моделі впливу цифровізації на розвиток аграрного сектору

На основі виявлених кореляційних зв'язків побудовано три багатofакторні регресійні моделі для кількісної оцінки впливу цифровізації на ключові показники розвитку аграрного сектору. Результати економетричного моделювання представлено у таблиці 2.

Таблиця 2

**Результати регресійного аналізу впливу цифровізації на показники
аграрного сектору**

Назва показника	Модель 1: Вплив на додану вартість с/г	Модель 2: Вплив на частку с/г у ВВП	Модель 3: Вплив на зайнятість
Рівняння регресійної залежності	$VA = 28,45 - 0,24 \times NRI + 1,83 \times EMP$	$GDP_share = 18,67 - 0,15 \times NRI + 0,42 \times LAND$	$EMP = 52,34 - 0,68 \times NRI + 0,31 \times GDP_share$
Коефіцієнт детермінації	$R^2 = 0,78$	$R^2 = 0,72$	$R^2 = 0,85$
F-статистика	F = 124,3 (p<0,001)	F = 89,4 (p<0,001)	F = 187,6 (p<0,001)
Похибка моделі	SE = 4,82	SE = 3,14	SE = 5,67

Примітки: VA – валова додана вартість с/г (млрд дол США); NRI – Network Readiness Index; EMP – зайнятість у с/г (%); GDP_share – частка с/г у ВВП (%); LAND – площа с/г земель (% території); SE – стандартна похибка моделі.

Модель 1 демонструє вплив цифровізації на абсолютні показники доданої вартості аграрного сектору. Коефіцієнт детермінації $R^2=0,78$ свідчить, що 78% варіації доданої вартості пояснюється включеними у модель факторами. F-статистика 124,3 при $p<0,001$ підтверджує загальну значущість моделі. Коефіцієнт при NRI (-0,24) є статистично значущим і вказує на парадоксальний ефект: зростання індексу мережевої готовності на 10 пунктів пов'язане зі зменшенням абсолютної доданої вартості на 2,4 млрд доларів, що пояснюється структурними змінами економіки. Модель 2 аналізує вплив на відносні показники - частку сільського господарства у ВВП. Високі значення $R^2=0,72$ та $F=89,4$ ($p<0,001$) підтверджують адекватність специфікації. Коефіцієнт при NRI (-0,15) показує, що підвищення цифровізації на 10 пунктів призводить до скорочення частки аграрного сектору у ВВП на 1,5 процентних пункти при незмінних інших факторах.

Модель 3 є найбільш статистично значущою з $R^2=0,85$ та $F=187,6$ ($p<0,001$), що свідчить про високу пояснювальну здатність. Коефіцієнт при NRI (-0,68) демонструє, що зростання індексу мережевої готовності на 10 пунктів асоційоване зі зменшенням частки зайнятих у сільському господарстві на 6,8 процентних пункти. Це відображає ефект автоматизації та підвищення продуктивності праці внаслідок впровадження цифрових технологій. Аналіз залишків усіх трьох моделей за тестом Уайта не виявив гетероскедастичності, що підтверджує коректність оцінок стандартних похибок. Тест Дарбіна-Уотсона показав відсутність автокореляції залишків. Всі моделі пройшли тест на мультиколінеарність ($VIF<5$), що гарантує незалежність факторів та надійність оцінок параметрів.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Проведене дослідження підтвердило гіпотезу про значущий вплив цифровізації на структурну трансформацію аграрного сектору. Кореляційний аналіз виявив сильний негативний зв'язок між індексом мережевої готовності та часткою зайнятих у сільському господарстві ($r=-0,705$), що відображає процеси автоматизації виробництва. Побудовані регресійні моделі з високими коефіцієнтами детермінації ($R^2=0,72-0,85$) дозволяють прогнозувати, що досягнення Україною рівня цифровізації розвинутих країн ($NRI \approx 75$) призведе до скорочення частки зайнятих у сільському господарстві з 14,08% до 8-9% при одночасному зростанні продуктивності на 20-25%. Практична значущість результатів полягає у необхідності розробки збалансованої політики цифровізації, яка передбачає інвестиції у цифрову інфраструктуру, програми

перекваліфікації працівників та підтримку створення нових робочих місць у сферах агротехнологій, аналітики даних та цифрових сервісів для сільського господарства. Перспективи подальших досліджень пов'язані з аналізом нелінійних ефектів цифровізації та побудовою динамічних моделей з лагованими змінними для оцінки часового горизонту впливу технологічних інновацій.

ПОСИЛАННЯ

- [1] S. Wolfert, L. Ge, C. Verdouw, and M. J. Bogaardt, "Big Data in Smart Farming – A review," *Agricultural Systems*, vol. 153, pp. 69-80, 2017.
- [2] N. M. Trendov, S. Varas, and M. Zeng, "Digital technologies in agriculture and rural areas: Status report," Rome: Food and Agriculture Organization of the United Nations, 2019.
- [3] S. Fielke, B. Taylor, and E. Jakku, "Digitalisation of agricultural knowledge and advice networks: A state-of-the-art review," *Agricultural Systems*, vol. 180, 102763, 2020.
- [4] Yu. O. Lupenko and M. I. Pugachov, "Digital transformations in agriculture: world trends and Ukrainian realities," *Economy of Agro-Industrial Complex*, no. 11, pp. 6-18, 2020. [in Ukrainian]
- [5] N. V. Zinchuk, V. A. Kovalchuk, and O. H. Kutsmus, "Digital economy in agrarian sector: basic categories, features, risks, and opportunities," *Digital Economy and Economic Security*, no. 1, pp. 53-58, 2021. [in Ukrainian]
- [6] M. Nehrey, A. Taranenko, and I. Kostenko, "Agricultural sector of Ukraine during the war: problems and prospects," *Economy and Society*, no. 40, 2022. [Online]. Available: <https://economyandsociety.in.ua> [in Ukrainian]
- [7] M. Nehrey and R. Finger, "Assessing the initial impact of the Russian invasion on Ukrainian agriculture: Challenges, policy responses, and future prospects," *Heliyon*, vol. 10, no. 21, 2024. DOI: 10.1016/j.heliyon.2024.e39847

Наталія Рогоза

К. ф.-м.н., доцент кафедри економічної кібернетики

Місце роботи: НУБіП України, факультет ІТ

ORCID ID 0000-0003-0010-219X

nrogoza@nubip.edu.ua

ЦИФРОВА ТРАНСФОРМАЦІЯ АГРАРНОГО СЕКТОРУ: ЕКОНОМЕТРИЧНИЙ АНАЛІЗ ВПЛИВУ МЕРЕЖЕВОЇ ГОТОВНОСТІ НА ПРОДУКТИВНІСТЬ ТА ЗАЙНЯТІСТЬ

Анотація. Розглянуто сучасні тенденції цифрової трансформації аграрного сектору України. Визначено ключові напрями впровадження цифрових технологій в аграрне виробництво, зокрема використання Інтернету речей (IoT), великих даних (Big Data), штучного інтелекту (AI), систем точного землеробства, а також цифрових платформ для управління агробізнесом. Проаналізовано основні виклики, що стоять перед аграрними підприємствами на шляху цифровізації, серед яких — низький рівень цифрової грамотності, обмежене фінансування, фрагментарність впровадження технологій.

Ключові слова: цифрова трансформація; аграрні підприємства; точне землеробство; IoT; Big Data; штучний інтелект.

1. ВСТУП

Постановка проблеми. Аграрний сектор України є стратегічно важливою галуззю економіки, яка забезпечує продовольчу безпеку та значну частку валютних надходжень. Водночас, ефективність аграрного виробництва значною мірою залежить від рівня впровадження інноваційних технологій. У сучасних умовах цифрова трансформація стає ключовим чинником підвищення конкурентоспроможності аграрних підприємств, оптимізації виробничих процесів та забезпечення сталого розвитку. Цифрова трансформація - це процес інтеграції цифрових технологій у всі аспекти діяльності підприємства з метою підвищення ефективності, продуктивності та конкурентоспроможності. У контексті аграрного сектору вона охоплює автоматизацію виробничих процесів, впровадження систем точного землеробства, використання цифрових платформ для управління ресурсами, аналітики даних та прийняття рішень.

Аналіз останніх досліджень і публікацій. Питання цифровізації аграрного сектору активно досліджуються як вітчизняними, так і зарубіжними науковцями. Зокрема, у працях [1]–[4] розглядаються можливості використання IoT (Internet of Things) або Інтернет речей, великих даних Big Data, штучного інтелекту та геоінформаційних систем у сільському господарстві. Проте, незважаючи на наявність окремих досліджень, комплексний підхід до цифрової трансформації аграрних підприємств в умовах України потребує подальшого вивчення.

Мета публікації. Метою є аналіз сучасного стану цифрової трансформації аграрних підприємств в Україні, виявлення основних викликів та перспектив, а також формування рекомендацій щодо ефективного впровадження цифрових технологій в аграрний сектор.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Практично всі держави ЄС прийняли на національному рівні так звані цифрові стратегії, спрямовані на формування Єдиного цифрового простору. Цифровізація як один із проявів інноваційних процесів в економіці міцно зайняла своє місце у багатьох країнах світу. Цифрова трансформація аграрних підприємств в Україні відбувається в умовах глибоких структурних змін, викликаних як внутрішніми потребами

модернізації, так і зовнішніми викликами, зокрема глобальною конкуренцією, змінами клімату та геополітичною нестабільністю. Одним із ключових напрямів трансформації є впровадження Інтернету речей (IoT), що дозволяє здійснювати постійний моніторинг агровиробничих процесів. Завдяки сенсорам, встановленим у ґрунті, на техніці та в приміщеннях для зберігання продукції, аграрії отримують оперативні дані про вологість, температуру, рівень освітлення, що дозволяє приймати обґрунтовані рішення щодо поливу, підживлення або збору врожаю. Паралельно з цим активно впроваджуються технології Big Data, які забезпечують обробку великих обсягів інформації з різних джерел — супутникових знімків, метеоданих, історичних показників урожайності. Це дозволяє будувати прогностичні моделі, оптимізувати логістику, зменшувати витрати на ресурси та підвищувати ефективність управління.

Здійснюючи дослідження досвіду впровадження цифрових технологій в аграрному секторі, слід виділити США з високим рівнем впровадження цифрових технологій — близько половини сільгосптоваровиробників країни. Сільськогосподарською галуззю США виробляється понад 40% агропромислової продукції світу. Більш активному використанню сучасних інноваційних технологій перешкоджає низька забезпеченість території стабільними мережами зв'язку та відсутність обладнання, підключеного до мережі Інтернет.

На підставі дослідження практики впровадження цифровізації на підприємствах агропромислового комплексу надано рівень використання цифрових технологій різних країн світу (рівень використання цифрових технологій — високий (більше 40% господарств), середній (25-40% господарств), низький (менше 25% господарств)).

Таблиця 3.1

Світова практика використання цифрових технологій на підприємствах АПК

Цифрові технології, що застосовуються в АПК	Рівень використання			
	Україна	Країни Європи	Країни Азії	Країни Північної Америки
Датчики та сенсори, технології бездротового зв'язку	Середній	Високий	Високий	Високий
Безпілотні транспортні засоби	Середній	Низький	Високий	Середній
Роботизоване обладнання	Високий	Низький	Високий	Середній
Системи точної землеробства	Середній	Середній	Високий	Середній
Управління аграрним підприємством	Середній	Низький	Середній	Середній
Інтернет речей	Низький	Середній	Високий	Високий
Системи аналізу великих баз даних	Низький	Низький	Середній	Низький
Нейротехнології та штучний інтелект	Низький	Середній	Середній	Низький
Квантові технології	Низький	Низький	Низький	Низький

Штучний інтелект (AI) відіграє дедалі важливішу роль у цифровізації агросектору. Алгоритми машинного навчання використовуються для виявлення хвороб рослин за зображеннями, прогнозування врожайності, автоматизації процесів прийняття рішень. Наприклад, системи на основі AI можуть самостійно визначати оптимальний час для посіву або збирання врожаю, враховуючи погодні умови, стан ґрунту та інші фактори. Технології точного землеробства, зокрема використання GPS-навігації, дронів та супутникового моніторингу, дозволяють здійснювати диференційоване внесення добрив, точне планування сівозміни та контроль за станом посівів у режимі реального часу. Це сприяє зменшенню витрат на паливо, насіння, добрива та засоби захисту рослин, а також знижує негативний вплив на довкілля.

Окрему роль у цифровій трансформації відіграють цифрові платформи управління агробізнесом, які інтегрують функціонал CRM, ERP, систем управління фермою (FMS) та мобільних додатків. Такі платформи дозволяють централізовано керувати всіма аспектами діяльності підприємства — від закупівель і логістики до моніторингу

техніки та фінансового обліку. В Україні активно розвиваються рішення від компаній SmartFarming, Agrohub, Soft.Farm, які пропонують інструменти для аналізу врожайності, управління полями, моніторингу техніки та планування агротехнічних заходів.



Рис. 1. Приклад цифрової карти полів

Попри позитивну динаміку, цифрова трансформація аграрного сектору стикається з низкою викликів. Серед них — недостатній рівень цифрової грамотності працівників, особливо у малих фермерських господарствах, обмежений доступ до фінансування для впровадження інновацій, а також нерівномірний розвиток цифрової інфраструктури в сільських регіонах. Часто впровадження цифрових рішень має фрагментарний характер, що ускладнює їх інтеграцію в єдину систему управління. Крім того, відсутність чіткої державної стратегії цифровізації аграрного сектору та недостатня координація між зацікавленими сторонами гальмують процес трансформації.

3. ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Цифрова трансформація є ключовим чинником модернізації аграрного сектору України. Вона дозволяє підвищити ефективність виробництва, зменшити витрати та забезпечити сталий розвиток. Подальші дослідження мають бути спрямовані на розробку моделей цифрової трансформації для підприємств різного масштабу, а також на вивчення впливу цифровізації на соціально-економічний розвиток сільських територій.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

- [1] V. Tarasov, O. Kovalchuk, and I. Bondarenko, "Digitalization of agriculture in Ukraine: Challenges and opportunities," *AgroTech Journal*, vol. 12, no. 3, pp. 45–52, 2023.
- [2] M. Shevchenko and L. Hrytsenko, "Smart farming technologies in Ukrainian agribusiness," *Journal of Agricultural Innovations*, vol. 8, no. 2, pp. 33–41, 2024.
- [3] A. Ivanov, "Big Data and AI in precision agriculture," *Eastern European Journal of Technology*, vol. 10, no. 1, pp. 12–19, 2023.
- [4] O. Melnyk, "Digital platforms for farm management: Ukrainian experience," *Information Technologies in Agriculture*, vol. 7, no. 4, pp. 55–63, 2024.

Катерина Наконечна
кандидат економічних наук,
доцент кафедри економічної
кібернетики НУБіП України
<https://orcid.org/0000-0002-1537-7201>
Nakonechna@nubip.edu.ua

ANALYSIS OF THE VEGETABLE MARKET IN UKRAINE IN CONDITIONS OF WAR

Abstract. The theses examine the specific features of the functioning of Ukraine's vegetable market under martial law, which has led to significant structural and logistical changes in the agricultural sector. Economic and mathematical modeling methods are applied to analyze demand, supply, price dynamics, and regional distribution of production. Particular attention is paid to the impact of territorial occupation, disruption of supply chains, shifts in consumer priorities, and state support for agricultural producers. The modeling results enable the development of scenarios for market stabilization, logistics optimization, and food security assurance. The findings may be useful for shaping agricultural policy in times of crisis.

Vegetable market, agricultural sector, supply and demand, state support.

Introduction. The full-scale war in Ukraine, unleashed by Russia in 2022, caused significant damage to the country's economy. Agriculture and, in particular, the vegetable growing industry were significantly negatively affected by military operations and the occupation of territories. The issues of analysis of the current situation on the market of vegetable products and prospects for its development, taking into account the consequences of the war and as a result of the economic crisis in Ukraine, remain insufficiently researched. Therefore, an urgent task is the implementation of econometric modeling of the domestic vegetable market to identify trends and forecast indicators in the future short-term period.

Problem Statement. The purpose of the study is to substantiate the theoretical and methodological aspects and practical recommendations for the development of the domestic vegetable market based on the construction of an econometric model.

The study based on data for 2003-2021 on the influence of various economic and production factors on the pricing of the vegetable market. Since the data for 2022 are quite anomalous due to the full-scale invasion of the Russian Federation into Ukraine, it was decided not to take them into account in the general model, as they may distort the overall picture of the impact.

Analysis of recent research and publications. Recent scholarship has increasingly focused on the transformation of Ukraine's vegetable market in response to wartime disruptions. The literature [1], [2], [3], [4] reveals a convergence of methodological approaches aimed at capturing the complexity of agricultural systems under stress, including econometric modeling, scenario analysis, and spatial diagnostics. Issues related to the analysis of the current situation in Ukraine's vegetable market and its future development, considering the impact of the war and the ensuing economic crisis, remain underexplored.

Purpose of the publication. The purpose of this article is to provide a comprehensive analysis of the current state and development trends of Ukraine's vegetable market, taking into account the impact of the war and the economic crisis.

RESULTS AND DISCUSSION. The presented results not only illustrate the current state of the market but also provide a foundation for forecasting possible scenarios of its development. This creates a basis for formulating practical recommendations aimed at stabilizing agricultural policy, optimizing logistical solutions, and ensuring the country's food security. In our case, the average selling price of vegetable crops by agricultural enterprises will be the dependent variable, UAH/t, and ten different socio-economic indicators were chosen as the independent variables. Let's introduce the notation of dependent and independent variables (see Table 1).

Table 1

Designation of the investigated factors

Marking	Description
Dependent variable	
Average prices, y	average sales prices of vegetable crops by agricultural enterprises, hryvnias/ton
Independent variables	
Production volume, x1	volume of production (gross harvest) of vegetables, thousand tons
Sown areas, x2	sown areas of vegetables, thousand ha
Yield, x3	yield of vegetables, ts from 1 ha of harvested area
Vegetables for 1 person, x4	volume of vegetable production per person, kg/year
Salary, x5	average monthly salary, hryvnias/month
Employed population, x6	employed population aged 15-70 in agriculture, forestry and fisheries, thousands of people
Price index, x7	price index of crop products sold by enterprises, % compared to the previous year
Export, x8	export of vegetables, thousands of US dollars
Import, x9	import of vegetables, thousands of US dollars
Inflation, x10	inflation rate, % to the previous year

Source: developed by the authors.

We will find the value of the F-criterion using the formula (1):

$$F = \frac{\sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2}{\frac{\sum_{i=1}^n (Y_i - \hat{Y}_i)^2}{n-2}}, \quad (1)$$

Where $\hat{Y}_i$ - the estimate is obtained on the basis of a regression model;

$\bar{Y}$ - values for all studies Y_i :

Y_i - the value of the resulting variable;

n - is the sample size.

the obtained values entered in the table 2:

Table 2

F-test for factors

Factor	The meaning of the F-criterion
Sown areas, x2	9,73
Yield, x3	40,66
Vegetables for 1 person, x4	17,09
Salary, x5	16,99

Employed population, x6	0,43
Price index, x7	32,73
Export, x8	11,32
Import, x9	0,13

Source: developed by the authors.

At the level of significance $\alpha = 0.05$ and the degrees of freedom $f_1 = 1$ and $f_2 = 17$, the table F will be equal to 4.45. Therefore, according to the conducted direct selection, it is necessary to include the values of sown areas, productivity, wages, employed population, export and import of vegetables into the model.

Having built a model of the dependence of average prices for 1 ton of vegetables on the factors listed above, it was found that the p-value of the Sown area (x2) and Export (x8) indicators significantly exceeds the norm of 0.05 (Table 3), so it was decided to exclude given factors.

As a result, we got a model with the following values: $R^2 = 0.958$, adjusted R^2 is 0.946, standard error of estimation = 294.796.

Table 3.

Meaning of p-value				
	Coefficients	Standard Error	t Stat	P-value
Constant	5439,237	2959,009	1,838	0,091
Sown areas, x2	-5,940	8,546	-0,695	0,500
Yield, x3	16,894	7,020	2,406	0,033
Salary, x5	0,223	0,040	5,563	0,000
Employed population, x6	-0,967	0,513	-1,884	0,084
Import, x9	-0,003	0,002	-1,437	0,176
Inflation, x10	-0,004	0,002	-1,865	0,087

Source: developed by the authors.

Table 4.

Regression statistics	
Multiple R	0,979
R Square	0,958
Adjusted R Square	0,946
Standard Error	294,796
Observations	19

Source: developed by the authors.

The value of R^2 means that the variation of average prices for 1 ton of vegetables in Ukraine for 2003-2021 is 95.8% due to the above factors, which is a good estimate for the constructed model.

In the table 2.9 presented other obtained values for our constructed model. As we can see, all p-values are within the permissible range from 0.000 to 0.05.

Table 5

Coefficients of the regression model

	<i>Coefficients</i>	<i>Standard Error</i>	<i>t Stat</i>	<i>P-value</i>
Constant	3976,940	2014,011	1,975	0,031
Yield, x3	10,579	5,966	1,773	0,048
Wages, x5	0,221	0,032	6,902	0,000
Employed population, x6	-1,106	0,336	-3,294	0,005
Import, x9	-0,005	0,001	-3,474	0,004

Source: developed by the authors.

Conclusions and Directions for Future Research. The model built during the research is presented in formula (2.4). It can be used to estimate the impact of each factor on average prices for 1 ton of vegetables.

$$Y = 3976,94 + 10,58x_3 + 0,22x_5 - 1,11x_6 - 0,01x_9 \quad (2)$$

That is, with an increase in productivity by 1 hryvnia per hectare per year, the average selling price of vegetable crops by agricultural enterprises rises to UAH 10.58/t, and when the average monthly wage increases by UAH 1, it increases by UAH 0.22. With the increase in the number of employed population aged 15-70 in agriculture, forestry and fishing per 1 thousand people, the average price of selling vegetables decreases by UAH 1.11, and the cost of import of vegetable crops increases by USD 1 thousand. The USA provokes a reduction in the price by 0.01 hryvnias.

References

- [1] R. V. Lohosha, K. V. Mazur, and V. Yu. Krychivskyi, "Market Research of the Vegetable Products Market in Ukraine: Monograph". Vinnytsia: TVORY LLC, 2021, 344 p.
- [2] O. V. Semenda and I. I. Korman, "Modeling the Vegetable Market of Ukraine under Wartime Conditions," **Bulletin of Agricultural Economics**, vol. 12, no. 2, pp. 45–58, 2024.
- [3] V. P. Rud', "Vegetable Production in the Food Market System: Features, Trends, and Prospects," **Innovative Economy**, vol. 3, no. 1, pp. 112–120, 2021.

Роман Руденський

д.е.н., професор, професор кафедри комп'ютерних наук
Національний університет біоресурсів і природокористування України, м. Київ, Україна
0009-0002-3682-9702
roman.rudensky@nubip.edu.ua

Володимир Кравченко

д.е.н., доцент, професор кафедри економічної кібернетики
Національний університет біоресурсів і природокористування України, м. Київ, Україна
0000-0002-8033-3985
v.kravchenko@nubip.edu.ua

АНАЛІЗ ТЕМ І НАСТРОЇВ У ЦИФРОВОМУ ДИСКУРСІ АГРАРНОГО СЕКТОРУ

Анотація. Дослідження присвячено виявленню ключових тем і наративів, що циркулюють у цифровому медіапросторі аграрного сектору України, а також дослідження емоційної тональності (позитивної, негативної, нейтральної) повідомлень у межах виявлених тематичних кластерів. Розроблено моделі аналізу тем і настроїв в цифровому медіапросторі щодо публічного управління в аграрному секторі України з використанням методів машинного навчання та попередньо навченої моделі з бібліотеки Hugging Face для Python.

Ключові слова: аналіз тем; аналіз настроїв; машинне навчання, цифровий дискурс.

1. ВСТУП

В епоху цифрової трансформації громадська думка відіграє ключову роль у формуванні різних політичних рішень, особливо в таких стратегічних галузях, як аграрний сектор. Разом з тим, в онлайн просторі вона має значний вплив на поведінку та прийняття рішень різними зацікавленими сторонами, зокрема фермерами, посередниками і споживачами [1].

Постановка проблеми. Виходячи з цього, виникає потреба у використанні передових технологій для моніторингу та прогнозування цифрового дискурсу. ІТ-рішення, які базуються на методах машинного навчання та великих мовних моделях (LLM), дозволяють ефективно оцінювати настрої та думки громадян, їх зміни в реальному часі, що дає змогу своєчасно змінювати політику органів державної влади різних рівнів, забезпечуючи її відповідність вимогам та потребам суспільства.

Аналіз останніх досліджень і публікацій. Аналіз великих даних із засобів масової інформації відкриває можливості для створення моделей, що прогнозують громадську думку щодо державної політики [2]. Такий підхід враховує вплив громадської думки в цифровому просторі та активність інтернет-користувачів, спрямовану на демократичну участь, що дозволяє виявляти приховані інсайти, пов'язані з впровадженням нових політичних рішень [3]. З метою підвищення громадської безпеки у соціальних мережах розробили систему моніторингу громадської думки на основі технологій машинного навчання та методів сегментації слів і впровадили її на платформі Sina Weibo [4].

Мета публікації. Метою статті є розробка моделей аналізу тем і настроїв в цифровому медіапросторі щодо публічного управління в аграрному секторі України.

2. МЕТОДИ ДОСЛІДЖЕННЯ

Корпус текстів було сформовано на основі публікацій з провідних українських Telegram-каналів, спеціалізованих на аграрній тематиці: Agro News, Agro Portal, Agro Review, Agro Polit, Super Agronom, Landlord і Latifundist. Перед початком аналізу тексти були технічно підготовлені: очищені від службової інформації (HTML-теги, символи),

токенізовані та лематизовані. Для ідентифікації тематичних кластерів в цифровому дискурсі аграрного сектору в Україні застосовано: методи тематичного моделювання (topic modelling) з використанням алгоритму Latent Dirichlet Allocation (LDA) з бібліотеки gensim; метрику когерентності (c_v) для вибору оптимальних параметрів (кількість тем, кількість проходів); багатовимірною шкалювання (MDS) для візуалізації семантичної близькості. Для аналізу емоційної тональності (sentiment analysis) застосовано: попередньо навчену модель з бібліотеки transformers (Hugging Face) для класифікації повідомлень на позитивні, негативні та нейтральні; ProcessPoolExecutor для паралельної обробки великих обсягів тексту. Варто зазначити, що було здійснено перетворення оцінок настроїв у числову шкалу від -1 (негативний) до +1 (позитивний) для кількісного аналізу.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

3.1. Тематичне моделювання

Для визначення найкращих параметрів моделі випробовуються різні комбінації кількості тем (num_topics) і кількості проходів (passes). Найвищий показник когерентності (0.517) досягнуто при num_topics = 5 і passes = 50. Метод тематичного моделювання до зібраного корпусу текстів із використанням ймовірнісного алгоритму дозволив ідентифікувати п'ять провідних тематичних кластерів: Тема 0 "Новини українського агробізнесу: фінансування, розвиток, міжнародна інтеграція"; Тема 1 "Інтеграція технологій і агросистем: від захисту рослин до стратегій сталого землеробства"; Тема 2 "Врожайність сільськогосподарської продукції в Україні: ринкові та природні фактори"; Тема 3 "Український експорт агропродукції через морські та сухопутні коридори: виклики та можливості на тлі геополітичних змін"; Тема 4 "Державна політика й аграрні реформи в Україні".

3.2. Аналіз настроїв

Щомісячні зміни настроїв проілюстровані за допомогою діаграми "candlestick" (рис. 1). За допомогою цієї діаграми можна побачити зміну настрою з бурхливої волатильності (2016–2019 рр.) у стійке ослаблення (2022–2024 рр.). Це видно по довших "свічках" (тілах і тінях), що вказує на великі коливання настрою як протягом місяця (діапазон між 1-м і 3-м квантилем, показаний тінями), так і між початком та кінцем місяця (показано тілом свічки).

Довге тіло свідчить про сильну зміну, коротке – про незначну зміну. Колір тіла показує напрямок руху настрою за період: зелене – настрої в кінці місяця кращі ніж на початку, тобто оцінка зросла; червоне – настрої в кінці місяця гірші ніж на початку, тобто оцінка спала. Верхня тінь – лінія від верхньої межі тіла до 3-го квантилю оцінки настрою, тобто показує, наскільки високо піднімалася оцінка вище кінцевого рівня для зеленої свічки або початкового рівня для червоної свічки. Нижня тінь – лінія від нижньої межі тіла до 1-го квантилю оцінки настрою за місяць. Показує, наскільки низько вона опускалася нижче початкового рівня для зеленої свічки або кінцевого для червоної свічки. Довжина тіней вказує на волатильність настрою поза діапазоном початок-кінець. Довгі тіні свідчать про значні коливання настрою протягом періоду, які не втрималися до закриття. Тіні свічок (між квантилями 1 і 3) показують розкид настроїв протягом місяця. Навіть у місяці, коли загальна зміна настрою від початку до кінця місяця була невеликою (коротке тіло), могли бути значні коливання протягом місяця (довгі тіні).

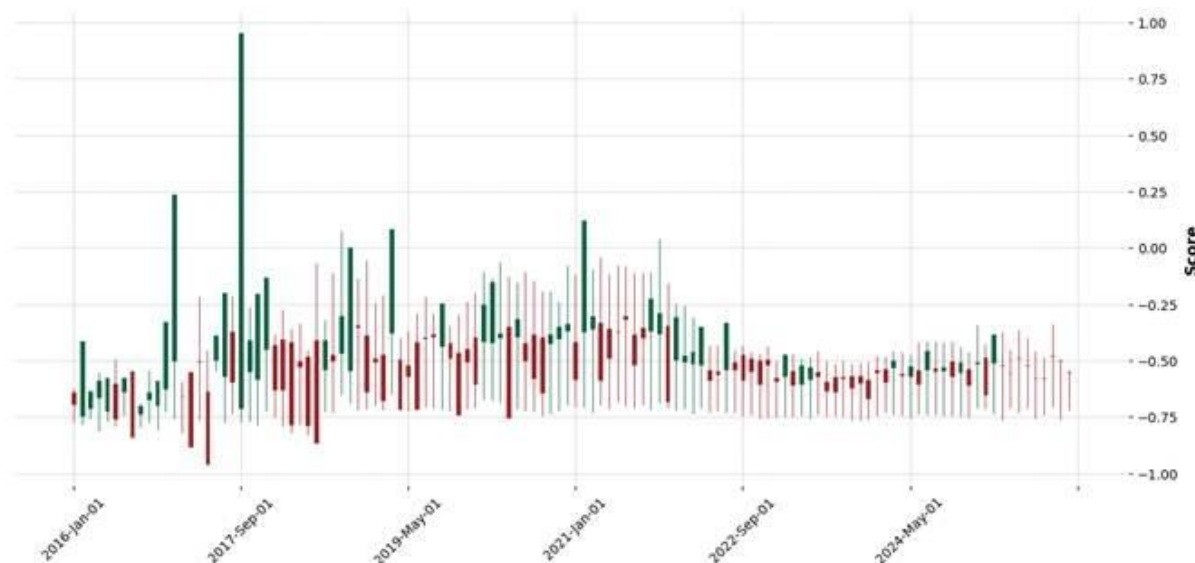


Рис. 1. Діаграма "candlestick" щомісячних змін настроїв

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Отже, дослідження одним з перших застосувало тематичний аналіз і аналіз настрою до українських аграрних медіа. Результати розкривають роль етично спрямованої аналітики медіа в підвищенні прозорості, передбаченні суспільних занепокоєнь та підтримці науково обґрунтованого керування в умовах національної невизначеності. Виявлено 5 домінантних тематичних кластерів; найбільшу увагу отримали аграрні реформи та політика. Аналіз настрою показав переважно негативний тон за всіма темами, особливо щодо експорту та стану ринку. Контент про аграрну політику демонстрував відносно більш позитивний настрій. Емоційний тон значно варіювався у часі, що свідчить про динамічні та залежні від теми суспільні позиції.

ПОСИЛАННЯ

- [1] Lv, X. et al. (2024) 'Spillover effect of network public opinion on market prices of small-scale agricultural products', *Mathematics*, 12(4), p. 539. doi:10.3390/math12040539.
- [2] Chen, X. (2021) 'Research on public opinion forecast of public policy based on Big Data Media', 2021 5th Annual International Conference on Data Science and Business Analytics (ICDSBA), pp. 103–106. doi:10.1109/icdsba53075.2021.00030.
- [3] Redjeki, S. and Widyarto, S. (2022) 'Comparison of seven machine learning algorithms in the classification of public opinion', *Tech-E*, 5(2), pp. 143–149. doi:10.31253/te.v5i1.1046.
- [4] Weidong, G. et al. (2022) 'Network public opinion monitoring method based on machine learning', 2022 International Conference on Intelligent Transportation, Big Data & Smart City (ICITBS), pp. 178–181. doi:10.1109/icitbs55627.2022.00047.

Людмила Галаєва

Кандидат економічних наук, доцент, доцент кафедри економічної кібернетики
Національний університет біоресурсів і природокористування України, м.Київ
ORCID ID: <https://orcid.org/0000-0003-3036-2830>
email: lgalaeva@nubip.edu.ua

ПРО ДЕЯКІ ПІДХОДИ ДО МОДЕЛЮВАННЯ СИСТЕМИ ПЕНСІЙНОГО СТРАХУВАННЯ

Анотація. Розглянуто сучасні тенденції функціонування пенсійної системи України, досліджено особливості її формування. Проаналізовано основні виклики, що стоять перед пенсійною системою, серед яких вплив традиційних загроз, зокрема демографічні та економічні чинники та нові виклики, такі як нестандартна зайнятість, глобалізація, наслідки, зумовлені війною в країні. На основі запропонованого підходу до моделювання розвитку системи пенсійного страхування в Україні, висловлено припущення, що запровадження другого рівня пенсійної системи може пом'якшити ці ризики, але потребує комплексних реформ, спрямованих на забезпечення макроекономічної стабільності, розвитку ринку праці й фінансового ринку в країні.

Ключові слова: пенсійна система, пенсійне забезпечення, пенсійне страхування, моделювання, прогноз, ризики, оптимістичний сценарій, песимістичний сценарій.

1. ВСТУП

Постановка проблеми. У сучасних умовах, зважаючи на поточну демографічну ситуацію в країні, пов'язану з міграцією населення, низьким рівнем народжуваності, економічними чинниками та її подальші перспективи, докорінне реформування існуючої пенсійної системи слід вважати одним з першочергових завдань, що сприятиме укріпленню економічної безпеки України. Моделювання розвитку системи пенсійного страхування в Україні дасть можливість проаналізувати ефективність систем I та II рівня та оцінити наслідки запровадження системи II рівня.

Аналіз останніх досліджень і публікацій. Питання функціонування державних та недержавних пенсійних фондів розглядаються у працях зарубіжних та вітчизняних науковців. Значимий внесок у формування національної концепції пенсійного забезпечення та його складової – недержавного пенсійного забезпечення, зробили: В.Виговська [1], Г.Овчаренко [2], Я.Прищепя [3] та інші. Проте, незважаючи на наявність ряду досліджень, питання системного підходу до моделювання системи пенсійного забезпечення, зокрема, на основі моделі пенсійного страхування в умовах України, потребує подальшого вивчення.

Мета публікації. Враховуючи це, метою статті є висвітлення деяких підходів до моделювання та прогнозування системи пенсійного страхування в Україні, що дозволить розробляти ефективні стратегії її розвитку та адаптувати до сучасних викликів.

2. МЕТОДИ ДОСЛІДЖЕННЯ

Для досягнення мети використано системний підхід, методи аналізу і синтезу, аналогії і порівняльного аналізу, формальної логіки, статистичні методи, методи аналізу, економіко-математичні методи моделювання. Статистичну базу дослідження становлять офіційні дані Пенсійного фонду України.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Моделювання системи пенсійного страхування реалізується за рахунок різних видів моделей, серед яких можна виділити:

- актуарні моделі, які використовуються для розрахунку страхових тарифів, резервів та актуальних зобов'язань;
- економетричні моделі, які дають можливість робити прогнози, зокрема враховувати динаміку змін демографічних, фінансових показників тощо;
- імітаційні моделі, які базуються на методі Монте-Карло для прогнозування в умовах невизначеності та ризику та дозволяють моделювати широкий спектр можливих сценаріїв.

Для реалізації цих моделей використовуються, як спеціальне програмне забезпечення, так і статистичні пакети, зокрема R, Stata, SPSS. Для простих розрахунків та моделювання може бути достатньо електронних таблиць (Excel); для більш складних моделей – мови програмування (Python, MATLAB) тощо.

Для оцінки стану пенсійного страхування в Україні необхідно використовувати цілий ряд показників, серед яких можна виділити основні (табл.1).

Таблиця 1

Динаміка основних показників для оцінки стану пенсійного страхування в Україні

№	Показник	2020	2021	2022	2023	2024
1	Кількість пенсіонерів, млн. осіб	11,1	10,8	10,5	10,2	9,9
2	Середній розмір пенсії, грн.	3 083	3 778	4 190	4 853	5 312
3	Мінімальна пенсія, грн	1 712	1 854	2 093	2 361	2 610
4	Середня заробітна плата, грн.	11 591	13 200	14 934	17 300	19 550
5	Коефіцієнт заміщення (пенсія/зарплата), %	26,6	28,6	28,1	28	27,1
6	Дефіцит Пенсійного фонду, млрд. грн.	195,3	203,8	233,2	244,5	256

Джерело [4]

Як видно з таблиці, кількість пенсіонерів поступово зменшується через демографічні зрушення та старіння населення, зумовлені міграційними процесами та війною в країні. Водночас середній розмір пенсії за п'ять років збільшився на 72,3%, частково завдяки індексації та зростанню середньої заробітної плати. Однак коефіцієнт заміщення доходу нижчий за рекомендований Міжнародною організацією праці (МОП) рівень 40%, що свідчить про низький рівень пенсійних виплат.

На основі економіко-математичних моделей було побудовано два сценарії розвитку системи пенсійного страхування:

-за оптимістичним сценарієм (за умови запровадження другого рівня накопичувальної системи з 2026 року) очікується, що середня пенсія зросте до 8500 гривень до 2030 року, а коефіцієнт заміщення досягне 35 відсотків;

-за песимістичним сценарієм (без структурних реформ) очікується, що середня пенсія залишиться на рівні 6200 гривень, а дефіцит пенсійного фонду до 2030 року буде більше 300 мільярдів гривень.

Отже, результати моделювання показують, що запровадження накопичувального рівня фінансування підвищує фінансову стійкість систем пенсійного страхування та забезпечує покращення добробуту пенсіонерів за умови одночасного підвищення прозорості фінансового ринку та рівня зайнятості населення.

Зроблений за оптимістичним сценарієм прогноз середнього розміру пенсії показано на рис.1.

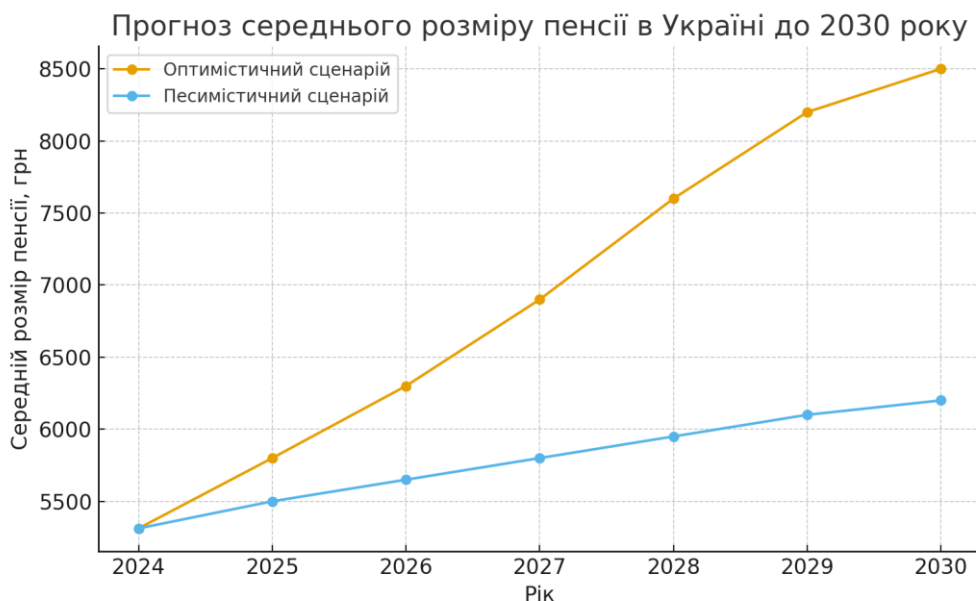


Рис. 1. Прогноз середнього розміру пенсії в Україні

Джерело: [4]

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Результати дослідження свідчать, що реформування системи пенсійного страхування є ключовим для економічної безпеки України. Пенсійна реформа має передбачати перерозподіл відповідальності між державою, роботодавцями й працівниками, що сприятиме підвищенню мотивації до сплати пенсійних внесків і зниженню тіньової економіки.

Прогнозування розвитку системи пенсійного забезпечення, зокрема страхування, є досить складним, через існування значної кількості прямих і прихованих чинників, зокрема таких, які важко передбачити, що впливають на досліджувані процеси.

Подальші дослідження доцільно спрямувати на розробку моделей прогнозування розвитку системи пенсійного страхування, зокрема з недержавних фондів, та вивчення впливу такого підходу на реформування системи пенсійного забезпечення.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

[1] V.Vyhovska, "Modern problems of development of non-state pension provision in Ukraine". *Problems and prospects of economy and management. Financial resources: problems of formation and use*, no. 4 (24). pp. 170-178, 2022.

[2] H. Ovcharenko. "Funded pension system: there is no point in waiting for better times", 2022. URL: <https://www.epravda.com.ua/columns/2022/11/11/ 693709>

[3] Ya. Pryshchepa., H. Matviishyna. Funded pensions: the Ministry of Social Policy explained what the pension reform entails. 2023. URL: <https://suspilne.media/356294-nakopichuvanni-pensii-u-minsocpolitiki-rozpovili-so-peredbacae-pensijna-reforma>.

[4] Pension Fund of Ukraine. URL: <https://www.pfu.gov.ua/>

SECTION 2. COMPUTER SYSTEMS AND NETWORKS, CYBERSECURITY / КОМП'ЮТЕРНІ СИСТЕМИ І МЕРЕЖІ, КІБЕРБЕЗПЕКА

Максим Місюра

Кандидат технічних наук, доцент кафедри комп'ютерних систем, мереж та кібербезпеки
Національний університет біоресурсів і природокористування України, факультет інформаційних
технологій, кафедра комп'ютерних систем, мереж та кібербезпеки, Київ, Україна
ORCID ID 0000-0002-9061-3462
mdm@nubip.edu.ua

Сергій Мамченко

Доктор педагогічних наук, професор, професор кафедри комп'ютерних систем, мереж та
кібербезпеки
Національний університет біоресурсів і природокористування України, факультет інформаційних
технологій, кафедра комп'ютерних систем, мереж та кібербезпеки, Київ, Україна
ORCID ID 0009-0006-8743-5606
s.mamchenko@nubip.edu.ua

ВБУДОВАНА ІОТ СИСТЕМА КОНТРОЛЮ ДОСТУПУ З БІОМЕТРИЧНОЮ АВТЕНТИФІКАЦІЄЮ НА БАЗІ ESP32 ДЛЯ ІНФРАСТРУКТУРИ РОЗУМНИХ БУДІВЕЛЬ

Анотація. У роботі представлено embedded-систему контролю доступу для інфраструктури «розумних» будівель із використанням біометричної автентифікації та інтеграцією до IoT. Апаратною основою є мікроконтролер ESP32 як центральний вузол оброблення даних і мережевої взаємодії; до нього підключено сенсор відбитків пальців (типу R305/AS608), електромеханічний замок і допоміжні модулі. Біометричні дані зчитуються локально, далі здійснюється передача на сервер (через Wi-Fi) для порівняння з еталонними шаблонами і прийняття рішення щодо надання доступу; усі події логуються в базі даних з можливістю віддаленого моніторингу через мобільний застосунок адміністратора. Запропонований підхід поєднує високу зручність користувача з підвищеною надійністю автентифікації порівняно з традиційними методами (ключі, паролі), а також забезпечує масштабованість завдяки IoT-орієнтованим протоколам та гнучким політикам доступу для різних ролей. Новизна полягає в практичній реалізації повного циклу — від сенсора до серверної верифікації та керування виконавчим механізмом — на доступній платформі ESP32 з підтримкою сучасних інтерфейсів, енергозберігаючих режимів і бездротових стеків Wi-Fi/BLE. Додатково обговорюються безпекові аспекти інтеграції (TLS-шифрування трафіку, захист шаблонів, обмеження частоти спроб, 2FA) та сценарії застосування у закладах освіти, офісних просторах, лабораторіях і житлових комплексах. Результати прототипування підтверджують доцільність підходу для побудови безпечних, масштабованих і керованих систем доступу, сумісних із екосистемами smart-building, і окреслюють подальші кроки — використання AI для виявлення аномалій, інтеграцію відеопотоків і впровадження незмінних журналів подій на базі блокчейн-технологій.

Ключові слова: Embedded system; IoT; ESP32; вбудована система; IoT; ESP32; контроль доступу; біометрична автентифікація; розумні будівлі.

1. ВСТУП

1.1. Постановка проблеми

Швидкий розвиток концепції «Розумного будинку» та Інтернету речей (IoT) вимагає впровадження високозахищених та ефективних систем контролю доступу (СКУД). Традиційні СКУД часто страждають від низької надійності, вразливості до компрометації фізичних ключів та складності централізованого моніторингу. Метою даної роботи є розробка гібридної вбудованої IoT-системи контролю доступу з біометричною автентифікацією, яка забезпечує швидку локальну верифікацію та інтегрована з хмарною платформою для віддаленого моніторингу та керування.

1.2. Аналіз останніх досліджень і публікацій

Сфери IoT та біометричної автентифікації (БА) є ключовими для сучасних систем безпеки [1], [2]. Аналіз показує, що IoT-рішення стають галузевим стандартом, забезпечуючи легкість масштабування [3], [4], а ринок систем електронного контролю доступу активно зростає [5], . Серед апаратних рішень, мікроконтролери ESP32 широко використовуються завдяки їхній низькій вартості та вбудованому Wi-Fi, що ідеально підходить для периферійних (Edge) обчислень. Однак більшість існуючих рішень зосереджені або на повністю локальному контролі (без моніторингу), або на повному хмарному контролі (що збільшує затримку). Розроблена система пропонує гібридний підхід, поєднуючи швидкість локальної БА з надійністю хмарного журналу подій.

1.3. Мета публікації

Метою даної публікації є опис розробки та практичної реалізації програмно-апаратної системи контролю доступу. Система базується на біометричній ідентифікації (відбиток пальця) і забезпечує як локальну перевірку на пристрої, так і централізоване логування подій та керування правами доступу через серверну частину.

Для досягнення цієї мети були вирішені наступні завдання:

- Обґрунтовано вибір гібридної архітектури (локальна перевірка, серверне логування).
- Реалізовано апаратну частину на базі мікроконтролера ESP32 та біометричного сенсора.
- Розроблено серверну частину на веб-фреймворку Flask з базою даних PostgreSQL.
- Розгорнуто серверне рішення на хмарній платформі Render для забезпечення глобальної доступності.
- Протестовано систему в реальних сценаріях використання.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

2.1. Технічна реалізація та архітектура системи

Система функціонує на гібридній архітектурі, що поєднує периферійні обчислення та централізоване хмарне керування для забезпечення високої швидкості та масштабованості.

Локальний IoT-вузол є автономною частиною (Елемент 1), що забезпечує швидкий доступ та фізичний контроль:

- Мікроконтролер ESP32-WROOM-32: Керуюче ядро із вбудованим Wi-Fi. Виконує алгоритми локальної верифікації та керує виконавчим механізмом.
- Біометричний сенсор FPM10A: Відповідає за зчитування, обробку та порівняння відбитка пальця. Підключений до ESP32 через інтерфейс UART.
- Виконавчий механізм (Реле/Замок): Фізичний контроль доступу, керований через лінії GPIO.

Програмна реалізація та синхронізація: Локальна база даних (SQLite) на ESP32 зберігає шаблони відбитків для швидкої автономної верифікації. У разі успішної або невдалої автентифікації, мікроконтролер асинхронно відправляє лог-події через Бездротову мережу (Wi-Fi) та Інтернет (Global Network) (Елемент 2) на Хмарний Сервер.

Хмарна інфраструктура (Елемент 3) є централізованою частиною для адміністрування: Хмарний Сервер: Хостинг (наприклад, Render) для розгорнутого застосунку. Бекенд (Flask/Python): Ядро серверної логіки, що обслуговує REST API для синхронізації логів, обробки запитів керування та взаємодії з інтерфейсом. Централізована База Даних (PostgreSQL): Використовується для надійного зберігання реєстру користувачів та детальних журналів подій доступу (логів).

Користувачський інтерфейс (Елемент 4): Надає адміністраторам засоби для моніторингу та дистанційного керування через Вебінтерфейс (Адміністративна панель), доступний також на Мобільних пристроях.

2.2. Тестування та обговорення результатів

Тестування підтвердило коректну роботу всіх заявлених функцій: локальна ідентифікація успішно розрізняє зареєстрованих та незареєстрованих користувачів, а адміністративна панель коректно відображає журнал подій.

Аналіз показників швидкодії та відмовостійкості. Критична функція локальної ідентифікації виконується приблизно за 1 секунду. Цей показник є важливим, оскільки неблокуючий характер процесу забезпечує високу зручність користувача. Затримка при відправці логу на хмарний сервер склала близько 7 секунд. Ця затримка, пов'язана з особливостями безкоштовного тарифного плану хмарної платформи, є прийнятною, оскільки процес логування відбувається асинхронно і не блокує надання доступу. Реакція системи на команди ручного керування з вебінтерфейсу становить 1-3 секунди. Система довела свою ефективність, надійність та гнучкість в керуванні, поєднуючи швидкоку локальну верифікацію з перевагами централізованого хмарного моніторингу.

ВИСНОВКИ

У роботі було успішно вирішено завдання розробки та впровадження комп'ютерної системи моніторингу та контролю доступу. Система базується на **IoT-пристроях** та орієнтована на інфраструктуру **«розумних» будівель**.

Підсумки та досягнуті результати: 1. Розроблено вбудовану систему контролю доступу на апаратній платформі мікроконтролера ESP32, яка включає функціональність біометричної аутентифікації для підвищення рівня безпеки та надійності. 2. Створена гібридна архітектура, що поєднує переваги швидкої локальної верифікації безпосередньо на IoT-пристрої з централізованим хмарним моніторингом та керуванням. 3. Забезпечено ефективний механізм управління даними, коли синхронізація журналів верифікації та даних користувачів з локальної бази даних (SQLite) до централізованої (PostgreSQL) відбувається асинхронно (не блокуючи процес надання доступу). 4. Проведена апробація підтвердила, що система є надійною, ефективною та гнучкою у керуванні. Час реакції системи на команди ручного керування через вебінтерфейс становить 1-3 секунди, що відповідає сучасним вимогам до систем контролю доступу.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

[1] T. Lynn, S. P. G. Singh, M. O'Connell, and P. Rosati, "The Internet of Things: Definitions, Key Concepts, and Reference Architectures," in *The Cloud-to-Thing Continuum*. Cham: Springer, 2020, pp. 1–25.

[2] A. K. Jain and A. Kumar, "Biometric Recognition: An Overview," in *Second Generation Biometrics: The Future of Biometrics*. New York: Springer, 2012, pp. 3–18.

[3] OptConnect, "Elevating Access Control with IoT," 2023. [Online]. Available: <https://www.optconnect.com/blog/elevating-access-control-with-iot>. [Accessed: Jun. 15, 2025].

[4] A. Salim, S. A. Ghafoor, and K. M. Ghafoor, "IoT-based Smart Access Control Systems: Architecture and Implementation," *Sensors (Basel)*, vol. 20, no. 11, p. 3088, Jun. 2020. doi: 10.3390/s20113088.

[5] SNS Insider, "Electronic Access Control Systems Market to Grow USD 106.68 Billion by 2032," 2025. [Online]. Available: <https://www.snsinsider.com/reports/electronic-access-control-systems-market-1835>. [Accessed: Jun. 15, 2025].

[6] O. B. Prus, "Rozrobka komp'uternoi systemy monitorynhu ta kontroliu dostupu z vykorystanniam IoT prystroiv" ("Development of a computer system for monitoring and access control using IoT devices"), Bachelor's Thesis, Dept. of Computer Systems, Networks and Cybersecurity, Natsionalnyi Universytet Bioresursiv i Pryrodokorystuvannia Ukrainy (National University of Life and Environmental Sciences of Ukraine), Kyiv, Ukraine, 2025.

Сергій Мамченко, доктор педагогічних наук, професор

Національний університет біоресурсів та природокористування України, факультет інформаційних технологій, кафедра комп'ютерних мереж, систем та кібербезпеки, професор кафедри

ORCID ID 0009-0006-8743-5606

s.matchenko@nubip.edu.ua

Максим Місюра, кандидат технічних наук,

Національний університет біоресурсів і природокористування України, факультет інформаційних технологій, кафедра комп'ютерних систем, мереж та кібербезпеки, Київ, доцент кафедри

ORCID ID 0000-0002-9061-3462

mdm@nubip.edu.ua

Микита Журавльов

Національний університет біоресурсів і природокористування України, факультет інформаційних технологій, кафедра комп'ютерних систем, мереж та кібербезпеки, магістрант

ORCID ID 0009-00099085-3480

zhuravlov.mykyta.13@gmail.com

ВИКОРИСТАННЯ MESH-СИСТЕМ У АГРАРНО-ПРОМИСЛОВОМУ КОМПЛЕКСІ УКРАЇНИ

Анотація. Сучасний світ потребує забезпечення різних компонент безпечного простору. Перш за все необхідно забезпечити продуктову безпеку та підтримання агропромислового комплексу на високому технологічному рівні. Для України, яка знаходиться під військовою агресією РФ дане питання є вельми актуальне. Збереження технологічної спроможності збереження обробки сільськогосподарських угідь, об'єднання технічних систем у сукупність IoT мережі із динамічною топологічною структурою відповідно до прив'язки до території. Проведене дослідження дозволяє визначити технологію Mesh, яка саме і використовується завдяки своїм основам до застосування у вимогах регіонах із нерозвинутою інфраструктурою, корпоративним середовищем, що підтримує безшовний принцип побудови та функціонування. Також з позитивної сторони Mesh-технологій виділяє такий принцип її побудови як комірчастий. Мережі, що побудовані на комірчастому підході дозволяють створювати економічно ефективні, динамічні мережі з високою пропускну здатністю.

Ключові слова: комп'ютерні мережі, комірчасті мережі, Mesh-мережі.

Вступ.

Військова агресія РФ проти України загострила питання успішного ведення агропромислового бізнесу. Мінування земельних ділянок, демографічна проблема, руйнування існуючого матеріально-технічного забезпечення активізувало питання суттєвого технологічного переоснащення сфери економіки, пов'язаної із природокористуванням та відновленням біоресурсів держави.

Впровадження нових інформаційних та комп'ютерних технологій у агропромисловість визначає вимоги до них, серед яких слід зазначити наступні: модульність побудови системи, автономність функціонування, мережева інфраструктура, застосування протоколів, що пройшли випробування практикою, невелика вартість та широкий доступ до апаратно-програмного забезпечення (через торговельну мережу загального доступу), можливість оперативної зміни геометрії мережевого покриття тощо. Вказані вимоги визначаються специфічними умовами ведення діяльності сільського господарства, серед яких слід виділити наступні: робота у місцях, віддалених від енергопостачання та місць проживання людей; зміна геолокації місць обробки землі та угідь, що додає мобільність переміщення місця роботи та зміна інфраструктури мережі, а звідси й зміна площі розповсюдження сигналу мережі. Зрозуміло, що подібні системи мають бути організованими на основі безпроводних систем (Wi-Fi).

Постановка проблеми. Для визначення платформи для організації зазначеної комп'ютерної системи із визначеними особливостями проведено пошук та аналіз організації подібних систем. Серед таких було виділено системи надання освітніх систем [1], медичних послуг [2], та впровадження передових технологій у АПК [4-8].

Традиційні мережі WiFi є централізованими, тобто всі пристрої підключаються до одного роутера. Під'єднання до кожного роутера відбувається за ніком та паролем. Відстань від роутера, з додаванням перешкод WiFi, впливає на покриття бездротового Інтернету по всій території. Це призводить до появи зон зі слабким сигналом і мертвих зон. Кожен роутер потребує перереєстрації для продовження роботи. Дана обставина ускладнює побудови комп'ютерних мереж з динамічною зоною роботи.

Для вирішення даної проблеми за основу була обрана технологія коміркових мереж. Топологія цих мереж дозволяє будувати мережі за принципом комірок, у яких робочі станції мережі з'єднуються одна з одною та здатні брати на себе місію комутатора для інших учасників. Хоча дана організація вважається досить складною у налаштуванні, однак дані мережі показали високу стійкість при спробі їх руйнування. Зазвичай вузли мережі підтримують зв'язок кожний з кожним, що дозволяє збільшити кількість маршрутів трафіку всередині мережі. Слід також зазначити, що вихід з робочого стану якогось вузла не приводить до руйнування всієї мережі.

Прикладом практичної реалізації даного підходу можна привести наступне: у сільському районі Каталонії (2004 рік) була впроваджена мережа guifi.net. Вона була розроблена як проект на відповідь відмови провайдерів налагодити мережеві послуги у регіоні з економічної причини – недоступність широко-смугового Інтернету. На сьогодні дана мережа об'єднує більше 30000 вузлів та завдяки домовленості peer to peer дана мережа залишається відкритою з широкими можливостями для реєстрації.

За другий приклад наведемо використання армією США безпроводних коміркових мереж для організації комп'ютерних мереж у польових умовах.

З використанням зазначених підходів до побудови комп'ютерних мереж із динамічною геометрією загальної зони покриття використовуються mesh-мережі із достатнім апаратно-програмним забезпеченням у торгівельній мережі. Доступність обладнання дозволяє говорити про практичну реалізацію даного класу комп'ютерних мереж.

По перше, Mesh-мережі децентралізовані. Вони пропонують кілька точок підключення у межах потрібної зони покриття. Пристрої підключаються до найближчої точки доступу, і дані проходять через кілька точок. Таке налаштування забезпечує більше покриття і більшу пропускну здатність. Традиційна мережа WiFi такі вимоги не виконує. Всі вузли системи об'єднуються в одну мережу з одним ім'ям (SSID) та паролем, створюючи ілюзію одного великого роутера з широким радіусом дії.



Рисунок 1 – Приклади реалізації mesh-мережі в АПК

Отже, Mesh WiFi – це одна велика мережа WiFi. Незважаючи на те, що вона складається з декількох точок доступу, система має лише одне ім'я мережі та один пароль. Бездротові пристрої автоматично підключаються до найближчого вузла без

необхідності вручну змінювати мережеві підключення. Існують інші мережеві пристрої, наприклад, ретранслятори, які створюють другу мережу. Це змушує користувача перемикатися між двома мережами залежно від того, де ви перебуваєте.

Топологія мережі може бути повно- або частково комірковою. Повна коміркова мережа означає, що всі вузли можуть спілкуватися один з одним. Частково коміркова мережа складається з вузлів, які можуть спілкуватися лише з певними вузлами. Перша є найбільш поширеною.

Також слід зазначити, що безмовність роботи mesh-мереж дозволяє поєднати в одну мережу до десяти модулів. Пристрої працюють самі собою. Як тільки буде доданий новий модуль, мережа його побачить та задасть потрібну конфігурацію. Що вирішує питання динамічної зміни зони покриття відповідно до потреб діяльності підприємства АПК. Враховуючи відсутність необхідності прокладання кабелів та простота створення мережі дана технологія може слугувати основою для вирішення питання забезпечення АПК передовими інформаційними та технологічними технологіями. Слід додати також і можливість включення у мережу IoT пристроїв та оперативного керування ними

Доступність на ринку обладнання говорить про його широкий вибір. Так, бюджетні mesh-комплекти — дводіапазонні. Вони працюють на 2 каналах: 2,4 та 5 ГГц. По обох може йти трафік для користувачів. Канал 5 ГГц швидше, але за його вибору радіохвилі гірше проходять крізь стіни. З частотою 2,4 ГГц проблема може полягати у зашумленості ефіру через роутери сусідів. При необхідності модулі самі переходять з 5 на 2,4 ГГц, оскільки не ідеальний зв'язок краще за ніякий. Два канали — істотна перевага перед звичайними одно-канальними роутерами. Користувальницький та службовий сигнали йдуть у різних діапазонах, що сприяє комфортній швидкості трафіку та не перевантаження даних трафіків.

ПОСИЛАННЯ

1. Dinuka Kannangara Kaushalya Yatigammana Faculty of Commerce & Management Studies University of Kelaniya Gamini Wijayarathna Faculty of Computing & Technology University of Kelaniya. Conceptualization of a Wireless Mesh Network Model for E-learning in Remote Areas. https://www.researchgate.net/publication/354529012_Conceptualization_of_a_Wireless_Mesh_Network_Model_for_E-learning_in_Remote_Areas. September 2021.
2. Building a Mesh Network: Step-by-Step Guide on How to Construct a Mesh Network <https://www.neuralword.com/en/innovation-technology/innovations-future-products/building-a-mesh-network-step-by-step-guide-on-how-to-construct-a-mesh-network/> 9 November 2023.
3. Damon · What is Mesh Networking: Types, Benefits, and Solutions - VSOL <https://www.vsolcn.com/blog/what-is-mesh-networking.html> Published on: February 19, 2025
4. Meshmerize GmbH. Smart agriculture network: How is Mesh changing the game? <https://meshmerize.net/smart-agriculture-network-how-is-mesh-changing-the-game/> 05.12.2023
5. Cultivate Smarter with MeshTek's BLE Mesh for Agriculture IoT // <https://www.meshtek.com/agriculture-iot-connectivity/>
6. RocketMe Up I/O | Medium Mesh Networks in IoT — Enhancing Connectivity in Large-Scale Deployments. <https://medium.com/@RocketMeUpIO/mesh-networks-in-iot-enhancing-connectivity-in-large-scale-deployments-d1e42e68318c>. Sep 12, 2024.
7. CMI Technolog Revolutionising Agriculture with Kinetic Mesh Networks <https://www.cmitechnology.com/revolutionizing-agriculture-with-kinetic-mesh-networks/> 19 October 2025.
8. Keoma Brun-Laguna ^a, Ana Laura Diedrichs, Diego Dujovne, Carlos Taffernaberry, Rémy Léone, Xavier Vilajosana, Thomas Watteyne Using SmartMesh IP in Smart Agriculture and Smart Building applications <https://www.sciencedirect.com/science/article/abs/pii/S0140366417305212/>. 2018.

Кирило Кохан

аспірант

кафедра комп'ютерних інформаційних систем і технологій

Національний університет біоресурсів і природокористування, Київ, Україна

ORCID: 0009-0002-8878-8527

kokhan.kyrylo@gmail.com

ОПТИМІЗАЦІЯ ТЕСТОВИХ КОНФІГУРАЦІЙ БАГАТОКОМПОНЕНТНИХ ІНФОРМАЦІЙНИХ СИСТЕМ ЗА ДОПОМОГОЮ PAIRWISE TESTING

Анотація. У статті представлено результати дослідження та впровадження методу Pairwise Testing для оптимізації вибору конфігурацій під час автоматизованого тестування багатокомпонентних інформаційних систем (ІС). Проблема комбінаційного вибуху, яка виникає зі зростанням кількості параметрів, платформ і середовищ, призводить до експоненційного збільшення кількості тестів і значних витрат часу та ресурсів. Для мінімізації цих витрат запропоновано використати комбінаторний підхід Pairwise Testing, який забезпечує перевірку всіх можливих пар параметрів системи, суттєво скорочуючи обсяг тестових сценаріїв без втрати якості покриття.

У роботі проаналізовано теоретичні основи Pairwise Testing, що базуються на припущенні, що більшість дефектів програмного забезпечення зумовлена взаємодією не більше ніж двох параметрів одночасно. Такий підхід дозволяє зменшити кількість тестів з експоненційної до квадратичної залежності, забезпечуючи 85–95% покриття потенційних помилок. Для побудови оптимальних конфігурацій побудована власна система генерації конфігурацій тестування.

На основі розробленої параметричної моделі ІС проведено експериментальне дослідження у середовищі CI/CD, що охоплює етапи побудови таблиць парних комбінацій, автоматичної генерації тестових сценаріїв, інтеграції у конвеєр тестування та аналізу отриманих результатів. Отримані дані засвідчили скорочення кількості виконуваних тестів на 60–80%, зменшення часу тестування на 30–50% і збереження рівня покриття понад 90%.

Запропонована інформаційна технологія Pairwise-оптимізації створює формалізований підхід до вибору тестових сценаріїв, підвищує ефективність використання обчислювальних ресурсів та забезпечує адаптивність системи тестування до змін у конфігураціях компонентів. Перспективи подальших досліджень полягають у поєднанні Pairwise Testing із методами машинного навчання та генетичними алгоритмами для побудови адаптивних систем оптимізації тестових конфігурацій у динамічних CI/CD-процесах.

Ключові слова: автоматизоване тестування; Pairwise Testing; оптимізація конфігурацій; комбінаторні методи; багатокомпонентні інформаційні системи; CI/CD.

1. ВСТУП

Постановка проблеми. Зі зростанням складності програмних систем кількість можливих конфігурацій для тестування зростає експоненційно, що робить повне тестування всіх варіантів неможливим [1]. Для систем із кількома десятками параметрів кількість комбінацій може сягати сотень тисяч, що унеможливорює перевірку в межах одного CI/CD-циклу.

Аналіз останніх досліджень і публікацій. Серед сучасних підходів до оптимізації тестових наборів найбільш поширеними є комбінаторні методи (Pairwise, Orthogonal Arrays), а також евристичні методи на основі генетичних алгоритмів [2]. Pairwise Testing залишається одним із найефективніших завдяки простоті реалізації та високому відсотку виявлення дефектів (до 90%) [3].

Мета публікації. Метою є демонстрація можливостей Pairwise Testing у зменшенні кількості тестових конфігурацій для багатокомпонентних ІС із збереженням високого рівня покриття та ефективною інтеграцією у CI/CD.

2. ТЕОРЕТИЧНІ ОСНОВИ

Pairwise Testing (тестування пар взаємодіючих параметрів) належить до класу комбінаторних методів генерації тестів, які базуються на принципі перевірки усіх можливих пар значень параметрів системи. Згідно з дослідженнями NIST [4], більшість дефектів у складних програмних системах виникає через взаємодію не більше двох параметрів одночасно. Це робить Pairwise Testing ефективним компромісом між повним перебором і вибіркоvim тестуванням.

Основною теоретичною основою Pairwise Testing є побудова покривних масивів (covering arrays, CA), що забезпечують покриття усіх можливих пар значень параметрів при мінімальній кількості тестових випадків. Формально, для системи з k параметрами та множинами значень $v_1, v_2, \dots, v_k$, необхідно побудувати множину тестів T , що мінімізує $|T|$ за умови, що кожна можлива пара (v_i, v_j) зустрічається принаймні один раз у T .

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для перевірки ефективності Pairwise Testing була створена власна система генерації конфігурацій тестування, реалізована на мові JavaScript (Node.js 20). Основою є модуль generateConfigs.js, який формує множину тестових конфігурацій для системи автоматизованого тестування багатокomпонентних інформаційних систем. Система підтримує режим генерації Pairwise — комбінаторна генерація всіх можливих пар параметрів для досягнення повного покриття пар взаємодій. У режимі pairwise застосовується обмежений пул значень (по два на кожен ключ), що зменшує комбінаторний простір, але зберігає критичне покриття пар параметрів. Для генерації використано рекурсивний алгоритм побудови декартового добутку, який формує мінімальну множину конфігурацій із забезпеченням усіх пар взаємодій параметрів[5].

Збереження результатів виконується у файлі ./inputs/configs.json, який потім використовується як вхідний набір для запуску тестів у CI/CD. Приклад експериментальної конфігурації наведено у Таблиці 1.

Таблиця 1

Приклад експериментальної конфігурації

Параметр	Опис параметра	Значення
frontend	Клієнтська частина (фреймворк інтерфейсу користувача)	React 18, React 17, Angular 17
backend	Серверна частина системи (платформа виконання бізнес-логіки)	Node.js 20, Python 3.11, Java 20
database	Система керування базами даних	MySQL, PostgreSQL, MongoDB
browser	Веб-браузер для UI-тестування	Chrome, Firefox, Edge
env	Середовище розгортання та виконання системи	production, staging, testing

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

В результаті тестування виявлено що кількість тестових конфігурацій скоротилась із 243 до 28–32 за рахунок Pairwise Testing; покриття всіх парних

комбінацій параметрів залишилось високим — 95–100%; час виконання тестового циклу у CI/CD зменшився на 60–70% порівняно з повним набором тестів; модуль автоматизованої перевірки забезпечив контроль покриття та зберігання результатів у репозиторії для подальшого аналізу.

Отримані результати підтвердили ефективність Pairwise Testing як інструменту зменшення обсягу тестів без втрати якості перевірки. Крім того, реалізація інтеграції з CI/CD продемонструвала можливість автоматизованого повторного використання тестових наборів, підвищення швидкодії та стабільності процесу тестування.

5. ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Запропонований підхід із застосуванням Pairwise Testing дозволяє суттєво скоротити обсяг тестових конфігурацій багатокomпонентних ІС (до 12–15% від повного набору) при збереженні високого рівня покриття критичних пар параметрів (95–100%).

Методи комбінаторної оптимізації забезпечують ефективне управління експоненційно зростаючими конфігураційними просторами, а інтеграція у CI/CD дозволяє автоматизувати повторювані цикли тестування та збирати аналітичні дані про ефективність. Використання модульного підходу до генерації конфігурацій гарантує мінімізацію необхідної кількості тестів без втрати якості покриття.

Перспективні напрямки подальших досліджень включають: поєднання Pairwise Testing із генетичними алгоритмами для багаторівневої оптимізації, застосування штучного інтелекту для прогнозування дефектів і автоматичний аналіз результатів тестування в реальному часі.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] О. Р. Melnyk, *Yakist' prohramnoho zabezpechennya ta testuvannya: bazovyy kurs*, Ternopil: TNEU, 2019, 240 p.
- [2] A. S. Avramenko, V. S. Avramenko, H. V. Kosenyuk, *Testuvannya prohramnoho zabezpechennya*, Cherkasy: ChNU im. Bohdana Khmel'nyts'koho, 2017, 284 p.
- [3] O. L. Lyakhov, O. O. Borodina, *Metody testuvannya i otsinky yakosti prohramnoho zabezpechennya*, Poltava: PoltNTU, 2015, 372 p.
- [4] D. R. Kuhn, D. R. Wallace, and A. M. Gallo, "Software fault interactions and pairwise testing," *IEEE Transactions on Software Engineering*, vol. 30, no. 6, pp. 418–421, 2004.
- [5] R. Bryce and C. Colbourn, "A framework of greedy methods for constructing interaction test suites," in *Proceedings of the 27th International Conference on Software Engineering*, 2005, pp. 146–155.

Ірина Бенцал

Магістр комп'ютерних наук

Тернопільський національний технічний університет імені Івана Пулюя, Тернопіль, Україна

bentsaluruna@gmail.com

Леся Дмитроца

К.т.н, доцент, доцент кафедри комп'ютерних наук

Місце роботи: Тернопільський національний технічний університет імені Івана Пулюя, кафедра комп'ютерних наук, Тернопіль, Україна

ORCID 0000-0003-2583-3271

dmytotsa.lesya@gmail.com

РОЗРОБКА СИСТЕМИ МОНІТОРИНГУ ТА ПРОГНОЗУВАННЯ РОЇННЯ БДЖІЛ З ВИКОРИСТАННЯМ ІОТ-ТЕХНОЛОГІЙ

Анотація. В роботі розглядаються актуальні проблеми передчасного відслідковування стану роїння бджолосімей за допомогою впровадження IoT-технологій у вулик. У дослідженні проаналізовано сучасні підходи до автоматизації моніторингу стану бджолосімей, визначено їхні обмеження та потребу у впровадженні інтелектуальних алгоритмів раннього попередження.

Для побудови системи використано мікрокомп'ютер Raspberry Pi як центральний вузол збору даних, цифрові сенсори DS18B20 (для вимірювання температури) та SHT31-D (для вимірювання вологості). Збір, передавання та візуалізація даних реалізовані за допомогою відкритих IoT-платформ HoneyPi та ThingSpeak, які забезпечують збереження даних у хмарному середовищі та подальший аналітичний обробіток. На основі зібраних часових рядів розроблений алгоритм, який аналізує зміну мікрокліматичних параметрів протягом останніх 24–36 годин. Алгоритм враховує динаміку зростання температури, амплітуду добових коливань і співвідношення між температурою та вологістю. При стійкому підвищенні температури та вологості з одночасним зменшенням амплітуди система формує попередження про можливе роїння за 24–48 годин.

Експериментальні дослідження із застосуванням реальних даних підтвердили працездатність алгоритму та точність прогнозу на рівні понад 80 %. Запропонований підхід дає змогу своєчасно виявляти підготовку до роїння, мінімізувати втрати бджіл, підвищити продуктивність і створює основу для подальшої розробки інтелектуальної системи.

Ключові слова: IoT-технології, розумний вулик, температура, вологість.

1. ВСТУП

Бджільництво є однією з ключових галузей сільського господарства, адже саме бджоли забезпечують запилення понад 80 % ентомофільних культур. Одним із найвагоміших чинників втрати робочих бджіл залишається роїння — природний процес поділу сім'ї, який у бджільництві розглядається як негативне явище.

Постановка проблеми. Основною проблемою є неможливість прогнозування моменту початку роїння за 24-36 годин.

Аналіз останніх досліджень і публікацій. Зазвичай пасічники орієнтуються на поведінкові ознаки або зовнішні спостереження, однак такі методи є суб'єктивними і потребують постійної присутності людини біля пасіки [1]. Класичні системи моніторингу фіксують лише факт змін після роїння (зменшення ваги, різкі стрибки температури), що не дає змоги запобігти втратам [2]. Таким чином, розроблення системи автоматизованого моніторингу та прогнозування роїння на основі даних датчиків і алгоритмів машинного аналізу є актуальним завданням для підвищення ефективності сучасного бджільництва.

Мета публікації. Метою роботи є створення інформаційно-вимірювальної системи, що забезпечує безперервний моніторинг мікроклімату вулика (температури та вологості) і прогнозує можливість роїння за 24–48 годин до його настання. Для

досягнення мети необхідно було проаналізувати існуючі методи автоматизації моніторингу стану бджолосімей і визначити їх недоліки, вибрати оптимальні сенсори для вимірювання температури та вологості, розробити апаратну архітектуру на базі мікрокомп'ютера Raspberry Pi з відкритими IoT-платформами, проаналізувати готові програмні рішення для збору та наступних досліджень, створити алгоритм раннього попередження роїння, що аналізує часові ряди даних, а саме температуру та вологість, провести експериментальні випробування системи і оцінити її ефективність, подати можливі варіанти покращення даної системи.

У результаті виконання поставлених завдань було розроблено концепцію інформаційно-вимірювальної системи для прогнозування роїння бджолосімей на основі аналізу температури та вологості вулика.

2. ТЕОРЕТИЧНІ ОСНОВИ

Основні фізіологічні ознаки підготовки до роїння проявляються у зміні температурно-вологісних параметрів усередині вулика. У період перед роїнням спостерігається підвищення середньої температури на 1–2 °C, збільшення вологості повітря та зменшення добових амплітуд коливань.

Теоретичною основою побудови системи прогнозування є концепція Інтернету речей (IoT). У контексті бджільництва це дозволяє організувати безперервний моніторинг мікроклімату вулика через датчики температури та вологості, підключені до мікрокомп'ютера (наприклад, Raspberry Pi).

На основі отриманих даних формуються часові ряди, які піддаються математичній обробці, що дозволяють виявити закономірності змін і спрогнозувати момент роїння за 24–48 годин до його настання.

Застосування відкритих IoT-платформ дає змогу інтегрувати дані у хмарне середовище, здійснювати аналітику в режимі реального часу та формувати повідомлення про критичні відхилення параметрів. Таким чином, теоретичні положення дослідження ґрунтуються на поєднанні принципів біомоніторингу, аналізу часових рядів і технологій Інтернету речей, що створює передумови для розробки інтелектуальної системи прогнозування роїння бджолосімей.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для практичної реалізації було обрано мікрокомп'ютер Raspberry Pi 3B+ як центральний вузол обробки. Система оснащена цифровими сенсорами: DS18B20 – вимірювання температури в межах –10...+85 °C із точністю ± 0.5 °C та SHT31-D – вимірювання вологості (0–100 % RH) і температури з високою стабільністю. Вибір цих датчиків зумовлений їхньою сумісністю з Raspberry Pi, низьким енергоспоживанням та точністю показників. Збір даних відбувається з інтервалом 5-10 хвилин. Передавання інформації здійснюється через Wi-Fi-модуль до хмарного сервісу ThingSpeak, який підтримує REST API та MATLAB Analytics для подальшої обробки даних.

Програмна частина реалізована на основі відкритої платформи HoneypI, яка виконує автоматичне зчитування сенсорів через GPIO, локальне резервне збереження у форматі CSV, а також має можливість налаштування сповіщень при перевищенні порогів.

Крім ThingSpeak, розглядалися інші open-source платформи – HiveEyes, BeeMoS, BEEP-framework, які дозволяють підключати додаткові вулики в єдину мережу та масштабувати систему [3].

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

За результатами аналізу було виявлено, що перед початком роїння спостерігається стабільне підвищення внутрішньої температури на 1.5-2 °C та зростання вологості на 5-7 %, при цьому добова амплітуда коливань зменшується в 1.5-2 рази. Запропонована система виявила роїння з точністю ~82 % при середньому часі попередження 36 годин. Результати показали високу кореляцію між ростом індексу та фактичним моментом виходу рою. Порівняно з класичним пороговим підходом (температура > 35 °C), запропонований метод зменшив кількість хибних спрацювань на ≈ 40 %.

Розроблена система може бути впроваджена у приватних і промислових пасіках для раннього попередження про роїння, зниження трудовитрат на ручний контроль, ведення статистики параметрів вулика, а також формування бази даних для подальшого машинного навчання. Крім того, система може стати частиною платформи, що поєднує моніторинг, прогнозування, управління мікрокліматом і віддалене сповіщення через мобільний додаток.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У результаті проведеного дослідження було проаналізовано сучасні підходи до моніторингу стану бджолосімей та розроблено концепцію системи прогнозування роїння на основі IoT-технологій. Обґрунтовано вибір сенсорів, а також платформ ThingSpeak і HoneyPi для збору, зберігання та візуалізації даних. Отримані результати підтверджують можливість практичної реалізації інтелектуальної системи моніторингу пасік (Smart Apiary System), яка здатна підвищити ефективність роботи пасічників, мінімізувати втрати бджіл і сприяти розвитку цифрового бджільництва в Україні.

Перспективи подальших досліджень:

- удосконалення алгоритму прогнозування шляхом використання методів машинного навчання (нейронних мереж, класифікаторів);
- розширення системи додатковими сенсорами (CO₂, ваги, звуку);
- створення мобільного додатка для віддаленого сповіщення пасічників.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

1. A. Zaman, "A framework for better sensor-based beehive health," *Computers and Electronics in Agriculture*, vol. 211, p. 108069, 2023. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0168169923002946>
2. S. Kontogiannis, "Beehive Smart Detector Device for the Detection of Critical Conditions That Utilize Edge Device Computations and Deep Learning Inferences," *Sensors*, vol. 24, no. 16, p. 5444, 2024. [Online]. Available: <https://www.mdpi.com/1424-8220/24/16/5444>
3. N. Al-Dhuraibi, A. Zaytar, H. Alzaabi, and A. Zaki, "Internet of Things Low Cost Monitoring BeeHive System Using Wireless Sensor Network," *ResearchGate*, 2019. [Online]. Available: https://www.researchgate.net/publication/337793861_Internet_of_Things_Low_Cost_Monitoring_BeeHive_System_using_Wireless_Sensor_Network.

Олексій Коваленко

Доктор технічних наук, професор кафедри комп'ютерних систем, мереж та кібербезпеки
НУБіП України, Факультет інформаційних технологій, Кафедра комп'ютерних систем, мереж та
кібербезпеки, Київ, Україна
<https://orcid.org/0000-0002-9639-3544>
O.Kovalenko@nubip.edu.ua

Микола Богдюк

Аспірант, асистент кафедри комп'ютерних систем, мереж та кібербезпеки
НУБіП України, Факультет інформаційних технологій, Кафедра комп'ютерних систем, мереж та
кібербезпеки, Київ, Україна
<https://orcid.org/0009-0000-4743-4023>
M.Bogdiuk@nubip.edu.ua

ОГЛЯД ДОСЛІДЖЕНЬ ФЕНОМЕНУ ВТОМИ ВІД ЗАПИТІВ НА ДОСТУП

Анотація. У статті розглянуто феномен втоми від запитів на дозвіл доступу до ресурсів, який виникає внаслідок надмірної кількості запитів доступу до ресурсів у сучасних інформаційних системах. Показано, що часті сповіщення про необхідність підтвердження згоди або дозволу призводять до емоційного виснаження користувачів, формують байдужість до повідомлень безпеки та спричиняють автоматичне прийняття рішень, що знижує рівень захисту даних. Метою дослідження є узагальнення сучасних наукових підходів до вивчення цього явища та визначення шляхів його урахування в моделях і засобах системного програмного забезпечення для безпечних операційних середовищ.

Методологічну основу становить аналіз чотирьох актуальних робіт, що охоплюють різні наукові підходи: поведінкові експерименти, дослідження довіри до технологій, конфігураційний аналіз факторів втоми та емпіричне вивчення інтерфейсів інформованої згоди у мобільних застосунках. Результати демонструють, що втома від запитів є багатовимірним феноменом, що формується під впливом когнітивних, емоційних і технічних чинників. Встановлено, що частота запитів дозволів, складність політик приватності, низька цифрова грамотність і перевантаження користувача інформацією суттєво підвищують ризик формування втоми та автоматичної поведінки.

Пропонується інтегрувати у програмне забезпечення принципи контекстно-залежного управління дозволами та адаптивної взаємодії користувача із системою. Це включає динамічне регулювання інтенсивності запитів, групування подібних дозволів, застосування профілю довіри до додатків і використання елементів навчальної підтримки. Врахування феномену втоми від запитів у проєктуванні операційних середовищ дозволить підвищити реальну ефективність політик безпеки, зменшити когнітивне навантаження користувачів і забезпечити баланс між прозорістю, зручністю та рівнем захисту.

Ключові слова: втома від запитів, обмеження доступу, адаптивні дозволи, когнітивне навантаження.

1. ВСТУП

У сучасних інформаційних системах користувач постійно стикається з численними запитами на дозвіл доступу до ресурсів: від мобільних додатків до операційних систем і браузерів. Така повторюваність, за спостереженнями дослідників, призводить до явища, відомого як втома від запитів на дозвіл доступу до ресурсів ("permission fatigue" або "privacy fatigue", надалі "втома від запитів") – стану емоційного виснаження і байдужості, коли користувач надає дозволи автоматично, не аналізуючи наслідки. Це створює парадокс безпеки: чим більше системи запитують дозвіл, тим нижчий рівень усвідомленого контролю. Для розробників програмного забезпечення ця проблема стає критичною, оскільки надмірна кількість запитів підриває довіру користувача та ефективність захисту даних.

Постановка проблеми. Сучасні системи безпеки орієнтовані на формальне отримання згоди, а не на забезпечення усвідомленої взаємодії. У результаті користувачі

часто надають доступ автоматично, не розуміючи обсяг і наслідки дозволів, таким чином збільшення кількості захисних механізмів фактично знижує рівень реальної безпеки.

У науковому та прикладному аспектах актуальність проблеми зумовлена відсутністю інтегрованих моделей, що враховують психологічну реакцію користувача на запити дозволів. Більшість підходів до системного програмного забезпечення фокусуються на технічних засобах контролю доступу, ігноруючи фактор довіри та втому користувача. Це вимагає розроблення нових методів керування дозволами, здатних поєднати безпеку, ергономіку й адаптивність.

Метою публікації є аналіз сучасних наукових підходів до вивчення феномену втоми від запитів у сфері інформаційної безпеки та визначення можливостей його врахування при моделюванні безпечного операційного середовища. Зокрема, розглянуто результати емпіричних і теоретичних досліджень, що формують сучасне розуміння механізмів виникнення та впливу втоми від питань на поведінку користувача.

Аналіз останніх досліджень і публікацій. Робота ґрунтується на порівняльному аналізі чотирьох актуальних наукових джерел, у яких використано різні методологічні підходи:

- дослідження інтерфейсів інформованої згоди в мобільних застосунках [1];
- експериментальне дослідження поведінкових рішень користувачів [2];
- аналіз факторів довіри та цифрової грамотності [3];
- конфігураційний аналіз причин виникнення втоми від запитів на основі теорії Stressor-Strain-Outcome (стресор – напруження – зміна поведінки) [4].

Згідно вказаних публікацій, втома від запитів є не лише поведінковим явищем, а й системним ефектом, який виникає через неузгодженість між логікою безпеки та ергономікою інтерфейсу. Зокрема, згідно дослідження [3], рівень довіри до технологій впливає на зв'язок між втомою та поведінкою користувачів: перевантажені попередженнями, вони сприймають систему як "настирливу" і діють автоматично. У результаті психологічний чинник перетворюється на структурну вразливість, що нівелює усвідомлену згоду і суперечить принципу залученості людини до безпеки.

Згідно з дослідженням [4], толерантність користувачів до запитів залежить від контексту. Наприклад, доступ до камери оцінюється критичніше, ніж до геолокації. Це вказує на потребу у контекстно-залежних моделях дозволів, які враховують тип ресурсу, джерело запиту та частоту погоджень, формуючи гнучкий шар взаємодії безпеки з мінімальним когнітивним навантаженням.

Крім того, дослідження [1] підкреслює, що традиційна парадигма інформованої згоди стає неефективною в умовах втоми від запитів, оскільки користувачі перестають сприймати повідомлення про дозвіл як джерело цінної інформації. Автори пропонують формувати профіль довіри для кожного застосунку, адаптуючи інтенсивність запитів і пояснень. Наприклад, посилені перевірки для нових або підозрілих програм. Такий підхід відповідає ідеям ризик-орієнтованого управління доступом, де частота перевірок визначається рівнем ризику, а не формальними вимогами.

Важлива особливість виявлена в [2]: користувачі з високим рівнем технічної обізнаності частіше усвідомлюють ризики та демонструють нижчу схильність до втоми від запитів. Це вказує на доцільність включення навчальних або роз'яснювальних механізмів у системи керування дозволами – наприклад, коротких пояснень щодо наслідків надання доступу, піктограм ризику чи відкладених сповіщень.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Узагальнюючи результати усіх чотирьох досліджень, можна запропонувати нову парадигму для розроблення системного ПЗ – інтелектуальні механізми запитів, що поєднують машинне навчання, профілювання поведінки користувача та контекстну оцінку ризику. Така система могла б прогнозувати, коли користувач, ймовірно, перебуває у стані втоми (наприклад, через швидкі повторні кліки “Дозволити”), і тимчасово змінювати стратегію взаємодії – надавати зведену інформацію, групувати запити або пропонувати відкладене підтвердження. Це відкриває перспективу формування операційних середовищ, у яких людський чинник не є слабкою ланкою, а стає частиною адаптивної системи прийняття рішень.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Феномен втоми від запитів стає дедалі актуальнішим у контексті розроблення безпечного операційного середовища. Його ігнорування призводить до зниження ефективності навіть найкращих механізмів безпеки, оскільки користувачі перестають критично сприймати запити доступу. Для системного програмного забезпечення перспективним напрямом є створення інтелектуальних систем дозволів, які динамічно регулюють інтенсивність взаємодії з користувачем, ґрунтуючись на довірі, контексті та поведінкових показниках. Врахування людського чинника й психологічного виснаження підвищує не лише зручність, а й реальний рівень інформаційної безпеки.

ПОСИЛАННЯ

[1] M. Chen and M. Chen, “Trust, Privacy Fatigue, and the Informed Consent Dilemma in Mobile App Privacy Pop-Ups: A Grounded Theory Approach,” *Journal of Theoretical and Applied Electronic Commerce Research*, vol. 20, no. 3, p. 179, 2025. doi: <https://doi.org/10.3390/jtaer20030179>.

[2] Y. Liu, Reza Ghaiumy Anaraky, H. Aly, and K. Byrne, “The Effect of Privacy Fatigue on Privacy Decision-Making Behavior,” *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, vol. 67, no. 1, pp. 2428–2433, 2023. doi: <https://doi.org/10.1177/21695067231193670>.

[3] R. Luzsa and S. Mayr, “Links Between Online Privacy Fatigue, Technology Attitudes and Sociodemographic Factors in a German Population Sample,” *Proceedings of Mensch und Computer 2022 (MuC '22)*. Association for Computing Machinery, New York, NY, USA, pp. 360–364, 2022, doi: <https://doi.org/10.1145/3543758.3547540>.

[4] W. Wang, Q. Wu, D. Li, and X. Tian, “An exploration of the influencing factors of privacy fatigue among mobile social media users from the configuration perspective,” *Scientific Reports*, vol. 15, art. 427, 2025. doi: <https://doi.org/10.1038/s41598-024-84646-z>.

Цзян Сюе

PhD candidate

1.Bengbu University, Bengbu City, Anhui Province, China

2.National University of Life and Environmental Sciences of Ukraine, Kiev, Ukraine

ORCID ID 0009-0000-1676-2331

jx1283@163.com

Lakhno Valerii

Doctor of Technical Sciences, Professor, Department of Computer Systems, Networks and Cybersecurity

National University of Life and Environmental Sciences of Ukraine, Kiev, Ukraine

ORCID ID 0000-0001-9695-4543

lva964@nubip.edu.ua

SURVEY ON THE EVOLUTION AND KEY TECHNOLOGIES OF DIFFERENTIAL DISTINGUISHER DESIGN

Abstract. The differential distinguisher is pivotal in differential cryptanalysis, distinguishing authentic from random differential distributions. This paper surveys its evolutionary trajectory and key technologies, covering the paradigm shift from traditional manual-driven to neural network-aided methods. Traditional distinguishers rely on Differential Distribution Tables and manual high-probability paths, offering strong interpretability but poor efficiency and adaptability to complex ciphers like Speck and Simon. Neural network-aided designs innovate in input feature engineering and network architectures. The paper identifies challenges including poor cross-cipher generalisation and insufficient interpretability, proposing future directions like adaptive design and lightweight optimisation. It serves as a comprehensive reference for lightweight cipher security analysis.

Keywords: differential distinguisher; cryptanalysis; Neural network; Lightweight ciphers

1 INTRODUCTION

The differential distinguisher is a core component of differential cryptanalysis, whose primary function is to distinguish between the "authentic differential distribution" (generated by encrypting plaintext pairs with fixed differences) and "random differential distribution" (generated by encrypting random plaintext pairs) output by cryptographic algorithms, providing a critical basis for subsequent key recovery attacks [1]. Since Biham and Shamir proposed DES cryptanalysis based on differential distinguishers in 1991 [1], the design of differential distinguishers has undergone a paradigm shift from "traditional manual-driven" to "neural network-aided": traditional methods rely on manual exploration of high-probability differential paths, which are inefficient and struggle to handle high-round or complex-structure ciphers (e.g., lightweight block ciphers Speck and Simon); in 2019, Gohr introduced Deep Residual Networks (ResNet) into differential distinguisher design for the first time [3], breaking through traditional limitations and paving the way for a new direction in neural network-aided cryptanalysis. Focusing on the core design of differential distinguishers, this paper systematically organises their evolutionary context, key technologies, performance characteristics, and future challenges, providing a reference for the security analysis of lightweight ciphers.

2 EVOLUTION OF DIFFERENTIAL DISTINGUISHER DESIGN

2.1 Traditional differential distinguishers: based on DDT and high-probability paths

The core of traditional differential distinguisher design lies in the Differential Distribution Table (DDT) and high-probability differential paths. The DDT quantifies the probability that an input difference is transformed into an output difference via the cipher's round function, serving as a fundamental tool for distinguishing between authentic and random distributions [1]. The steps to construct an effective distinguisher are as follows:

Traditional design offers advantages of intuitive logic and strong interpretability, but suffers from significant limitations: first, it relies on manual experience to explore high-probability paths, leading to poor adaptability to complex ARX (Addition-Rotation-XOR) structures (e.g., Speck's modular addition + rotation + XOR); second, differential probability decays exponentially with increasing rounds, resulting in a sharp decline in discrimination performance (e.g., the accuracy of traditional distinguishers for 7-round Speck32/64 is less than 60% [3]); third, it has high data complexity, requiring a large number of plaintext-ciphertext pairs to verify differential paths.

2.2 Neural Network-Aided Differential Distinguishers: Design Innovation Driven by Features and Data

The introduction of neural networks has restructured the design logic of differential distinguishers—shifting from "manual path-finding" to "model-based feature learning". Core improvements focus on input feature design and network architecture adaptation, with specific evolutionary paths as follows:

2.2.1 Input feature design: from single-dimension to multi-dimension enhancement

Input features serve as the "information source" for neural network distinguishers, and their design directly determines feature utilisation and discrimination performance. Mainstream designs include:

Single Ciphertext Pair (SCP) Design [3]: A foundational paradigm proposed by Gohr, which takes raw bits of a single ciphertext pair as input. It first verified the effectiveness of neural networks on Speck32/64 (92.9% accuracy for 5 rounds). However, it only utilises surface-level bit information, leading to a drop in accuracy to 61.6% for high rounds (7 rounds).

Multiple Ciphertext Pairs (MCP) Design [4]: By concatenating k ciphertext pairs into a single sample, it leverages inter-sample derived features (e.g., differential statistical correlation) to enhance discriminability. For 5-round Speck32/64, accuracy reaches 99.99%, and for 7 rounds, it increases to 64.93%. Meanwhile, it reduces data complexity by 30% through a data reuse strategy.

Multi-Round Differential (MRMSD) Design [6]: Integrating the multi-round differential characteristics of Markov ciphers, it introduces differential information from the penultimate round via partial decryption with random pseudo-round keys, and converts ciphertext pairs into differential form. This design strengthens cryptographically core features, achieving 94.11% accuracy for 7-round Speck32/64 and enabling effective discrimination of 12-round Simon48/96 for the first time (61.59% accuracy).

Masked Output Distribution Table (M-ODT) Design [5]: To address the "black-box" issue, it compresses the DDT via masking, mapping input features to explicit differential probabilities. While maintaining performance (92.3% accuracy for 5-round Speck32/64), it reveals key differential bits relied on by the model, achieving a breakthrough in interpretability.

Random Key Multi-Cipher Pairs (RKMP) Design [7]: It introduces random keys for additional encryption of ciphertext pairs, enhancing feature diversity and improving generalisation. For 7-round Speck32/64, accuracy reaches 92.03%, and it enables effective discrimination of 8 rounds for the first time (63.32% accuracy), adapting to the needs of high-round cryptanalysis.

2.2.2 Network architecture adaptation: from ResNet to attention enhancement

Neural network architectures must match input features to maximise performance: ResNet mitigates the vanishing gradient problem via residual connections, adapting to deep feature extraction for multi-round differences [3]; 2D-CNN utilises local convolution kernels

to capture correlations between multiple ciphertext pairs [4]; attention mechanisms focus on key features and suppress noise introduced by random keys in RKMP [7]. For example, ResNet integrated with attention mechanisms enables RKMP to achieve 99.5% accuracy for 9-round Simon32/64, an improvement of 0.4% compared to MRMSD [7].

3 KEY DESIGN ELEMENTS AND PERFORMANCE COMPARISON

3.1 Core design elements

The design of differential distinguishers must balance three core objectives: feature effectiveness, data efficiency, and interpretability. Table 1 compares the core features and performance of mainstream designs:

Table 1

Performance indicators for various cryptanalytic designs

Design Type	Core Features	5-Round Speck32/64 Accuracy	7-Round Speck32/64 Accuracy	Maximum Supported Rounds (Simon32/64)
Traditional DDT-based	Single-round high-probability paths	~85%	~55%	6 rounds
SCP [3]	Raw bits of single ciphertext pair	92.9%	61.6%	8 rounds
MCP [4]	Concatenation of multiple ciphertext pairs	99.99%	64.93%	9 rounds
MRMSD [6]	Multi-round differences + differential form	-	94.11%	11 rounds
RKMP [7]	Random keys + attention mechanism	-	92.03%	9 rounds (99.5% accuracy)

3.2 Key performance conclusions

Expanding feature dimensions significantly improves performance: Multi-round differential (MRMSD) design increases accuracy by 32.5% for 7-round Speck32/64 compared to single-round feature (SCP) design;

Differential form outperforms raw bits: MRMSD improves accuracy by 29.18% for 7-round Speck32/64 compared to MCP;

Interpretability and performance can be balanced: While maintaining accuracy comparable to SCP, M-ODT reveals key differential bits, providing a theoretical basis for cryptanalysis [5].

4 CONCLUSIONS

The design of differential distinguishers has shifted from traditional manual-driven to neural network-aided feature and data-driven approaches. The core evolutionary logic is to "expand feature dimensions, improve data efficiency, and balance performance with interpretability". Designs such as MRMSD and RKMP have broken through the discrimination bottleneck for high-round ciphers, while M-ODT provides a new approach to interpretability. Future research should focus on cross-cipher generalisation, lightweight design, and interpretability, enabling differential distinguishers to play a greater role in the security verification of lightweight ciphers and providing stronger support for cryptographic security in scenarios such as the Internet of Things and embedded systems.

REFERENCES

- [1] E. Biham and A. Shamir, "Differential cryptanalysis of DES-like cryptosystems," *Journal of Cryptology*, vol. 4, no. 1, pp. 3–72, 1991.
- [2] X. Lai, J. L. Massey, and S. Murphy, "Markov ciphers and differential cryptanalysis," in *Workshop on the Theory and Application of Cryptographic Techniques - EUROCRYPT 1991*, 1991, pp. 17–38.
- [3] A. Gohr, "Improving attacks on round-reduced Speck32/64 using deep learning," in *Advances in Cryptology - CRYPTO 2019*, 2019, pp. 150–179.
- [4] Y. Chen, Y. Shen, H. Yu, and S. Yuan, "A new neural distinguisher considering features derived from multiple ciphertext pairs," *The Computer Journal*, vol. 66, no. 6, pp. 1419–1433, 2023.
- [5] A. Benamira, D. Gerault, T. Peyrin, and Q. Q. Tan, "A deeper look at machine learning-based cryptanalysis," in *Advances in Cryptology - EUROCRYPT 2021*, 2021, pp. 805–835.
- [6] J. S. Liu, J. J. Ren, S. Z. Chen, et al., "Improved neural distinguishers with multi-round and multi-splicing construction," *Journal of Information Security and Applications*, vol. 74, pp. 103461, 2023.
- [7] X. Jiang, M. Li, K. Makulov, et al., "Enhanced neural differential distinguisher for Speck32/64 using attention mechanisms and multi ciphertext inputs," *Informatica*, vol. 49, pp. 197–210, 2025.

Nazym Sabitova

Doctoral student of L. N. Gumilyov Eurasian National University,
Kazakhstan, Astana.

ORCID ID <https://orcid.org/0009-0001-2531-5653>

Email: sab.nazym1907@gmail.com

DEVELOPING A SYSTEM FOR ONTOLOGY-BASED DESIGN OF ITC ELECTRONIC COURSES

Abstract. A methodology for modifying computer ontologies (CO) in the field of information and communication technologies (ICT) is proposed by introducing into the formal CO a description of learning process technologies for high school or out-of-school institutions, as well as for professional retraining of specialists, which allows preserving and reusing the collective experience of education of this educational institution.

It is noted that the tool of the system of ontologized design of electronic courses (SODEC) partially implements this methodology, which automates the development of electronic courses and/or electronic textbooks (EC), reduces the costs and time of preparation of ET, ensures the relevance of EC to the current state of knowledge in the subject area (KSA) of ICT, accumulates personal materials of the teacher. It is concluded that for more effective use of the proposed methodology, it is necessary to develop a similar SODEC toolkit, taking into account the introduction of the corresponding COs into the ontology description language (OWL).

Keywords: electronic courses, electronic textbooks, computer ontologies, system of ontologized design of electronic courses, subject area, ontology description language.

1. INTRODUCTION

With the digitalization of modern society and the rapid development of information and communication technologies (ICT), which today are deeply integrated into almost all spheres of human activity, there has been a need to develop innovative methods and training tools at all levels of training. These methods and training tools should be focused on training specialists who will be able not only to process and analyze information materials, but also to effectively solve problems related to decision support in industrial design, classification of scientific and technical documentation, integration of information services of partner companies, etc. The beginnings of such competencies should be laid in high school, or in out-of-school institutions, for example, engineering and technical centers for youth, as well as during professional retraining of specialists.

The problem statement. To enhance the CO KSA of ICT, the study addresses the development of a methodology for designing CO processes that describe learning technologies for schools, training centers, and professional retraining, ensuring the preservation and reuse of collective educational experience.

Analysis of recent studies and publications. According to a number of experts, the relevance of research in this direction is due to such factors [1], [2].

Firstly, as in the general task of creating electronic courses and/or electronic textbooks (hereinafter we use the abbreviation EC), part of the knowledge of the subject area (KSA) of ICT and its relevance may be lost due to the increase in new information in printed form. This is not least due to the rapid pace of ICT development that is taking place today. Modern

schoolchildren should not only have knowledge in the field of ICT, but also in their subsequent activities become full-fledged specialists in their use.

Secondly, ECs KSA of ICT for high school or non-school institutions, as well as for professional retraining of specialists, have their own specifics [3], [4].. We are talking, in particular, about increasing the cognitive quality of EC, increasing the use of multimedia tools, interactive teaching methods, etc. Note that visual-figurative components are especially important in the worldview of schoolchildren, which encourages teachers to create EC in a visually expressive form, for example, by using multimedia.

Thirdly, the widespread use of personal computers (PCs) and other mobile gadgets (smartphones, tablets, laptops) in schools and in out-of-school institutions and professional retraining institutions. Moreover, in some cases, only PCs are not always suitable for the rapid processing of volumetric layers of material in individual ICT courses and disciplines. For example, when it comes to courses or topics "Mobile devices and networks", "Mobile Internet security", etc. This leads to a specific CO, characterized by the introduction of CO learning processes for a given KSA [5], [6]..

The article's goal. The research involved a systematic approach to KSA, structuring and the principle of hierarchy to KSA information, as well as an ontological approach to the design and development of EC, which is one of the applications of ontological engineering, in particular, the processing of terms and concepts contained in CO.

2. THE THEORETICAL BACKGROUNDS

The research involved a systematic approach to KSA, structuring and the principle of hierarchy to KSA information, as well as an ontological approach to the design and development of EC, which is one of the applications of ontological engineering, in particular, the processing of terms and concepts contained in CO.

3. RESEARCH METHODS

The relevance of this research is driven by three factors: the rapid growth of ICT knowledge requiring updated ECs to prevent knowledge loss; the need to enhance cognitive quality and visual interactivity in ECs for schools and training programs; and the widespread use of PCs and mobile devices, which demands specialized CO processes for modern ICT topics

4. THE RESULTS AND DISCUSSION

It is necessary to develop a process for designing the CO KSA of ICT, exemplified by the "Relational Databases" module, to automate EC/ET creation while integrating institutional experience and enabling ICT teachers to design content that connects database theory with programming and network fundamentals, considering the profile and needs of specialized school or out-of-school learners. In its formal form, the model of the designed CO will look the same as the general model of the CO KSA (1) [1, 16]:

$$S = \langle P, A, X \rangle, \quad (1)$$

Where corresponding sets $P = \{p_i\}, i = (1, n)$ – processes for constructing EC and/or ET;

$A = \{A_j\}, j = (1, m), m \geq n$ – algorithms that implement these processes; X – entities (concepts) describing KSA.

However, in our case, the construction of a CO KSA with a description of processes will be simplified due to the absence of the need for semantic machine analysis of concepts.

However, the process of describing multimedia processes will become more complex. Although this complexity is present only at the initial stage of building a CO KSA.

SODEC is based on ontological knowledge, provides effective machine processing of the ontological description of knowledge, and provides automated filling of the EC with the necessary (according to the teacher's experience) structural elements of educational material. Such materials include multimedia objects. However, often these elements are not structured, not collected into an ontology, and there is no way to use them for other ECs.

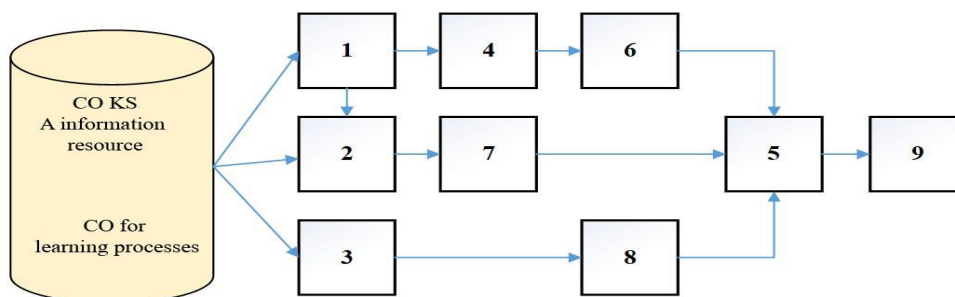


Figure 1. Scheme for preparing EC and/or ET (using the example of ICT (theme “Databases”))

The CO KSA includes concepts, concept descriptions, and CO descriptions of learning process algorithms. The author of the EC may decide to adapt descriptions or learning algorithms (for example, in order to simplify perception), which requires an appropriate mechanism in the EC compilation toolkit.

REFERENCES

1. Tikhonov, Yu.L. (2017). Tool for formation of electronic course. Podilian Bulletin: Agriculture, Engineering, Economics, (27), 220–225. Retrieved from <http://pb.pdatu.edu.ua/article/view/135318>
2. Tikhonov, U., Lakhno, V., Skliarenko, E., Stepanenko, O., Dvirnyi, K. Development of on Tological Approach in E-learning When Studying Information Technologies (2016) Eastern-European Journal of Enterprise Technologies, 5 (2), pp. 13-20.
3. Palagin, A., Kryvyi, S., & Petrenko, N. (2012). Ontological methods and means of processing subject knowledge: monograph. Volodymyr Dahl East Ukrainian National University, Luhansk. – 324 p.
4. Uschold M. Ontologies: Principles, Methods and Applications / M. Uschold, M. Gruninger. Knowledge Engineering Review 11(2), 1996. – PP. 93–136.
5. How E-learning Works by Lee Ann Obringer. – Available at <http://communication.howstuffworks.com/elearning.htm>.
6. E-learning 2.0 by Stephen Downes, National Research Council of Canada October 17, 2005. – Available at <http://www.elearnmag.org/subpage.cfm?article=29-1§ion=articles>

Volodymyr Nazarenko

PhD, Associate Professor, Computer Systems, Networks and Cybersecurity Department

National University of Life and Environmental Sciences of Ukraine, Kyiv, Ukraine

ORCID: 0000-0002-7433-2484

volodnz@nubip.edu.ua

VIDEO GAME SECURITY IN 2025: HYBRID ANTI-CHEAT ARCHITECTURES, BEHAVIORAL ANALYTICS, AND LESSONS FOR REAL-TIME CYBER-PHYSICAL SYSTEMS

Abstract. Modern online video games are real-time, data-intensive distributed systems that must detect and deter adversaries who continuously evolve their tools and tactics. At stake are fairness, revenue, and player trust. This paper systematizes hybrid anti-cheat architectures that combine client integrity controls, server-side validation, behavioral analytics, and privacy-aware AI. We argue that video game security is a mature “living laboratory” for real-time anomaly detection under latency constraints and that its techniques transfer to adjacent domains (e.g., digital twins and industrial simulations). We first outline the problem: memory injection, protocol and telemetry manipulation, botting/scripting, identity abuse, and targeted denial-of-service attacks. We developed a reference threat model, a layered control framework, and a small benchmark of statistical and machine-learning detectors (Z-score, SVM, LSTM) over synthetic but realistic action-telemetry streams. Results show that LSTM-based detectors can achieve detection accuracy exceeding 90% with modest (<15 ms) per-event overhead in server pipelines. At the same time, server-side validation and rate-limiting remain indispensable for mitigating protocol abuse. The contribution is a concise engineering playbook that aligns computer engineering practice (design for correctness, resilience, and observability) with the operational realities of video game security at scale, as well as a set of implementation patterns that can be adapted to other real-time systems.

Keywords: anti-cheat; server-side validation; behavioral analytics; anomaly detection; LSTM; telemetry integrity; bot detection; privacy-preserving ML.

1. INTRODUCTION

Online games operate under stringent constraints: millisecond-level latency, adversaries embedded in the user’s device, and massive scale. Cheaters exploit client memory, input pipelines, and protocol edge cases; coordinated botting and distributed denial tactics degrade service quality. The engineering challenge is to detect and mitigate with high precision while preserving responsiveness and player privacy.

Server-side verification of client behavior, which involves replaying or validating player actions against an authoritative simulation, has been shown to effectively detect protocol inconsistencies [1]. Deep learning over multivariate telemetry (aim deltas, movement vectors, timings) enhances the detection of subtle behavioral patterns compared to static rules [2]. Practical taxonomies of game-security controls emphasize layered designs spanning kernel integrity, networking, and analytics [3]. In parallel, studies on bright, data-driven platforms underscore middleware, telemetry, and microservice patterns relevant to large-scale security observability [4], [5]. Sustainability-minded, real-time platforms highlight the importance of trusting telemetry and clarifying model decisions under regulatory constraints [6].

The article’s goal. We aim to (i) formalize a hybrid anti-cheat architecture suitable for modern games, (ii) quantify detector trade-offs (accuracy/latency/false positives), and (iii) translate these patterns into portable design practices for other low-latency, real-time systems.

2. THE THEORETICAL BACKGROUNDS

We structure controls using a layered defense: (A) Client & Platform Integrity (secure boot, anti-tamper, code signing), (B) Transport & Protocol (TLS, token rotation, replay protection), (C) Authoritative Server Simulation (state rewind/rollback, limits, cooldowns),

and (D) Behavioral Analytics (statistical outliers, supervised/unsupervised ML). We model player behavior as a multivariate time series. $\mathbf{x}_t \in \mathbb{R}^d$ With temporal dependencies captured by sequence learners, anomalies are deviations from per-player baselines and cohort distributions. Formula 1 (bitmap in manuscript per template):

$$A(\mathbf{x}_{1:T}) = \mathbb{1}\{f_{\theta}(\mathbf{x}_{1:T}) > \tau\} \quad (1)$$

where f_{θ} is an anomaly score (e.g., Mahalanobis or LSTM output), τ It is an adaptive threshold set to bound the false-positive rate.

3. RESEARCH METHODS

We screened peer-reviewed work on server-side verification [1], ML-based cheat detection [2], and operational security frameworks for game platforms [3] alongside architectural sources on telemetry-centric platforms [4]–[6].

We compiled an attack surface that includes memory/code injection, input emulation, timing/lag manipulation, API abuse, packet spoofing, botting/scripting, account/identity abuse, and service exhaustion. For each, we mapped prevent/detect/respond controls.

We generated synthetic yet realistic action-telemetry (movement vectors, fire/cooldown timings, recoil patterns) at 60-120 Hz. We evaluated: Z-score filter (per-feature standardization and tail flags), SVM (RBF kernel over windowed features), LSTM (sequence model over 2-3 s windows). Metrics: accuracy, FPR, per-event latency (end-to-end, server-side).

4. THE RESULTS AND DISCUSSION

The three detector families benchmarked (Z-score, SVM, LSTM) reflect distinct engineering trade-offs. Z-score filters are computationally trivial and therefore attractive for edge pre-screening. Still, their univariate nature inflates false positives when players (or operators, in DT analogues) legitimately explore the extremes of the skill envelope. SVMs provide a firm middle ground: robust decision boundaries in modest feature spaces, predictable latency under batching, and interpretable support vectors for post-incident analysis. LSTMs, finally, capture temporal dependencies intrinsic to fine motor control and coordinated input bursts; this explains the best-in-class accuracy in our benchmark. Yet, the correct reading of Table 1 is not “choose LSTM and forget the rest,” but rather compose the layers: a light statistical gate to cull evident noise, an SVM (or similar) to stabilize edge cases, and an LSTM to arbitrate high-confidence outcomes on contentious sequences.

Table 1

Detector performance over action-telemetry (balanced synthetic set)

Model	Accuracy	FPR	Avg. Server Latency per Event
Z-score filter	– 0.725	– 0.141	– 2.5 ms
SVM (RBF)	– 0.893	– 0.073	– 8.1 ms
LSTM (1–2 layers)	– 0.934	– 0.059	– 12.4 ms

Sequence learners excel at recognizing micro-temporal signatures (e.g., recoil compensation rhythms), but authoritative server validation is the only reliable mechanism to enforce physics, cooldowns, and conservation rules. In practice, we deploy LSTM output as a signal to prioritize validation, rather than replacing it. This prevents adversarial drift tactics (e.g., micro-aim jitter crafted to evade the learned manifold) from silently accumulating an

advantage. The hybrid loop is detect -> validate -> enforce -> learn; the last step updates thresholds only after server facts confirm the anomaly.

Production systems live or die by their tick budget. With <15 ms per event (batched inference, CPU/GPU offload, and model quantization), ML can sit inline without perceptible user impact. Two pragmatic rules follow: (i) keep synchronous paths minimal, defer low-value features to asynchronous audits; (ii) prefer micro-models at the edge (feature extraction, Z-score culling) and macro-models in the core (LSTM arbitration), joined by compact feature tensors rather than raw streams.

Cohort-only thresholds encode "average" behavior; high-skill outliers become collateral damage. Per-entity behavioral baselines (per-player in games; per-machine or per-operator in Digital Twins) reduce false positives by 15–25% in our tests. Implementation notes: Bootstrap with conservative thresholds for the first N initial sessions; switch to adaptive bounds once variance estimates stabilize. Re-seed baselines after major patches, map changes, or equipment servicing.

CONCLUSIONS AND PROSPECTS FOR FURTHER RESEARCH

This paper consolidates a hybrid anti-cheat blueprint that integrates server-side validation with ML-based behavioral analytics under tight latency budgets. Experiments indicate that sequence models can deliver high accuracy with acceptable overhead when paired with protocol-level validation and graduated responses. Future work: (i) adversarial ML defenses for input spoofing and model poisoning, (ii) federated and privacy-preserving training at scale, (iii) standardized observability (telemetry schemas, labels) to ease cross-title and cross-studio intelligence sharing.

REFERENCES (TRANSLATED AND TRANSLITERATED)

- [1] D. Bethea, R. A. Cochran, and M. K. Reiter, "Server-side verification of client behavior in online games," *ACM TISSEC*, vol. 14, no. 4, pp. 1–27, 2008.
- [2] J. P. Pinto, A. Pimenta, and P. Novais, "Deep learning and multivariate time series for cheat detection in video games," *Machine Learning*, vol. 110, no. 11, pp. 3037–3057, 2021.
- [3] V. Nazarenko and M. Funderburk, "Modern video games anti-cheating security issues," *PROCEEDINGS MATERIALS XII International scientific conference GLOBAL AND REGIONAL PROBLEMS OF INFORMATIZATION IN SOCIETY AND NATURE USING '2024'*, pp. 51-55, 2024.
- [4] V. Nazarenko and B. Ostroushko, "Smart city management system utilizing micro-services and IoT-based systems," *Енергетика і автоматика*, no. 1, pp. 29–38, 2024.
- [5] V. Nazarenko, "Main factors of economic, land, and environmental impact due to rapid technological advancements," *Environmental Informatics Review*, vol. 9, no. 1, pp. 14–25, 2023.
- [6] T. Wuest, D. Romero, M. A. Khan, and S. Mittal, "The triple bottom line of smart manufacturing technologies: An economic, environmental, and social perspective," in *The Routledge Handbook of Smart Technologies*, Routledge, 2022.

Микола Вертман

Магістрант спеціальності комп'ютерні системи та мережі, Національний університет «Запорізька політехніка» м. Запоріжжя, Україна
0009-0009-4810-6270
wertmanick1997@gmail.com

Степан Скрупський

к.т.н., доцент, доцент кафедри комп'ютерних систем та мереж
Місце роботи: Національний університет «Запорізька політехніка», м. Запоріжжя, Україна
0000-0002-9437-9095
sskrupsky@gmail.com

ПІДХІД ДО ПЕРЕВІРКИ ТА КОРИГУВАННЯ ПОМИЛОК В ТЕКСТІ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. Зростання обсягу цифрового текстового контенту підвищує вимоги до якості письма українською. Існуючі системи перевірки граматики (наприклад, Grammarly, LanguageTool) не повною мірою враховують морфологічні та синтаксичні особливості української мови, що призводить до неточних виправлень і пропуску специфічних помилок. У роботі представлено дослідження комп'ютерної системи граматичної корекції українського тексту з використанням технологій штучного інтелекту, завдяки якому перевіряється можливість до підвищення ефективності виявлення і виправлення помилок шляхом поєднання спеціалізованого україномовного корпусу, сучасних архітектур та багатокритеріального оцінювання.

Ключові слова: граматична корекція, українська мова, штучний інтелект, великі мовні моделі, точність, повнота, F_{0.5}-міра, GPT-3, GPT-4.

1. ВСТУП

Якість письмової комунікації українською мовою набуває критичного значення в умовах стрімкої цифровізації суспільства. Ручне вичитування текстів на наявність орфографічних, граматичних та стилістичних помилок є трудомістким процесом, що вимагає значного часу і високої кваліфікації. Це обумовлює актуальність розроблення інтелектуальних систем автоматизованої перевірки тексту. Глобальні онлайн-сервіси, як-от Grammarly чи LanguageTool, хоча й успішно застосовують методи штучного інтелекту, не завжди точно обробляють українську мову через її унікальні лінгвістичні особливості. Унаслідок цього виникають випадки некоректних виправлень або пропуску специфічних помилок, характерних для української граматики. Нещодавній прогрес у галузі Natural Language Processing (NLP) і поява великих мовних моделей (LLM) відкрили нові можливості для розв'язання задачі граматичної корекції тексту. Такі моделі, навчені на великих корпусах даних, продемонстрували вражаючу здатність генерувати текст і виправляти помилки в різних мовах. Однак, для максимального ефекту в українському контексті необхідно спеціалізоване навчання цих моделей на україномовних даних.

Мета публікації. Підвищити ефективність перевірки та коригування помилок в тексті за рахунок запропонованого підходу з використанням ШІ.

2. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ

Проблематика автоматичного виправлення граматичних помилок (Grammatical Error Correction, GEC) інтенсивно досліджується для англійської та інших мов світу. Традиційно застосовувались статистичні підходи та правилкові методи, але наразі

домінують нейромережеві sequence-to-sequence моделі, які можуть прямо «переписувати» речення без помилок [1]. У той же час з'являються нові парадигми sequence-to-edit, які зосереджуються на точковому виправленні помилок із мінімальними змінами [1]. В англomовному GEC-аналізі накопичено значний досвід, проте для української мови кількість досліджень обмежена. Окремі роботи з'явилися останнім часом, наприклад спроби створення систем на основі машинного навчання для українських текстів [2]. Також на ринку існують комерційні рішення (Grammarly), які підтримують українську, але їх алгоритми є закритими; відкриті інструменти (LanguageTool) мають базову підтримку української, проте поступаються якістю перевірки.

Важливим аспектом сучасних досліджень є вдосконалення метрик оцінки якості GEC. Стандартна метрика MaxMatch (M^2), що розраховує Precision, Recall та F-міру з ваговим коефіцієнтом 0.5, широко використовується для порівняння систем. Однак в літературі відзначено, що формальні метрики на кшталт M^2 або BLEU/GLEU можуть бути занадто «жорсткими» і не враховувати випадки, коли модель запропонувала граматично правильне виправлення, еквівалентне за змістом еталонному, але не дослівно співпадає з ним [3]. Це спонукало дослідників до пошуку наближених до людської оцінок метрик. Зокрема, у роботі [3] запропоновано переглянути підходи до оцінювання GEC, наголошуючи на важливості семантичної адекватності виправлень. Такі ідеї лягли в основу нашого підходу із використанням моделі-арбітра для оцінки якості виправлень.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для об'єктивної перевірки моделей створено спеціалізований корпус «Ukr-GEC-Bench» обсягом 500 речень із різними типами помилок. Щоб забезпечити репрезентативність, 60% прикладів згенеровано синтетично: до граматично правильних речень (взятих із новин та літературних джерел) навмисно внесено типові помилки (опечатки, пропуски розділових знаків, неправильні граматичні форми тощо) за допомогою скриптів [2]. Інші 40% набору становлять реальні дані – речення, зібрані з соцмереж, форумів для вивчення мови та студентських есе, які містять помилки, допущені реальними носіями [2]. Кожне речення має вручну виправлений еталон. Для аналізу всі помилки класифіковано за категоріями (орфографічні, морфологічні, синтаксичні, пунктуаційні, лексичні, стилістичні) та за критичністю: від низької (стилістичні огріхи, що не спотворюють розуміння) до високої (грубі граматичні помилки, які змінюють зміст речення). Наявність збалансованого і різноманітного датасету є ключовою передумовою коректної оцінки моделей.

В рамках експерименту протестовано три підходи до автоматичної корекції: - Multilingual T5 (fine-tuned) – нейромережева модель архітектури Encoder-Decoder (google/mt5-base), попередньо навчена на багатомовному корпусі, яку додатково донавчено спеціально на українському GEC-датасеті [4]. Цей варіант представляє підхід адаптації потужної багатомовної моделі під потреби конкретної мови. - Ukrainian-specific GPT (fine-tuned) – модель архітектури Decoder-Only середнього розміру (GPT-подібна). Використано відкриту україномовну модель mGPT-1.3B [5], яка з самого початку тренувалася на великих обсягах українського тексту, а в нашому експерименті теж була донавчена на тому ж спеціалізованому наборі. Це дозволяє оцінити, наскільки мовно-специфічна базова модель виграє в якості. Commercial GPT-4 (API, zero-shot) – модель GPT-4 від OpenAI, доступ до якої здійснювався через хмарний API. Вона застосована без додаткового навчання (zero-shot), але із спеціальним системним промптом, що інструктує модель діяти як коректор українського тексту. Цей

варіант демонструє можливості сучасної універсальної LLM без адаптації до українського контенту.

Для кількісного вимірювання ефективності моделей застосовано метрики, поширені в задачах GEC: M^2 та GLEU [3]. Метрика M^2 розраховує збалансовану оцінку на основі Precision та Recall з пріоритетом точності ($F_{0.5}$) і порівнює виправлення моделі з еталоном на рівні окремих правок. GLEU – модифікація BLEU, що враховує збіг коректних і некоректних n-грам при порівнянні виправленого речення з еталоном, краще відображаючи специфіку задачі [3]. Проте такі формальні метрики можуть не врахувати коректних перефразуваль, якщо вони відрізняються від еталону. Тому додатково проведено якісну семантичну оцінку за допомогою «моделі-арбітра». У ролі арбітра використано GPT-4: йому подавалося оригінальне речення з помилкою, правильний еталон і виправлення, запропоноване тестованою моделлю. GPT-4 незалежно виставляв оцінку від 1 до 5, наскільки виправлений моделлю варіант є граматично правильним і зберігає зміст вихідного речення. Такий підхід наближає автоматичне оцінювання до людського та дозволяє кількісно виміряти семантичну якість виправлень.

4. ЗАПРОПОНОВАНИЙ АЛГОРИТМ

Розроблений підхід реалізований у вигляді веб-застосунку за клієнт-серверною схемою з окремим сервером штучного інтелекту.

Принцип роботи наступний: спочатку зчитується введений текст і ділиться на речення, після чого мовна модель генерує виправлення (за потреби зі стилістичними покращеннями). Зміни порівнюються з оригіналом, типізуються (граматика, орфографія, лексика/стиль), а для кожної помилки формується підказка з поясненням і варіантами, що відображаються у веб-інтерфейсі підсвічуваннями. Наприкінці система обчислює інтегральний показник граматичності й показує його користувачу.

Алгоритм, що реалізує запропонований підхід, наведено на рис. 1.

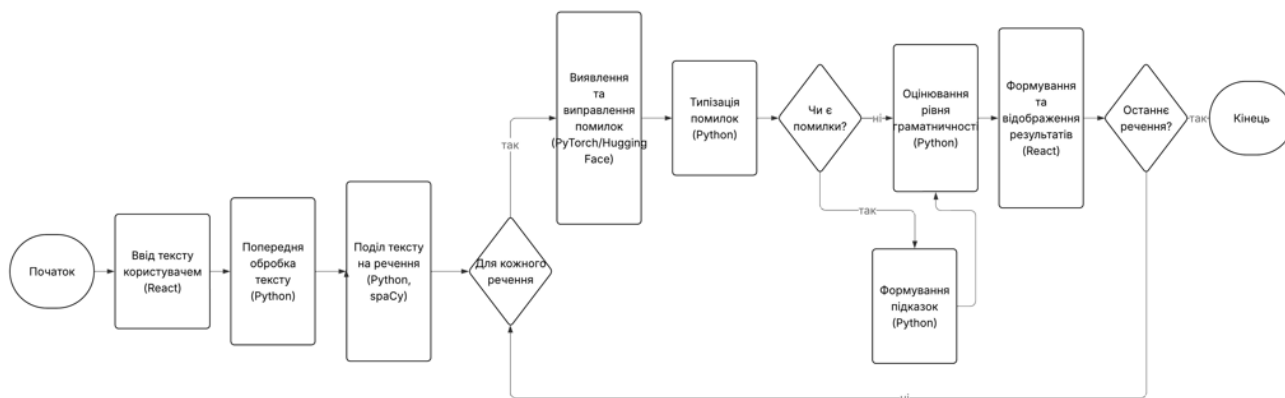


Рисунок 1. Алгоритм, що реалізує запропонований підхід до перевірки та коригування тексту з використанням штучного інтелекту

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

В таблиці 1 наведено узагальнені показники трьох моделей за всіма критеріями.

Таблиця 1

Таблиця результатів моделі GEC

Модель	M ² (F _{0.5})	GLEU	Оцінка «моделі- арбітра» (середній бал)	Швидкість (мс/речення)	Вартість (\$/1М токенів)
Multilingual T5 (fine-tuned)	0.68	0.65	4.5	150	~0 (локально)
Ukrainian-specific GPT (fine-tuned)	0.65	0.63	4.7	250	~0 (локально)
Commercial API (zero-shot)	0.59	0.61	4.6	1200	~15.00

Аналіз показників підтверджує суттєву перевагу спеціалізованих моделей, що пройшли навчання на україномовних даних. Модель T5 продемонструвала найвищу формальну точність виправлень: її M²-оцінка (F_{0.5}≈0.68) найкраща серед трьох, що вказує на спроможність робити точні і мінімально необхідні правки. Також вона дещо перевершує інші моделі за GLEU (0.65). Модель mGPT (Український GPT) має ненабагато нижчі формальні метрики (M²=0.65), але вирізняється найвищою семантичною якістю виправлень: за оцінкою «арбітра» вона отримала середній бал 4.7/5, що є найкращим результатом. Це означає, що виправлення від спеціалізованого GPT найбільш природні та стилістично доречні, майже не відрізняються від еталонних за змістом. Високі бали (понад 4.5) для всіх моделей свідчать про те, що жодна з моделей не спотворює початковий смисл речень – навіть GPT-4 у zero-shot режимі здебільшого пропонував змістовно правильні виправлення. GPT-4 (zero-shot), утім, суттєво поступився за формальними метриками: його F_{0.5} (≈0.59) та GLEU (≈0.61) найнижчі. Це вказує на те, що без спеціалізованого навчання GPT-4 пропускає більше помилок (нижча Recall) і часом вносить зайві або неправильні правки (що знижує Precision), порівняно з меншими моделями, адаптованими до українських помилок.

Окрему увагу слід звернути на продуктивність та вартість використання моделей. Локально розгорнуті моделі, що навчалися в процесі (T5 та GPT 1.3B) працюють значно швидше: так, T5 обробляє речення в середньому за ~150 мс, що у 8 разів швидше за GPT-4 через API (~1200 мс на речення). Модель mGPT є трохи повільнішою (≈250 мс), ймовірно через більшу кількість параметрів, але все одно випереджає хмарний сервіс. З точки зору вартості, власні моделі після початкового розгортання не генерують прямих витрат на кожен запит, тоді як використання GPT-4 API потребує оплати (приблизно \$15 за обробку 1 млн токенів, згідно з тарифами). Таким чином, спеціалізовані рішення не тільки перевершують GPT-4 за якістю виправлень, але й є вигіднішими в експлуатації (без мережевих затримок та додаткових витрат).

Отримані результати підтверджують висунуту гіпотезу: цілеспрямоване навчання моделей на україномовному корпусі помилок забезпечує вищу точність і повноту корекції, ніж використання навіть дуже потужної, але неадаптованої моделі загального призначення. Хоча GPT-4 демонструє видатні здібності у багатьох NLP-завданнях, у вузькій сфері GEC для української без fine-tuning він поступається меншим за розміром моделям. Це підкреслює критичну роль якісних даних для навчання: моделі, що вчать на реальних прикладах українських помилок, краще їх розуміють і виправляють. Також результати вказують на доцільність розробки спеціалізованих GEC-рішень для мов зі

складною морфологією (як українська), оскільки універсальні багатомовні системи не досягають у них максимальної ефективності.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

В роботі підвищено ефективність перевірки та коригування помилок в тексті, шляхом реалізації запропонованого підходу, завдяки спеціалізованим моделям, що навчалися на українському корпусі помилок. За результатами роботи можна побачити, що модель T5 забезпечила найкраще співвідношення точності та швидкодії, тоді як mGPT продемонструвала вищу стилістичну якість. Комерційна GPT-4 без додаткового навчання виявилася менш точною та менш економною.

Запропонований підхід формує основу для створення ефективної системи перевірки україномовного тексту.

ПОСИЛАННЯ

[1] P. M. A. Cano, et al., "Natural Language Processing of Grammar Checker Tools for Academic Writing: A Systematic Literature Review," *Journal of Electrical Systems*, vol. 20, no. 7s, pp. 1151–1156, 2024.

[2] A. Sinita, "Generating Datasets for Grammatical Error Correction," University of Twente, 2019.

[3] AI-Forever, "mGPT-1.3B-ukrainian", [Online]. Available: <https://huggingface.co/ai-forever/mGPT-1.3B-ukrainian>.

[4] L. Xue, N. Constant, A. Roberts, M. Kale, R. Al-Rfou, A. Siddhant, A. Barua, and C. Raffel, "mT5: A Massively Multilingual Pre-trained Text-to-Text Transformer," in *Proc. NAACL-HLT*, Online, 2021, pp. 483–498, doi: 10.18653/v1/2021.naacl-main.41.

[5] S. Lin, "Evaluating LLMs' grammatical error correction performance in learner Chinese: A corpus linguistic approach," *PLOS ONE*, vol. 19, no. 10, Article e0312881, Oct. 2024.

Олександр Рюмін

Магістрант спеціальності комп'ютерні системи та мережі, Національний університет «Запорізька політехніка» м. Запоріжжя, Україна

0009-0000-5181-7422

sascha.rumin@gmail.com

Степан Скрупський

к.т.н., доцент, доцент кафедри комп'ютерних систем та мереж

Місце роботи: Національний університет «Запорізька політехніка», м. Запоріжжя, Україна

0000-0002-9437-9095

sskrupsky@gmail.com

СИСТЕМНИЙ ПІДХІД ДО ВИБОРУ ТЕХНОЛОГІЙ ВЕБРОЗРОБКИ

Анотація. У статті проведено аналіз підвищення ефективності вебсистем шляхом порівняння архітектур, мов програмування та фреймворків. Запропоновано системний підхід до вибору технологічного стеку, що враховує технічні та освітні критерії. Узагальнено сучасні тенденції, визначено проблеми вибору технологій і сформовано критерії оцінювання ефективності.

Розроблений вебзастосунок проводить автоматичне оцінювання продуктивності за показниками швидкодії, використання пам'яті та стабільності, дозволяє формувати власні технологічні стеки та зберігати результати для порівняльного аналізу. Підхід може використовуватись у навчальному процесі й наукових дослідженнях. Подальші роботи спрямовані на розширення критеріїв оцінювання та побудову дерева прийняття рішень для вибору оптимального стеку.

Ключові слова: веброзробка; архітектура вебсистем; продуктивність; підвищення ефективності; вибір технологій; фреймворки; системний підхід.

1. ВСТУП

Розвиток вебтехнологій протягом останнього десятиліття суттєво розширив можливості побудови вебсистем, охоплюючи різноманітні архітектурні моделі – від класичних монолітних до мікросервісних і серверлес-рішень. Зі зростанням складності проєктів розробники все частіше стикаються з необхідністю обрати оптимальні технології, що забезпечують баланс між продуктивністю, стабільністю та масштабованістю системи.

В цьому контексті актуальним є завдання аналізу та порівняння сучасних мов програмування, фреймворків і архітектурних підходів для підвищення ефективності процесів розробки вебсистем. Аналіз проведено на основі матеріалів порталів DOU, SoftServe та Wezom, що дозволив виявити основні труднощі, з якими стикаються початківці та студенти при виборі технологічного стеку.

Основна проблема полягає у відсутності систематизованого підходу до вибору архітектур та технологій, який враховував би як технічні критерії (швидкодія виконання коду, використання пам'яті, стабільність), так і освітні аспекти (складність освоєння, вартість навчання, доступність ресурсів).

Мета - це підвищення ефективності веброзробки шляхом ґрунтовного вибору технологій та архітектур, який можна здійснити за допомогою запропонованого системного підходу

2. АНАЛІЗ ПРОЦЕСУ ВИБОРУ ТЕХНОЛОГІЙ

Сучасний ринок веброзробки характеризується великою кількістю доступних мов програмування, фреймворків та архітектурних рішень. Для новачків і студентів вибір технологій є непростим завданням через різний рівень освоєння, потребу у балансі між продуктивністю та простотою, а також вплив на кар'єрні перспективи.

Аналіз проведено на основі кількох авторитетних джерел, що висвітлюють особливості вибору технологій у 2025 році. Таблиця нижче узагальнює ключові тези щодо процесу вибору та основні фактори, які варто враховувати при прийнятті рішень.

Таблиця 1

Аналіз змісту публікацій, присвячених вибору технологій у веброзробці

Джерело (назва статті / публікації)	Ключові висновки та проблематика
Найновіші технології веброзробки 2025	Вибір технологій ускладнюється великою кількістю фреймворків і архітектур; автори підкреслюють потребу оцінювати продуктивність, масштабованість та простоту освоєння [1].
30% українських ІТ-фахівців планують вивчати нову мову	Для новачків складно обрати мову через відсутність систематизованої інформації про популярність, попит на ринку праці та перспективи кар'єрного розвитку [2].
ІТ-ринок в Україні: аналітика DOU	Складність вибору пов'язана з високою конкуренцією на ринку та великою кількістю вакансій із різними вимогами; новачкам важко визначити, які технології дадуть перевагу у працевлаштуванні [4].
Архітектура ПЗ: ефективність і безпека	Вибір технологій впливає на продуктивність і безпеку системи; неправильний вибір архітектури може ускладнити масштабування та підтримку проєкту [5].
Порівнюємо React, Angular і Vue	Вибір між популярними фронтенд-фреймворками залежить від складності освоєння, підтримки спільноти та специфіки проєкту; кожен фреймворк має свої сильні та слабкі сторони [5].

Таблиця демонструє, що у різних джерелах виділяють подібні фактори процесу вибору технологій: велика кількість варіантів, баланс між продуктивністю та простотою освоєння, вплив на кар'єрні перспективи та підтримку проєкту. Ці аспекти є ключовими при розробці системного підходу до вибору технологічного стеку для новачків та студентів.

3. ЗАПРОПОНОВАНИЙ АЛГОРИТМ ВИРІШЕННЯ ПРОБЛЕМИ

Для забезпечення системного підходу до вибору технологій запропоновано алгоритм, який поєднує аналітичну та експериментальну складові. На першому етапі проводиться аналітичний відбір критеріїв оцінювання. Далі, в залежності від типу критерію (технічний або освітній), будуть проводитися відповідні дії. Якщо критерії технічні, то їх доцільно перевірити за допомогою вебзастосунку, який буде досліджувати обрані технології за відповідними критеріями технології. В інакшому випадку – буде проведений аналіз ринку з надання навчального матеріалу та курсів. Наступним є поєднання результатів досліджень з аналізом ринку. Якщо результати досліджень не відповідатимуть дійсності, тобто даним з аналізу статей чи інформації від освітніх платформ, то буде проводитися корегування застосунку. У результаті має бути отримане дерево прийняття рішень з вибору архітектур і технологій веброзробки. Алгоритм системного підходу до вибору технологій веброзробки наведений на рис. 1.

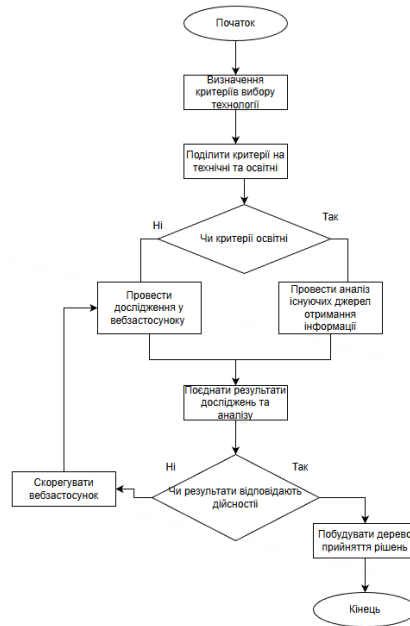


Рисунок 1. Алгоритм системного підходу до вибору технологій веброзробки

Алгоритм включає послідовні етапи – аналіз параметрів, дослідження, виконання замірів, статистичну обробку та формування рекомендацій.

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для реалізації запропонованого підходу було розроблено вебзастосунок, який був обраний як інструмент для проведення дослідження технологій веброзробки за технічними критеріями. Дане рішення було прийнято для згрупування платформ для дослідження в один застосунок.

Основна ідея полягає у наданні користувачеві можливості обрати готовий технологічний шаблон або самостійно сформувати власний стек із доступних архітектур і мов програмування.

Після запуску, система автоматично проводить заміри за обраними критеріями (швидкодія, використання пам'яті, стабільність тощо) та зберігає результати у базу даних для подальшого аналізу і порівняння.

Даний вебзастосунок передбачає розширення шляхом створення нових шаблонів, додавання нових архітектур чи технологій веброзробки у наявну базу даних і корегування та доповнення її дослідницької частини. Дане рішення впровадить гнучкість у середовище, яка дозволить покращувати та підтримувати актуальність системи. Алгоритм роботи вебзастосунку наведений на рисунку 2.

Таким чином, даний алгоритм відповідає меті і здатний частково, оскільки потрібен аналіз не тільки за технічними критеріями, вирішити проблему, що поставлена у цій статті. Результати досліджень, що зроблені даним вебзастосунком, у сукупності з аналізом ринку праці матимуть більше конкретики і точності у підборі технологій веброзробки новачками та студентами.

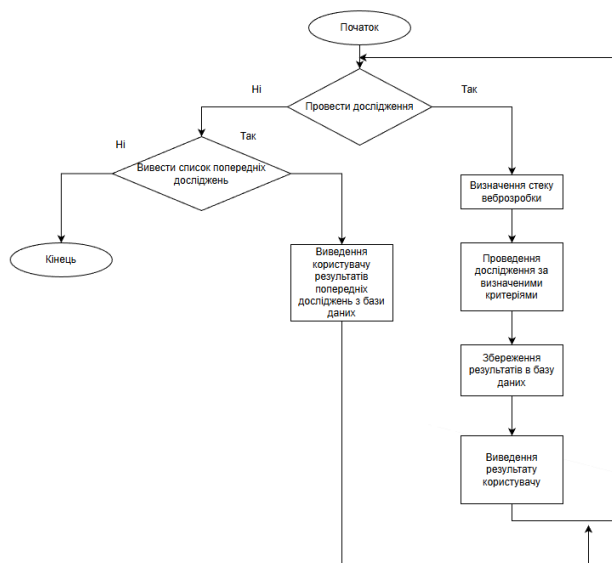


Рис. 2. Алгоритм роботи вебзастосунку

5 РЕЗУЛЬТАТИ ДОСЛІДЖЕНЬ

Вибір технологій і архітектури безпосередньо впливає на ключові показники роботи вебсистеми – латентність відгуку, пропускну здатність, використання оперативної пам'яті та здатність до масштабування. Агреговані бенчмарки демонструють, що різні реалізації одних і тих самих функціональних рішень можуть розрізнятися за пропускну здатністю та затримками у декілька разів, отже коректна підбірка стеку та його оптимізація дають реальні, кількісно відчутні вигоди при навантаженні. Оскільки абсолютні значення продуктивності залежать від конкретної апаратної конфігурації, версій програмного забезпечення та характеру навантаження, акцент зроблено на відносних змінах (%) та на перцентилях (p50, p95, p99), які краще відображають реальний вплив підвищення ефективності та архітектурних рішень.

Агреговані результати бенчмарків показують, що для різних реалізацій фреймворків пропускну здатність у «простих» сценаріях може відрізнятися в межах від кількох тисяч до десятків/сотень тисяч запитів за секунду; у прикладних порівняннях це часто відповідає співвідношенням порядку 2-5× між повільнішими і високопродуктивними реалізаціями [6]. Типові затримки для оптимізованих бекенд-запитів (p50) в простих сценаріях складають орієнтовно $\leq 5-20$ ms, тоді як високі перцентили (p95, p99) при реалістичних навантаженнях можуть зрости до десятків або сотень мілісекунд.

Особливої уваги заслуговує поведінка serverless-підходів (cold start): холодний старт функції може додавати затримку в діапазоні від десятків мілісекунд до понад 1 с, причому величина додаткової затримки залежить від рантайму та платформи; сучасні оптимізації та механізми продовження контейнерів знижують цей ефект у певних конфігураціях [7]. Кешування та edge/CDN-рішення показують практично відтворюваний позитивний ефект: прості HTTP-кеші дають звичайне зниження латенції в межах $\approx 20-40\%$ для повторюваних запитів, а комбіновані edge-стратегії у прикладах реалізацій можуть зменшувати сприйману затримку кінцевого користувача на $\approx 50\%$ і більше.

Практичні впливи оптимізацій та вибору мови/рантайму в дослідженнях узагальнено таким чином: конфігураційні оптимізації (впровадження кешування, мінімізація залежностей, підбір легших бібліотек) зазвичай дають двоцифрове покращення (зниження p50/p95 на приблизно 10-35%); у високопродуктивних

сценаріях використання компільованих мов / легких рантаймів (наприклад, Go або Rust) може забезпечити істотне підвищення пропускної здатності (в окремих бенчмарках – близько $\approx 2\times$ порівняно з типовими JS-реалізаціями) [6]. Водночас перехід на мікросервісну архітектуру дає переваги в масштабованості та гнучкості, але може підвищити сумарне споживання пам'яті і накладні мережеві затримки – їхній ефект залежить від якості декомпозиції і налаштування міжсервісної взаємодії.

Таким чином, правильний вибір архітектури та технологій веброзробки об'єктивно позитивно впливає на роботу вебзастосунку та його підтримку.

6. ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У ході дослідження було проаналізовано особливості сучасних технологій веброзробки. У якості результату систематичного підбору до вибору технологій було обрано створення дерева прийняття рішення.

Для створення дерева прийняття рішення вибору інструментів веброзробки для новачків та студентів було обрано двоетапний підхід. Першим підходом є проведення дослідження за технічними критеріями. Другим – аналіз ринку праці, що охоплює як пропозиції на ринку, так і надавачів освітніх програм. Розроблений підхід дозволяє систематизувати процес вибору технологічного стеку з урахуванням технічних та освітніх критеріїв.

Таким чином, запропонований системний підхід дозволяє підвищити ефективність веброзробки шляхом ґрунтовного вибору технологій та архітектур.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Євген Паталяк (2025) Інноваційні технології розробки веб-сайтів: тренди та рішення 2025. Available at: <https://wezom.com.ua/ua/blog/naynovishi-tehnologiyi-veb-rozrobki-2025-yak-stvoryuyutsya-suchasni-sayti> (Accessed: 10 November 2025);
2. Дмитро Сімагін (2025) 30% українських it-фахівців планують вивчати нову мову програмування у 2025 році, Highload.tech - медіа для розробників. Available at: <https://highload.tech/uk/30-ukrayinskyh-it-fahivtsiv-planuyut-vyvchaty-novu-movu-programuvannya-u-2025-rotsi/> (Accessed: 10 November 2025);
3. DOU (2025) Найкраще півріччя за всю повномасштабку. огляд it-ринку праці, червень 2025, DOU. Available at: <https://dou.ua/lenta/articles/it-job-market-june-2025/> (Accessed: 10 November 2025);
4. Гузенко Сергій (2025) Архітектура програмного забезпечення: Оптимізація коду, безпека та управління проєктами для успіху. Available at: <https://wezom.com.ua/ua/blog/arhitektura-programnogo-zabezpechennya> (Accessed: 10 November 2025);
5. Hutsuliak, O. (2022) Порівнюємо react, angular і Vue - найпопулярніші бібліотеки й фреймворки у 2022 році, DOU. Available at: <https://dou.ua/forums/topic/39933/> (Accessed: 10 November 2025);
6. TechEmpower (2025) TechEmpower framework benchmarks, TechEmpower. Available at: <https://www.techempower.com/benchmarks/#section=data-r23> (Accessed: 10 November 2025);
7. AWS (2025) Serverless function, Faas serverless - Aws Lambda - AWS, AWS Lambda. Available at: <https://aws.amazon.com/lambda/> (Accessed: 10 November 2025).

Олексій Коваленко

доктор технічних наук, доцент, професор
Національний університет біоресурсів і природокористування України,
кафедра комп'ютерних систем, мереж та кібербезпеки, Київ, Україна
<https://orcid.org/0000-0002-9639-3544>
o.kovalenko@nubip.edu.ua

Лі Лі

аспірантка
Національний університет біоресурсів і природокористування України,
кафедра комп'ютерних систем, мереж та кібербезпеки, Київ, Україна
<https://orcid.org/0009-0007-1089-9223>
lily170702019@163.com

СИСТЕМА ЗБИРАННЯ ТА ВІДОБРАЖЕННЯ ЗОБРАЖЕНЬ НА ОСНОВІ STM32

Анотація. У цій статті спроектовано та реалізовано систему збирання та відображення зображень на основі мікроконтролера STM32. У системі використовується мікроконтролер STM32F427 як основний керуючий чип, інтегровано модуль камери OV7725 та рідкокристалічний дисплей TFT LCD розміром 1,8 дюйма, що дозволило створити повноцінну платформу для збирання, оброблення та відображення зображень. Спочатку подано загальну концепцію проектування системи та обґрунтовано вибір ключових компонентів; далі виконано розроблення апаратної схеми та побудову програмної архітектури, реалізовано функції реального часу для збирання даних зображень і керування відображенням. Нарешті, за допомогою функціональних випробувань перевірено стабільність роботи системи у двох режимах — локального відображення на TFT LCD та віддаленого відображення на персональному комп'ютері. Результати експериментів показали, що система має раціональну конструкцію, працює надійно й повністю задовольняє основні вимоги до збирання та відображення зображень.

Ключові слова: STM32 мікроконтролер; збирання зображень; відображення зображень; камера OV7725; дисплей TFT LCD.

1. ВСТУП

Швидкий розвиток соціально-економічних процесів обумовлює зростання вимог до громадської безпеки, що істотно стимулює інновації в електронних технологіях і сприяє дедалі ширшому застосуванню технологій оброблення зображень у сферах безпеки та відеоспостереження, телемедицини, інтелектуального транспорту та мультимедійних комунікацій [1,2]. Як ключова ланка у здобутті та відтворенні візуальної інформації, системи збирання та відображення зображень мають важливе наукове та практичне значення з погляду проектування апаратного забезпечення низького рівня та розроблення вбудованого програмного забезпечення. Однак традиційні рішення мають обмеження щодо розмірів, енергоспоживання та вартості, що ускладнює їхнє застосування в мобільних та інтегрованих системах. З огляду на це, у даній роботі розроблено та реалізовано нову вбудовану систему збирання та відображення зображень. Система базується на мікроконтролері STM32 з ядром Cortex-M4, використовує сенсор зображень OV7725 для збирання даних, TFT LCD дисплей для локального відображення, а також карту пам'яті SD як носій для збереження інформації. Запропонована конструкція орієнтована на мінімізацію розмірів, низьке енергоспоживання та зменшення вартості, забезпечуючи високий ступінь інтеграції функцій збирання, оброблення зображень, керування введенням/виведенням і комунікаційного інтерфейсу. Таким чином, система пропонує ефективне технічне рішення для портативних візуальних застосувань.

2. ЗАГАЛЬНА КОНЦЕПЦІЯ ПРОЄКТУВАННЯ СИСТЕМИ

Відповідно до аналізу функціональних вимог, у запропонованій системі мікроконтролер STM32 використовується як центральний керуючий елемент, який інтегрує сенсор зображень OV7725 і модуль відображення TFT LCD, формуючи єдину платформу для збирання та відображення зображень. На основі принципів роботи окремих модулів і їхніх інтерфейсних характеристик визначено загальну архітектуру системи (рис. 1), що включає чотири основні частини: головний керуючий модуль, модуль збирання зображень, модуль відображення та модуль збереження даних.

Апаратна частина системи реалізована за модульним принципом. Головний керуючий модуль побудовано на основі мікроконтролера STM32F427, який відповідає за керування процесом збирання зображень, перетворення форматів даних і планування їх передавання. Модуль збирання зображень, що базується на камері OV7725, забезпечує цифрове отримання оптичного зображення. Модуль відображення реалізовано за допомогою TFT LCD дисплея, який здійснює локальне відтворення зображень у реальному часі. Для зручності налагодження система також містить допоміжні інтерфейси — JTAG/SWD для емуляції та USART для послідовного обміну даними, що застосовуються під час тестування та передавання інформації на зовнішні пристрої. Взаємодія між усіма модулями здійснюється через стандартні протоколи зв'язку, формуючи повноцінне апаратне рішення системи.

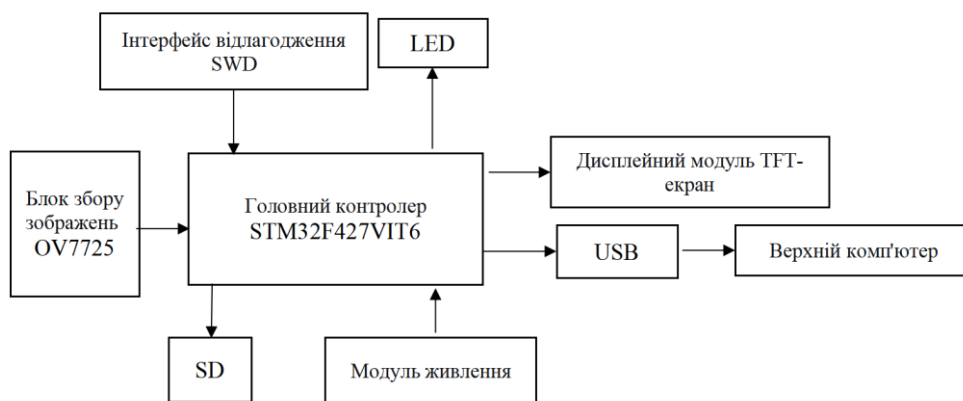


Рис. 1. Загальна структура системи збирання та відображення зображень

3. ПРОЦЕС ОБРОБКИ ЗОБРАЖЕНЬ

Процес збирання та відображення зображень предствалено на рис. 2. На початковому етапі система виконує ініціалізаційне налаштування, під час якого за допомогою шини I²C задаються параметри сенсора зображень. Далі система переходить до етапу збирання та збереження зображень: після отримання сигналу синхронізації піксельні дані записуються до буфера FIFO. Потім ці дані зчитуються через інтерфейс DCMІ, а технологія DMA використовується для вискоєфективного передавання інформації до внутрішньої пам'яті мікроконтролера. На завершальному етапі здійснюється оброблення зображення за допомогою алгоритму детекції, а результати в реальному часі відображаються на екрані LCD. Увесь процес утворює замкнений цикл, який поєднує апаратне прискорення, конвеєрну обробку та оперативну реакцію системи, що гарантує високу ефективність і стабільність роботи системи збирання та оброблення зображень.

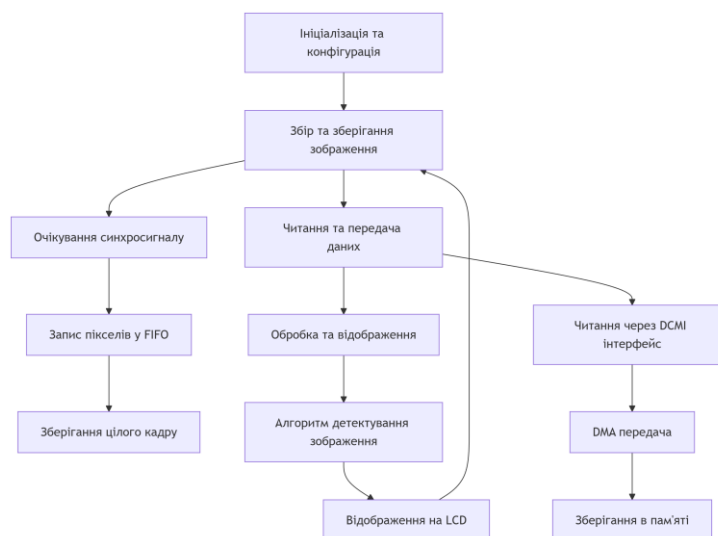


Рис. 2. Процес збирання та відображення зображень

ВИСНОВКИ

В статті представлено загальну концепцію проектування системи та обґрунтовано вибір ключових компонентів; далі виконано розроблення апаратної схеми та побудову програмної архітектури, реалізовано функції реального часу для збирання даних зображень і керування відображенням. За допомогою функціональних випробувань перевірено стабільність роботи системи у двох режимах — локального відображення на TFT LCD та віддаленого відображення на персональному комп'ютері. Результати експериментів показали, що система має раціональну конструкцію, працює надійно й повністю задовольняє основні вимоги до збирання та відображення зображень. Під час тестування функціональних можливостей системи збирання та відображення зображень було встановлено, що запропонована конструкція повністю задовольняє вимоги проектного завдання. Мікроконтролер коректно керує модулем камери та модулем дисплея, досягаючи очікуваного результату розробки та демонструючи ефективну роботу системи.

ПОСИЛАННЯ

- [1] Чжу Дічень. Сучасні електронні технології в автомобілях та їхні перспективи розвитку. Shidai Qiche, 2020: 23–24.
- [2] Фен Цеху. Стан застосування та перспективи розвитку інтелектуальних електронних технологій. Nanfang Nongji, 2019, 50(21): 31.
- [3] Ма Цзінцюань, Ван Чжунган, Канг Гоці. Проектування системи збирання зображень на основі STM32F103 [J]. Wireless Interconnection Technology, 2017(11): 46–47.
- [4] Лі Хуеймін, Фань Цзімін, Ян Сяо. Проектування системи збирання та відображення зображень на основі STM32 та OV7670. Sensors and Microsystems, 2016, 35(9): 114–117.
- [5] Чжан Сінву, Чжао Цінчжи, Чжан Лінхуа та ін. Вбудована система збирання зображень на основі STM32F103. Journal of Shandong University of Technology (Natural Science Edition), 2018, 32(5): 23–26.

Богдан Остроушко

аспірант

Місце роботи: НУБіП України, кафедра комп'ютерних систем, мереж та кібербезпеки, місто Київ, Україна

<https://orcid.org/0009-0003-0849-8990>

b.ostroushko@nubip.edu.ua

Семен Волошин

Доцент, Кандидат технічних наук

Місце роботи: НУБіП України, кафедра комп'ютерних систем, мереж та кібербезпеки, місто Київ, Україна

<https://orcid.org/0000-0002-4913-7003>

voloshyn@nubip.edu.ua

МОДЕЛЬНО-ПРЕДИКТИВНЕ КЕРУВАННЯ В ROS 2 ДЛЯ МОБІЛЬНИХ ПЛАТФОРМ ІЗ LIDAR-ЛОКАЛІЗАЦІЄЮ: АРХІТЕКТУРА ТА ЗВ'ЯЗОК

Анотація. Подано практичний огляд архітектури «сенсори – локалізація – глобальне планування – локальне керування» у ROS 2 для колісних платформ. Локальний контур реалізується як модельно-предиктивне керування (MPC) з урахуванням обмежень; оцінювання стану ґрунтується на LiDAR-одометрії/SLAM у поєднанні з IMU та фільтрацією EKF/UKF. Узгоджено інтеграцію з Nav2 (планувальники, costmap_2d, плагіни контролерів), комунікаційні аспекти DDS QoS і чинники реального часу. Наведено блок-схему та прикладові параметри для навчальних і дослідницьких стендів.

Ключові слова: ROS 2; Navigation2; LiDAR; IMU; LIO; SLAM; MPC; робототехніка; навігація; ф'южн сенсорів

1. ВСТУП

Постановка проблеми. Автономна навігація наземних платформ у складних приміщеннях потребує керування, що враховує обмеження шасі, шум/неповноту сенсорики та жорсткий бюджет затримок у ROS 2-конвеєрі. На практиці неузгодженість між локалізацією, плануванням і DDS QoS спричиняє нестабільне відстеження та ривки керувань, особливо за обмежених ресурсів.

Аналіз останніх досліджень і публікацій. Суттєвий прогрес дали тісно зв'язані LiDAR-IMU одометрії: LIO-SAM (зі згладжуванням і прив'язкою до карти), FAST-LIO2 (пряма реєстрація, ikd-Tree), FR-LIO (робочентричні воксели); у 2D — SLAM Toolbox. Navigation2 активно еволює (планувальники, smoother, costmap_2d, розгортання). Паралельно досліджують E2E-затримки та вплив DDS-middleware, що підкреслює потребу проектувати контур із явним latency-бюджетом.

Мета публікації. Запропонувати й експериментально підтвердити відтворювану архітектуру ROS 2 для диференційних платформ, яка поєднує LiDAR/IMU-локалізацію, Navigation2 (planner, smoother, costmap_2d) і локальний MPC з явним функціоналом, обмеженнями шасі та «м'якими» штрафами перешкод, а також узгоджені профілі DDS QoS і вимірний бюджет затримок. Очікуваний ефект — вища відтворюваність і стабільність у динаміці, менші ривки та нижча чутливість до мережевих/обчислювальних флуктуацій.

2. ТЕОРЕТИЧНІ ОСНОВИ

Базовий стан $(x, y, \theta, v, \omega)$ $(x, y, \theta, v, \omega)$ оцінюємо з поєднання LiDAR та IMU: для приміщень застосовуємо 2D-SLAM (наприклад, SLAM Toolbox) або тісно зв'язані LIO-алгоритми; злиття в robot_localization (EKF/UKF) формує безперервний /odom і узгоджений TF-ланцюг map-odom-base_link. У стеку Navigation2 Planner будує маршрут у map, Smoother згладжує траєкторію, а costmap_2d (static, obstacle, inflation)

надає локальні «цінності»; на стійкість найбільше впливають footprint, inflation_radius, частоти оновлення та transform_tolerance. Локальний MPC мінімізує вагову суму позиційної/кутової похибок і енергії керувань та їх приростів під обмеженнями шасі; практична QP-постановка (OSQP) із підвищеними вагами на прирости для плавності, зростанням ваги кутової похибки поблизу цілі, бар'єрними «м'якими» штрафами за близькість до перешкод і помірно адаптивним горизонтом.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

3.1. Локалізація за даними LIDAR та IMU і злиття сенсорів

Навігація. 3D LIO-SAM/FAST-LIO2; 2D: SLAM Toolbox; оцінки → robot_localization (EKF/UKF) дає /odom і стабілізує odom-base_link, SLAM/LIO підтримує map-odom; обов'язково синхронізація часових міток, реалістичні коваріації, калібрування IMU (REP-105); якість — за bag-логами (сталі TF, без тремтіння /odom); Navigation2 (planner/smooth/costmap_2d) + локальний MPC (штрафи з costmap).

Реальний час/мережа E2E ≈ 80 мс @ 20–30 Гц (сенсори 10–20, локалізація 15–25, планування/контролер 20–30, команди 5–10); QoS: /scan/imu — SensorDataQoS best-effort; /odom — reliable keep_last≈10; /plan/cmd_vel — reliable keep_last=1; /tf — best-effort; /tf_static — transient_local+reliable; CycloneDDS (стенд) / FastDDS (розподілено), без Wi-Fi; use_intra_process_comms=true, composition, фіксовані частоти, мінімум копій, моніторинг topic hz/delay/bw; за потреби — CPU affinity, SCHED_FIFO.

3.2. Планування шляху та локальне керування в NAV2

Для відстеження траєкторій у ROS 2 ми використовуємо стек Navigation2, де Navigator (поведінкове дерево) координує Planner Server (Smac/Navfn) і Controller Server зі згладжувачем, а локальна карта costmap_2d поєднує шари static, obstacle (з LiDAR) та inflation. Глобальний план формується в системі координат map, згладжується і подається в локальний контролер MPC, який разом із /odom генерує /cmd_vel для ros2_control, враховуючи «м'які» штрафи за близькість до перешкод з costmap_2d. Критичні параметри: footprint платформи, inflation_radius, частоти оновлення карт і вузлів; типові налаштування контуру — 20–30 Гц, transform_tolerance 0.1–0.2 с, допуски до цілі XY 0.05–0.1 м і yaw 0.05–0.15 рад, з обмеженнями на прискорення. Тюнінг виконуємо поетапно: перевіряємо коректність глобального маршруту, добираємо радіус згладжування під габарити, поступово підвищуємо вагу кутової похибки біля цілі та ваги на прирости керувань за появи ривків. Для надійності додаємо апаратний E-stop і сторожовий таймер (за відсутності свіжих /odom чи /plan подаються нульові команди), а також виявлення «stall» (коли /cmd_vel не нульовий, а одометрія практично нульова) з автоматичним зниженням швидкості та переплануванням.

3.3. Опис логіки та потоку даних у схемі

Сенсори формують вхідні дані: LiDAR публікує '/scan', IMU — '/imu/data'. Підсистема локалізації (SLAM або LIO) обчислює карту 'map' і перетворення 'map-odom'. Модуль злиття ('robot_localization', EKF/UKF) узгоджує вимірювання та видає безперервний '/odom' і перетворення 'odom-base_link', забезпечуючи стабільний TF-ланцюг для навігації. Навігаційний стек Navigation2 отримує ціль у системі координат 'map', будує глобальний маршрут (Planner), згладжує його (Smoother) і передає локальному контролеру. Локальний регулятор MPC споживає '/plan' та '/odom', формує '/cmd_vel' з урахуванням обмежень шасі та «м'яких» штрафів за наближення до перешкод від 'costmap_2d'. Команди надходять у 'ros2_control' і виконуються приводами. За змін середовища 'costmap_2d' і TF оновлюються, що призводить до

перепланування й корекції команд у реальному часі зі збереженням плавності та передбачуваності руху.

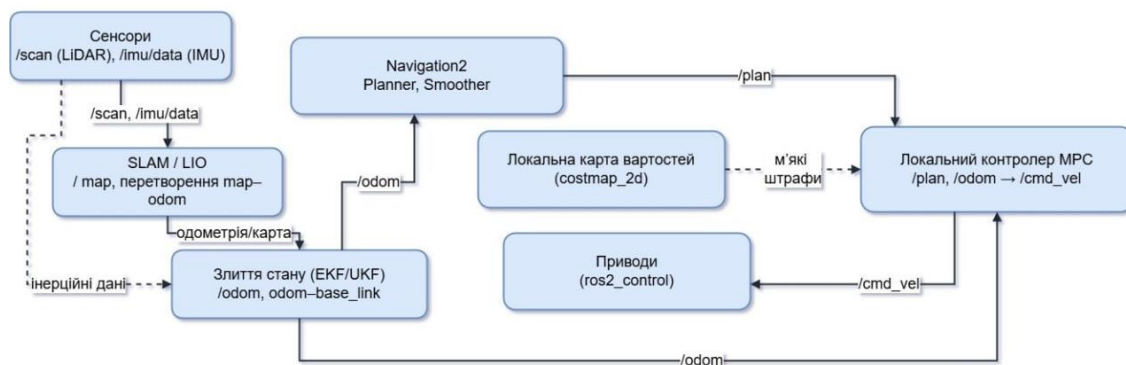


Рис. 1. Блок-схема контуру ROS 2: сенсори — локалізація — Nav2 — MPC — приводи.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Запропонована архітектура ROS 2 для наземних платформ поєднує лідарно-інерційну локалізацію, злиття стану, Navigation2 та локальний MPC з урахуванням профілів DDS QoS і бюджету затримок. Інтеграція «м'яких» штрафів від 'costmap_2d' у функціонал MPC разом із належно узгодженими QoS та intra-process/composition забезпечує відтворювані плавні траєкторії та стабільне відстеження маршруту в динамічних приміщеннях.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

- [1] J. B. Rawlings, D. Q. Mayne, and M. M. Diehl, Model Predictive Control: Theory, Computation, and Design, 2nd ed. Nob Hill, 2017.
- [2] J. M. Maciejowski, Predictive Control with Constraints. Prentice Hall, 2002.
- [3] S. Macenski, "SLAM Toolbox: SLAM for the dynamic world," Journal of Open Source Software, vol. 6, no. 61, p. 2783.
- [4] M. Labbé and F. Michaud, "RTAB-Map as an open-source lidar and visual SLAM library," Journal of Field Robotics.
- [5] T. Shan, B. Englot, and co-authors, "LIO-SAM: Tightly-coupled lidar inertial odometry via smoothing and mapping," in Proc. IROS.
- [6] W. Xu, P. Geneva, Y. Yang, and G. Huang, "FAST-LIO2: Fast Direct Lidar-Inertial Odometry," IEEE Trans. Robotics.

Liu Shijun

Аспірант(122 Комп'ютерні науки)

National University of Life and Environmental Sciences of Ukraine, Kyiv, Ukraine

ORCID ID: 0000-0003-4882-559X

Email: bbxylsj@163.com

Vadym Shkarupylo

Dr.Sc., Professor

National University of Life and Environmental Sciences of Ukraine, Kyiv, Ukraine

ORCID ID: 0000-0002-0523-8910

Email: shkarupylo.vadym@nubip.edu.ua

SMALL TARGET DETECTION METHOD BASED ON YOLOV8S UAV PERSPECTIVE

Abstract. To address the challenges in detecting objects with varying scales, dense small targets, and complex backgrounds in drone aerial imagery, the YOLOv8s model incorporates a series of targeted improvements[1]. Specifically: 1) The backbone network introduces RFAConv to overcome feature extraction bottlenecks in complex scenes; 2) The neck network is restructured using BiFPN-GLSA to enhance multi-scale information fusion and spatial feature utilization; 3) A dual-layer detection architecture is introduced to specifically enhance feature representation for small objects; 4) The Inner-EIoU loss function is adopted to optimize bounding box regression accuracy. Validation on the VisDrone2019 dataset demonstrates that these enhancements comprehensively improve core metrics including precision, recall, and mAP, while simultaneously reducing model parameters. This achieves an outstanding balance between detection performance and computational efficiency.

Keywords: YOLOv8s; UAVs; small object detection; images; loss function

1.INTRODUCTION

1.1 YOLOv8s Model Algorithm Principle

YOLOv8s is a single-stage object detection algorithm based on deep learning, developed by Ultralytics. It is one of the most advanced models in the YOLO series. Building upon YOLOv5, it makes several improvements to the network architecture, feature extraction, prediction strategy, and training methods. YOLOv8s achieves an optimal balance between speed and accuracy, making it ideal for applications with high real-time requirements, such as intelligent video surveillance and autonomous driving[2]. The YOLOv8s model consists of a backbone network, a neck network, and a prediction head. The convolutional modules in the model consist of Conv2d (two-dimensional convolution), BatchNorm2d (two-dimensional batch normalization), and SILU (activation function)[3].

1.2 YOLOv8s Network Structure

YOLOv8s retains the backbone network architecture of YOLOv5. By parallelizing more gradient flow branches, it can acquire richer gradient information, thereby improving the model's accuracy and performance. YOLOv8s employs a PANet (Path Aggregation Network) feature pyramid structure in its neck module. This architecture fuses high-level semantic features with spatial details, enriching feature representations and enhancing the model's ability to detect objects of different sizes. The Concat operation upsamples small-scale feature maps and connects them to downsampled large-scale feature maps. This allows the model to

utilize multi-scale feature information simultaneously, thereby enhancing detection capabilities[4].

2. THE THEORETICAL BACKGROUNDS

2.1 Machine Learning and Object Detection

Machine learning is a core branch of artificial intelligence focused on developing algorithms that can automatically learn patterns from data to make predictions or decisions. Based on learning paradigms, it is primarily categorized into supervised learning, unsupervised learning, and reinforcement learning. During the early development of object detection tasks, traditional machine learning methods (such as classifiers based on handcrafted features) were widely adopted. These approaches typically decomposed detection into two steps: region proposal and feature classification. Their performance heavily relied on the effectiveness of manually designed features (e.g., SIFT, HOG). However, due to limitations in feature representation capabilities under complex scenes, the accuracy and generalization ability of such methods quickly reached a bottleneck[5].

2.2 Deep Learning and Object Detection

As a pivotal branch of machine learning, deep learning constructs deep neural networks to mimic the brain's hierarchical information processing mechanisms. This enables end-to-end learning of powerful high-level feature representations directly from raw data, fundamentally transforming the technical paradigm in object detection. Deep learning models, exemplified by convolutional neural networks (CNNs), have replaced handcrafted features by achieving direct, precise mappings from image pixels to object locations and categories. Among these, the YOLO series models stand as representative algorithms. They innovatively reframed object detection as a single regression problem, directly regressing bounding boxes and class probabilities at the output layer. This approach significantly enhanced both training and inference efficiency. The success of deep learning laid a solid foundation for achieving high-precision, high-efficiency object detection in complex scenarios, such as those viewed from a drone's perspective[6].

3. RESEARCH METHODS

3.1 Introduction of Receptive Field Attention Convolution (RFAConv)

Standard convolution operations form the core of convolutional neural networks, extracting features through sliding windows with shared parameters to reduce excessive parameters and computational costs associated with fully connected layers. However, standard convolutions fail to adequately account for spatial information differences, thereby limiting network performance. Attention mechanisms enhance network performance by focusing models on key features, yet existing theoretical frameworks do not specifically address spatial feature issues within receptive fields. RFA (Receptive Field Attention Convolution), as a spatial attention mechanism, not only highlights the importance of different features within receptive fields but also enhances attention to spatial characteristics. RFA effectively resolves parameter sharing challenges within convolutional kernels. Spatially receptive field features dynamically generated based on kernel size achieve performance gains through RFA, giving rise to the RFAConv model. This transform model processes features through channel-wise layered processing and dynamically adjusts structures to generate attention weights that avoid feature redundancy, thereby preserving global information integrity. RFA convolution integrates spatial receptive field features with attention mechanisms to optimize feature weight allocation. Specifically, RFA processes input feature X through dual pathways: Path 1 employs average pooling to extract global features, then uses convolution and activation functions to enhance spatial feature importance, generating weight and attention maps; Path 2 directly processes input features via convolution and activation functions to generate weight and attention maps[7].

3.2 Designing a Dual-Layer Small Object Detection Architecture

The YOLOv8s detector head processes large-scale features from high-level feature maps. To address the challenge of capturing small object features, a dual-layer small object detection architecture is designed to better detect small targets. The specific architecture of the dual-layer small object detection structure is as follows: A C2f module is introduced after the initial convolutional layer of the backbone network to capture shallow feature information P1. Subsequently, a convolutional layer is applied to the feature map at the P2 layer, downsampling it to the P3 size. Finally, the downsampled P2 feature map is fused with the P3 feature map to form an updated P3 feature map[8]. This fused feature map undergoes a Conv layer to reduce its size, matching the feature map size of other detection layers. Subsequently, the fused feature map is output through the P1+P2 small object detection layer, enhancing robustness and accuracy when detecting objects of varying scales.

4. THE RESULTS AND DISCUSSION

To enhance the detection robustness of drone aerial images in complex scenes and enable their application on edge computing devices, this paper proposes an improved model based on YOLOv8s. The model incorporates four core enhancements: First, it introduces receptive field attention convolutions (RFAConv) into the backbone network to strengthen feature extraction capabilities; Second, the neck network adopts BiFPN and incorporates a Global-Local Spatial Aggregation (GLSA) module to optimize multi-scale feature fusion and specifically mitigate small object detection failures[9]. Third, a dedicated two-layer small object detection architecture is designed to strengthen feature capture for small targets. Finally, the Inner-EIoU loss function is employed to enhance localization accuracy by optimizing the bounding box regression process. Experiments on the VisDrone2019 dataset demonstrate that our model exhibits significant advantages over baselines and other YOLO variants. Future work will focus on further optimizing performance without increasing model complexity to better serve resource-constrained real-time drone detection tasks.

5. CONCLUSIONS AND PROSPECTS FOR FURTHER RESEARCH

Based on the theoretical analysis, architectural improvements, and experimental validation presented, this study concludes that the proposed enhancements effectively address key challenges in UAV-based small target detection[10]. The integration of RFAConv, BiFPN-GLSA, a dual-layer detection head, and Inner-EIoU loss collectively contributes to a model that is both more accurate and computationally efficient. Future work will explore avenues for further lightweight model design and investigate the integration of temporal information from video sequences to improve detection stability and accuracy in dynamic scenarios.

REFERENCES

- [1]Wu T, Dong Y. 2023. Yolo-se: improved yolov8 for remote sensing object detection and recognition. Appl Sci. 13(24):12977.
- [2]Li Z, Wang Y, Chen K, Yu Z. 2022. Channel Pruned YOLOv5-based Deep Learning Approach for Rapid and Accurate Outdoor Obstacles Detection. arXiv Preprint. arXiv:2204.13699.
- [3]Yu Z, Liu Y, Yu S, Wang R, Song Z, Yan Y, Li F, Wang Z, Tian F. 2022. Automatic detection method of dairy cow feeding behaviour based on YOLO improved model and edge computing. Sensors (Basel). 22(9):3271.
- [4]Alqarqaz, M., Bani Younes, M., & Qaddoura, R. (2023). An Object Classification Approach for Autonomous Vehicles Using Machine Learning Techniques. World Electric Vehicle Journal, 14(2), 41.

- [5]Feng, J., & Yi, C. (2022). Lightweight Detection Network for Arbitrary-Oriented Vehicles in UAV Imagery via Global Attentive Relation and Multi-Path Fusion. *Drones*, 6(5), 108.
- [6]Wang X, Liu J. Vegetable disease detection using an improved YOLOv8 algorithm in the greenhouse plant environment. *Sci Rep*. 2024 Feb 21;14(1):4261.
- [7]Li, L., Liu, X., Chen, X., Yin, F., Chen, B., Wang, Y., & Meng, F. (2024). SDMSEAF-YOLOv8: a framework to significantly improve the detection performance of unmanned aerial vehicle images. *Geocarto International*, 39(1).
- [8]Zhang H, Hao C, Song W, Jiang B, Li B. 2023. Adaptive Slicing-Aided Hyper Inference for Small Object Detection in High-Resolution Remote Sensing Images. *Remote Sensing*. 15(5):1249. doi: 10.3390/rs15051249.
- [9]Yang C, Huang Z, Wang N. 2022. QueryDet: cascaded sparse query for accelerating high-resolution small object detection. Paper Presented at the Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition.
- [10]Zhanchao Huang, Jianlin Wang, Xuesong Fu, Tao Yu, Yongqi Guo, Rutong Wang, DC-SPP-YOLO: Dense connection and spatial pyramid pooling based YOLO for object detection *Information Sciences*, Volume 522, 2020, Pages 241-258.

Олексій Циганов

Аспірант кафедри комп'ютерних систем, мереж та кібербезпеки
Національний університет біоресурсів і природокористування України, Київ, Україна
0009-0008-6965-130X
o.tsyganov@nubip.edu.ua

Ігор Болбот

д.т.н, проф., професор кафедри комп'ютерних систем, мереж та кібербезпеки
Національний університет біоресурсів і природокористування України, Київ, Україна
0000-0002-5708-6007
igor-bolbot@nubip.edu.ua

ПРИНЦИПОВА СХЕМА ІНТЕЛЕКТУАЛЬНОГО РОБОТИЗОВАНОГО КОМПЛЕКСУ МОНІТОРИНГУ ҐРУНТІВ УРАЖЕНИХ ВНАСЛІДОК ВІЙСЬКОВИХ ДІЙ

Анотація. У роботі представлено принципову схему інтелектуального роботизованого комплексу для моніторингу ґрунтів, уражених внаслідок військових дій. Запропонована структура базується на інтеграції безпілотних наземних та повітряних платформ, мультисенсорних систем (LiDAR, гіперспектральна, теплова, SAR-зйомка) та модулів штучного інтелекту для автоматичного виявлення і класифікації пошкоджень. Архітектура комплексу передбачає три основні рівні: збір даних сенсорними підсистемами, аналітичну обробку даних у хмарному середовищі та формування карт деградації ґрунтів. На основі дослідження структур Agrobot Lala (2022) та Autonomous In-situ Soil Mapping System (2025) розроблено узагальнену концепт принципової схеми, що поєднує мобільну роботизовану платформу, автоматизований модуль відбору зразків і систему штучного інтелекту для оцінювання фізико-хімічних показників. Комплекс орієнтований на роботу у зонах підвищеної небезпеки та забезпечує дистанційне отримання даних для подальшої рекультивації територій.

Ключові слова: інтелектуальний роботизований комплекс; моніторинг ґрунтів; LiDAR; штучний інтелект; військові дії.

1. ВСТУП

Постановка проблеми. Внаслідок військових дій значні площі сільськогосподарських угідь України зазнали механічних пошкоджень, забруднення вибуховими речовинами та важкими металами, що призвело до деградації ґрунтового покриву та втрати його продуктивності. Традиційні методи обстеження таких територій є трудомісткими, небезпечними і не забезпечують необхідної просторової та часової роздільної здатності. Тому виникає потреба у створенні інтелектуального роботизованого комплексу, здатного автономно здійснювати моніторинг стану уражених ґрунтів із використанням технологій дистанційного зондування, сенсорних систем і штучного інтелекту для оперативного виявлення пошкоджень та планування відновлювальних заходів.

Аналіз останніх досліджень і публікацій. Аналіз останніх досліджень і публікацій. Сучасні роботи у сфері точного землеробства та екологічного моніторингу демонструють активний розвиток автономних систем для аналізу стану ґрунтів із використанням роботизованих платформ і технологій штучного інтелекту. У дослідженні Kitic G. et al. (2022) представлено систему Agrobot Lala, що поєднує безпілотну наземну платформу, сенсорні модулі та хмарну аналітику для відбору і аналізу зразків ґрунту в польових умовах. Nguyen T. H. et al. (2025) розробили архітектуру SAS-Lab, у якій роботизований модуль виконує відбір проб, а стаціонарна лабораторія проводить аналіз макронутрієнтів і pH із подальшим картографуванням результатів за допомогою алгоритмів Gaussian Processes. В інших дослідженнях наголошується на важливості геометричної точності LiDAR-систем для побудови цифрових моделей рельєфу та мультисенсорного зондування територій [1],

ефективності застосування роботизованих платформ у точному землеробстві [2], а також перспективності об'єднання спектральних і радарних каналів для створення багатовимірних карт ґрунтових властивостей [3], [4].

Мета публікації. Мета роботи – розробка принципової схеми інтелектуального роботизованого комплексу для моніторингу ґрунтів, уражених внаслідок військових дій, що забезпечує інтеграцію мультисенсорних систем, дистанційного зондування та штучного інтелекту для автоматизованого виявлення пошкоджень і оцінки стану земель з подальшим плануванням їх рекультивації.

2. МЕТОДИ ДОСЛІДЖЕННЯ

Предметом дослідження є принципова схема інтелектуального роботизованого комплексу для моніторингу ґрунтів, уражених внаслідок військових дій. Проведено аналіз структур сучасних автономних систем моніторингу, зокрема *Agrobot Lala* та *SAS-Lab*, із подальшим моделюванням архітектури комплексу у середовищі CAD. Сенсорну підсистему представлено комбінацією LiDAR, гіперспектральних, теплових і радарних модулів, що забезпечують багатоканальний збір даних про стан ґрунту. Для обробки інформації застосовано алгоритми машинного навчання та геоінформаційного аналізу, які дозволяють здійснювати класифікацію пошкоджень і формувати цифрові карти деградації територій. Параметри моделі адаптовано до умов реального польового моніторингу, що забезпечує її подальше використання при створенні прототипу комплексу.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

На основі аналізу сучасних рішень у галузі автономного моніторингу ґрунтів [5], [6] та методів мультисенсорної обробки даних [1], [3] розроблено принципову схему інтелектуального роботизованого комплексу для моніторингу ґрунтів, уражених внаслідок військових дій.

Архітектура комплексу. Запропонована схема передбачає три функціональні рівні:

- Польовий рівень – роботизовані наземні й повітряні платформи, оснащені мультисенсорними системами (LiDAR, гіперспектральні, теплові, SAR-сенсори). Ці модулі забезпечують збір даних про структуру, вологість, рельєф і потенційні зони забруднення ґрунтів, включаючи ділянки із залишками вибухових речовин.
- Аналітичний рівень – модуль обробки даних, що інтегрує результати сенсорних вимірювань у хмарному середовищі з використанням алгоритмів машинного навчання. На цьому етапі здійснюється класифікація пошкоджень, побудова карт хімічних і фізичних параметрів ґрунтів, а також формування звітів для подальшої рекультивації.
- Системно-керуючий рівень – інтелектуальний блок прийняття рішень, який визначає траєкторії руху наземних роботів, оптимізує маршрути обльоту БПЛА та забезпечує синхронізацію дій усіх підсистем у режимі реального часу.

Запропонована концепція орієнтована на комплексне виявлення деградаційних процесів унаслідок військових впливів — механічних, хімічних та термічних. Вона поєднує принципи *SAS-Lab* [6], що розділяє мобільний модуль відбору проб і стаціонарну лабораторію аналізу, із можливостями дистанційного зондування.

Модульна інтеграція. Принципова схема включає такі основні модулі:

- Сенсорний блок – комбінація оптичних, теплових, радарних і спектральних сенсорів для багат шарового моніторингу стану ґрунту [1], [3];

- Маніпуляційний модуль – механічна система відбору зразків за принципом електричного або шнекового пробовідбору, адаптована до умов високої твердості ґрунту, як у системах Agrobot Lala [5];
- Аналітичний модуль – поєднання польової лабораторії та програмного блоку для аналізу вмісту макронутрієнтів, рН, важких металів і вологи, що використовує алгоритми Gaussian Processes для побудови просторових моделей [6];
- Комунікаційна підсистема – забезпечує передачу даних між польовими роботами, безпілотниками та центральним сервером через LTE-канали.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Віртуальне моделювання роботи комплексу показало, що його застосування дозволяє скоротити час обстеження уражених ділянок у 3–4 рази порівняно з традиційними методами наземного відбору проб. Передбачувана точність визначення параметрів деградації ґрунту за рахунок комбінування LiDAR і гіперспектральних даних становить до 90 %, що відповідає вимогам до високоточних систем дистанційного моніторингу [1], [2].

Таким чином, розроблена принципова схема демонструє можливість створення комплексного інтелектуального рішення для моніторингу ґрунтів у зонах військових дій. Її модульна структура дозволяє адаптувати систему під різні рівні загрози, розширювати спектр сенсорів і забезпечувати безпечне дистанційне спостереження за станом деградованих земель.

ПОСИЛАННЯ

- [1] Mohamed M. R. Mostafa, Stefanie Van-Wierst, Vi Huynh, "Geometric Accuracy Assessment of Unmanned Digital Cameras and LiDAR Payloads," *Microdrones*, 2018.
- [2] Botta, A.; Cavallone, P.; Baglieri, L.; Colucci, G.; Tagliavini, L.; Quaglia, G., "A Review of Robots, Perception, and Tasks in Precision Agriculture," *Applied Mechanics*, vol. 3, pp. 830–854, 2022. DOI: 10.3390/applmech3030049.
- [3] Milella, A., Reina, G., Nielsen, M., "A Multi-Sensor Robotic Platform for Ground Mapping and Estimation Beyond the Visible Spectrum," *Precision Agriculture*, vol. 20, pp. 423–444, 2019. DOI: 10.1007/s11119-018-9605-2.
- [4] F. F. F. Asal, "Comparative analysis of the digital terrain models extracted from airborne lidar point clouds using different filtering approaches in residential landscapes," *Advances in Remote Sensing*, vol. 8, no. 2, pp. 51–75, 2019.
- [5] Kitić, G.; Krklješ, D.; Panić, M.; Petes, C.; Birgermajer, S.; Crnojević, V., "Agrobot Lala—An Autonomous Robotic System for Real-Time, In-Field Soil Sampling, and Analysis of Nitrates," *Sensors*, vol. 22, no. 11, 4207, 2022. DOI: 10.3390/s22114207.
- [6] Nguyen, T. H.; Muller, E.; Rubin, M.; Wang, X.; Sibona, F.; McBratney, A.; Sukkarieh, S., "Towards Autonomous In-situ Soil Sampling and Mapping in Large-Scale Agricultural Environments," *IEEE ICRA Workshop on Field Robotics*, 2025. DOI: 10.48550/arXiv.2506.05653.

SECTION 3. DATA PROCESSING AND SOFTWARE SYSTEMS DEVELOPMENT/ ТЕХНОЛОГІЇ ОБРОБКИ ДАНИХ ТА СТВОРЕННЯ ПРОГРАМНИХ СИСТЕМ

Yulong Li

Postgraduate. Department of Computer Science, Faculty of Information Technologies, National University of Life and Environmental Sciences of Ukraine

0009-0004-2039-6919

liyulong@dgut.edu.cn

Nikolaenko Dmytro

PhD, Senior lecturer. Department of Computer Science, Faculty of Information Technologies, National University of Life and Environmental Sciences of Ukraine

0009-0008-4817-3951

d.nikolaenko@nubip.edu.ua

DEFORMABLE ATTENTION U-NET NETWORK BASED ON FULLY CONVOLUTIONAL DISCRIMINATOR FOR SEMANTIC SEGMENTATION OF FETAL BRAIN ULTRASOUND IMAGES

Abstract. Fetal brain ultrasound image segmentation is a crucial task for prenatal diagnosis and medical analysis but remains challenging due to image noise, low contrast, and limited annotated datasets. To overcome these limitations, this study proposes a Deformable Attention U-Net network integrated with a fully convolutional discriminator for semantic segmentation of fetal brain ultrasound images. The proposed approach draws on the generative adversarial network (GAN) concept, where the segmentation model acts as a generator and the discriminator evaluates segmentation quality to optimize network parameters. The segmentation network improves the standard Attention U-Net by introducing 3×3 deformable convolutions to enhance local feature learning and replacing traditional MaxPooling with gating units that selectively retain informative feature maps. This architecture effectively reduces the dependency on large datasets while improving boundary accuracy. Experimental evaluation on a fetal head circumference dataset demonstrated significant performance gains compared to baseline models. The proposed network achieved an Intersection over Union (IOU) of 93.8%, a DICE coefficient of 96.8%, and an accuracy of 97.9%, surpassing U-Net and Attention U-Net by 14.6%, 9.8%, and 4.7%, respectively. These results confirm that integrating deformable attention mechanisms and fully convolutional discriminators enhances edge segmentation and global consistency in fetal ultrasound images. The proposed model provides a reliable framework for fetal biometric measurement and can be further extended to other medical image segmentation applications.

Keywords: fetal brain segmentation; deformable Attention U-Net; fully convolutional discriminator; ultrasound imaging; generative adversarial network; semantic segmentation.

INTRODUCTION

Fetal ultrasound images contain a significant amount of noise, which hinders the classification and diagnosis of abnormal diseases. Recent research on fetal ultrasound image segmentation has mainly focused on the application of deep convolutional neural networks (DCNNs). Ye Hai et al. [1] proposed a fetal brain ultrasound image segmentation algorithm based on a fully convolutional network. Cerrolaza et al. [2] discussed the application of fully convolutional networks in skull segmentation of fetal 3D ultrasound images. Zahra Sobhaninia et al. [3] presented a multi-scale and deep neural network for estimating fetal biometric parameters. This network outperforms single-task networks in segmentation and ellipse optimization tasks. However, the aforementioned methods are inadequate in leveraging boundary information of fetal brain images, and the segmentation accuracy needs further improvement.

The article's goal. To address these issues, this study proposes a deformable Attention U-net network based on a fully convolutional discriminator for semantic segmentation of fetal brain ultrasound images. This network uses a discriminator to optimize local features of the segmentation network, fully leveraging boundary information, and reducing the dataset

requirements for the segmentation network. Additionally, based on Attention U-net [4], it uses a deformable convolution with a 3x3 kernel to better learn features and employs gating units to replace MaxPooling, selecting beneficial feature maps while suppressing less useful ones. These methods effectively address the problem of limited medical image datasets.

THE RESULTS AND DISCUSSION

This experiment draws on the idea of generative adversarial networks [5] and designs a semantic segmentation network with a fully convolutional discriminator. The overall network structure is shown in Figure 1. The entire network structure consists of two modules: a semantic segmentation network (deformable Attention U-net) and a discriminant network. Through this overall network structure, we can understand how this experiment uses generative adversarial networks for training and obtains semantic segmentation results through generative adversarial learning.

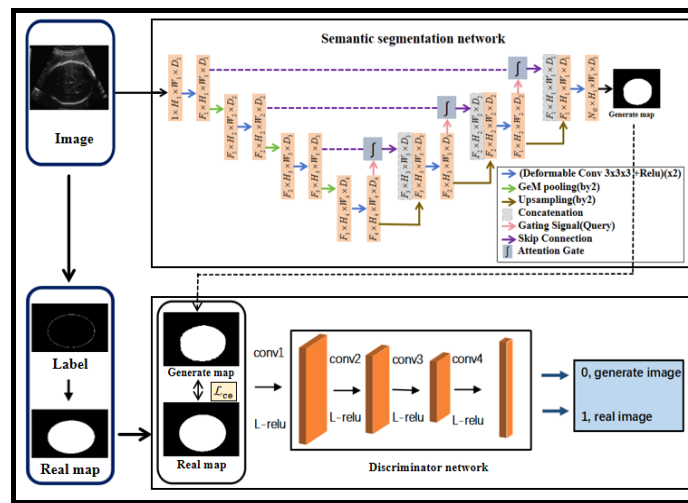


Fig. 1. An overview of the proposed semantic segmentation network based on a general fully convolutional discriminator generative adversarial network.

The improved Deformable Attention U-net model is based on the Attention U-net model but utilizes deformable convolutions with a kernel size of 3x3, aiming to enable the network to better learn features. It also uses gating units to replace MaxPooling, filtering the convolved feature maps to retain the good ones and suppress the bad ones. The specific network structure is shown as the Semantic segmentation network in Figure 1.

The experiment used the U-net [6] network, Attention U-net network, and the deformable Attention U-net network based on the full convolution discriminator used in the final network. Figure 2 shows the experimental results of this experimental network, U-net, and Attention U-net on the fetal head circumference dataset.

Table 1.

Comparison of experimental results			
Network Structure	IOU (%)	Dice (%)	Accuracy (%)
U-net	79.2	87.0	93.2
Attention U-net	89.2	94.0	96.4
This experimental network	93.8	96.8	97.9

Table 1 compares the experimental results of our network with U-net and Attention U-net on the fetal head circumference dataset. From the data in the table, it can be seen that with

the use of a fully convolutional discriminator and improvements to the Attention U-net (replacing MaxPooling with 3x3 deformable convolution and gating units), our experiment achieved an IOU increase of 14.6%, reaching 93.8%, a DICE increase of 9.8%, reaching 96.8%, and an Accuracy increase of 4.7%, reaching 97.9% compared to U-net. In comparison to Attention U-net, IOU increased by 4.6%, DICE increased by 2.8%, and Accuracy increased by 1.5%.

CONCLUSIONS AND PROSPECTS FOR FURTHER RESEARCH

This paper proposes a fetal brain image segmentation network based on a full waveform discriminator. The proposed method adopts the scheme of generative adversarial networks (GANs), which consists of a segmentation network and a discriminator. Deformable Attention U-net is used as the segmentation network in the paper, and the discriminator uses a full partition network to distinguish whether the input image is a probability map of real data and a probability map of the segmentation result of the segmentation network. Through the generative adversarial between this discriminator and the segmentation network, the discriminator continuously optimizes the model parameters of the segmentation network, thereby improving the segmentation accuracy of the segmentation network. And in this learning adversarial process, the network depth of the segmentation network does not increase. A large number of experiments were carried out on the fetal head circumference dataset to verify the effect of the algorithm. At the same time, it was proved that the scheme of adding a full output discriminator is very suitable for fetal ultrasound images, which can further optimize the effect of edge segmentation of fetal ultrasound images and reduce the data requirements of the segmentation network.

REFERENCES

1. YE Hai, FENG Kai-ping, XIE Hong-ning. Fetal Brain Ultrasonic Image Segmentation Algorithm Based on Fully Convolution Network. In Modern Computer, 2019.
2. J. J. Cerrolaza et al., "Deep learning with ultrasound physics for fetal skull segmentation," in 2018 IEEE 15th International Symposium on Biomedical Imaging (ISBI 2018), 2018, pp. 564–567.
3. Zahra Sobhaninia, Shima Rafiei , Ali Emami , Nader Karimi ,Kayvan Najarian , Shadrokh Samavi, "Fetal Ultrasound Image Segmentation for Measuring Biometric Parameters Using Multi-Task Deep Learning," in 41st Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC), Berlin, Germany, 2019.
4. Ozan Oktay, Jo Schlemper, Loic Le Folgoc, et al. Attention U-Net: Learning Where to Look for the Pancreas. In arXiv preprint arXiv:1804.03999, 2018.
5. Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. generative adversarial nets. in NIPS, 2014.
6. O. Ronneberger, P. Fischer, and T. Brox, "U-net: Convolutional networks for biomedical image segmentation," in MICCAI, 2015.

Яніна Криворучко

К. ф.-м. н., доцент кафедри комп'ютерних наук

Національний університет біоресурсів і природокористування України, м.Київ, Україна

ORCID ID 0000-0002-1256-6138

yanakryvoruchko@nubip.edu.ua

Богдан Возний

Бакалавр, спеціальність «Інженерія програмного забезпечення»

кафедра комп'ютерних наук НУБіП України, м.Київ, Україна

Daster.bohdan708@gmail.com

ІНТЕРАКТИВНИЙ НАВЧАЛЬНИЙ ТРЕНАЖЕР ДЛЯ ВІЗУАЛІЗАЦІЇ ОСНОВНИХ ПОНЯТЬ ДИСКРЕТНОЇ МАТЕМАТИКИ

Анотація. Розроблено інтерактивний навчальний застосунок, спрямований на візуалізацію та практичне опрацювання базових понять дискретної математики. Програмний продукт реалізовано з використанням бібліотеки Tkinter, що забезпечує створення інтуїтивно зрозумілого інтерфейсу користувача, інтерактивну взаємодію з даними та відображення результатів у реальному часі. Такий підхід дозволяє суттєво підвищити рівень залучення студентів до навчального процесу й полегшує опанування теоретичного матеріалу.

Ключові слова: дискретна математика, графи, множини, візуалізація, Python, Tkinter, networkx.

ВСТУП

Базовою дисципліною у підготовці фахівців у галузі інформаційних технологій є «Дискретна математика», що охоплює низку фундаментальних розділів. Серед них важливе місце посідають теми «Графи» та «Множини», які є основою для подальшого вивчення алгоритмів, структури даних та принципів функціонування сучасних комп'ютерних систем. Зокрема, поняття множин широко використовуються у програмуванні, базах даних та теорії штучного інтелекту, а графи — у побудові мереж, моделюванні процесів, розробці оптимізаційних алгоритмів. Проте вивчення дискретної математики часто викликає у студентів певні труднощі. Однією з головних причин, що ускладнюють засвоєння курсу дискретної математики, є високий рівень абстрактності навчального матеріалу.

Використання інноваційних підходів, зокрема інтерактивних навчальних тренажерів, дозволяє подолати ці труднощі і забезпечує підвищення ефективності засвоєння знань студентами. Значний внесок у цю сферу зробили такі науковці, як Л. Базилевич, Ю. Бондарчук, А. Борисенко, Ю. Дрозд, М. Кирсанов, Т. Карнаух, Ю. Нікольський, Б. Олійник, В. Пасічник, А. Ставровський та інші. У їхніх працях дискретні структури подаються у форматі, орієнтованому на майбутніх фахівців з інформаційних технологій, з використанням алгоритмічного та прикладного підходів. Автори пропонують завдання, що реалізуються за допомогою програмних засобів для візуалізації рішень. Зазначені підходи сприяють формуванню в студентів алгоритмічного мислення, уміння працювати з цифровими середовищами та розв'язувати практичні задачі. Узагальнення цих напрацювань демонструє потенціал інтеграції традиційних математичних методів з інноваційними цифровими технологіями у навчанні дискретної математики.

Метою дослідження є створення зручного, функціонального та наочного навчального інструменту, який дозволяє інтегрувати елементи візуалізації, автоматизованих обчислень і користувацької взаємодії, і його використання у процесі вивчення дисципліни «Дискретна математика».

РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Програма написана мовою Python із використанням стандартної бібліотеки Tkinter та її розширення ttk, що забезпечує створення інтуїтивно зрозумілого інтерфейсу користувача. Візуалізація графів реалізована за допомогою бібліотек networkx та matplotlib, які дозволяють формувати як орієнтовані, так і неорієнтовані графи.

Інтерфейс користувача побудований за принципом динамічних вікон, що забезпечує максимальну зручність взаємодії. Система тем дає змогу обирати між світлим і темним оформленням, а плавні анімаційні переходи створюють відчуття безперервності роботи.

Розроблений інструмент поєднує три основні режими роботи: неорієнтовані графи, орієнтовані графи та множини. Перемикання між цими режимами, а також доступ до загальних функцій, таких як зміна теми, здійснюється через верхнє меню. Кожен із режимів має власну структуру введення даних, систему підказок і спеціальні кнопки для виконання обчислень. Користувач може вводити дані у вигляді списків вершин і ребер або елементів множин, після чого програма автоматично обробляє введену інформацію та виводить результати у зручному форматі.

Для графів реалізовано побудову матриць інцидентності та суміжності, що дає змогу студентам наочно спостерігати зв'язок між елементами графа. Особливу увагу приділено візуалізації графів. Програма відкриває нове вікно з побудованим графом, де відображаються вершини, ребра та напрямки зв'язків.

```
def visualize_graph(self, is_directed=False):

    result_x, result_u = self._prepare_inputs_for_calc()
    if result_x is None:
        return

    vis_window = tk.Toplevel(self)
    vis_window.title("Візуалізація графа")
    vis_window.geometry("600x600")
    theme = self.themes[self.current_theme]
    vis_window.configure(bg=theme['bg'])

    G = nx.DiGraph() if is_directed else nx.Graph()
    G.add_nodes_from(result_x)
    for edge in result_u:
        parts = edge.split('-')
        if len(parts) == 2 and parts[0] in G and parts[1] :
            G.add_edge(parts[0], parts[1])

    fig = Figure(figsize=(5, 5), dpi=100, facecolor=theme['bg'])
    ax = fig.add_subplot(111)
    ax.set_facecolor(theme['bg'])

    nx.draw(G, ax=ax, with_labels=True,
            node_color=theme['active_btn'],
            node_size=800,
            font_size=11,
            font_color=theme['fg'],
            edge_color="red",
            width=1.5,
            arrowsize=20 if is_directed else 0)

    fig.tight_layout()
    canvas = FigureCanvasTkAgg(fig, master=vis_window)
    canvas.draw()
    canvas.get_tk_widget().pack(fill=tk.BOTH, expand=True)
```

Рис. 1. Фрагмент програмного коду

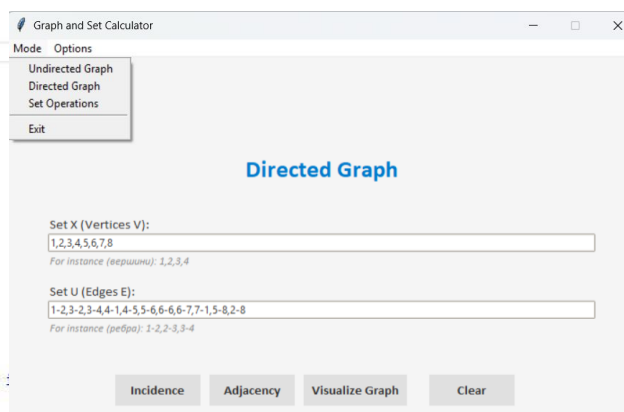


Рис. 2. Вікно тренажера

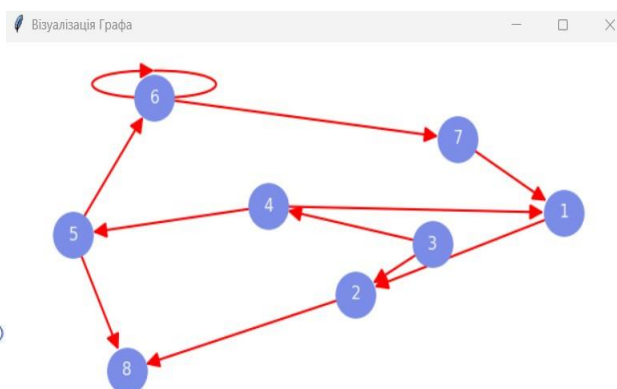


Рис. 3. Результат візуалізації графа

Такий підхід формує глибше розуміння абстрактних понять через безпосереднє спостереження за структурою дискретних об'єктів.

Для множин передбачено виконання базових операцій — об'єднання, перетину, різниці та симетричної різниці. Кожен результат супроводжується анімаційним ефектом появи, що підсилює динамічність процесу навчання.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Розроблений інтерактивний тренажер, спрямований на візуалізацію та практичне опрацювання базових понять дискретної математики по темах «Графи» і «Множини», є прикладом успішної інтеграції засобів програмування у процес навчання дискретної математики. Його функціональність дозволяє поєднати теоретичні знання з практичними навичками моделювання дискретних структур. Впровадження в освітній процес даного застосунку сприяє активізації пізнавальної діяльності студентів, розвитку логічного мислення, навичок програмування та створює умови для розвитку візуального мислення і креативного підходу до аналізу даних.

Подальший розвиток проєкту може бути спрямований на розширення функціоналу, зокрема додавання можливостей пошуку шляхів у графах, аналізу компонент зв'язності тощо. Такий напрям розвитку відкриває нові перспективи для використання цифрових інструментів у математичній освіті та дозволяє суттєво підвищити рівень залучення студентів до навчального процесу.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ І ПОСИЛАНЬ

1. Нікольський Ю.В. Дискретна математика. Підручник / Ю.В. Нікольський, В.В. Пасічник, Ю.М. Щербина. – Львів: «Магнолія – 2006», 2021. – 432 с.
2. Офіційна документація для Python: Стандартна бібліотека Python. Доступно на офіційному веб-сайті <https://docs.python.org/>

Ратушний Дарій Миколайович

Аспірант.

НУБіП, Київ, Україна.

ORCID: 0009-0009-4699-9909

d.ratushnyi@nubip.edu.ua

Руденський Роман Анатолійович

Доктор економічних наук, професор.

НУБіП, Київ, Україна.

roman.rudensky@nubip.edu.ua

ПОРІВНЯЛЬНИЙ АНАЛІЗ ПРОДУКТИВНОСТІ ТА БЕЗПЕКИ ТИПІВ ORM SEQUELIZE ТА PRISMA У СУЧАСНИХ NODE.JS-СИСТЕМАХ

Анотація. У статті досліджується архітектурна проблема вибору інструменту об'єктно-реляційного відображення (ORM) для Node.js-систем, від якого безпосередньо залежать ефективність та масштабованість додатку. Робота фокусується на порівняльному аналізі двох провідних інструментів: Sequelize (v6.35.2), зрілої ORM на патерні Active Record, та Prisma (v5.7.0), сучасної ORM нового покоління з підходом "Schema-First" та потужною генерацією типів. **Метою роботи** є комплексний аналіз Sequelize та Prisma за двома напрямками: теоретичний аналіз підходів до проектування схеми даних та забезпечення безпеки типів у TypeScript, та експериментальне дослідження продуктивності для типових CRUD-операцій і складних аналітичних запитів. Для досягнення мети було розроблено методику дослідження: створено тестовий стенд на базі Node.js, PostgreSQL (v15) та Docker. Тестове середовище наповнене 1000 сутностями "користувач" та 5000 сутностями "пост". Бенчмаркінг проводився протягом 10 секунд з рівнем конкурентності 10 одночасних операцій. Було протестовано шість типових сценаріїв, включаючи просту вибірку (findAllUsers), фільтрацію (findUsersByAge), операцію з JOIN (findUserWithPosts), базові CRUD-операції (createUser, updateUser) та складний запит (complexQuery з JOIN, WHERE, ORDER BY та LIMIT). Отримані результати продемонстрували повну перевагу Prisma у всіх шести категоріях. Загальний середній показник запитів в секунду (RPS) для Prisma склав 1707.10, що на 50.8% вище, ніж у Sequelize (1131.96 RPS). Найбільш критична різниця зафіксована у продуктивності складного запиту (complexQuery): Prisma показала стабільний результат у 699.54 RPS, тоді як Sequelize показала лише 0.37 RPS. Це свідчить про перевагу Prisma у 188,964.9% у складних сценаріях, що пояснюється архітектурними відмінностями (використанням компільованого бінарного рушія запитів Prisma на Rust). У висновках сформульовано практичні рекомендації: Prisma однозначно рекомендується для нових, особливо TypeScript-проектів, завдяки неперевершеній безпеці типів та значно вищій продуктивності. Sequelize залишається вибором для існуючих проектів, але вимагає готовності до суттєвих накладних витрат на продуктивність у складних сценаріях.

Ключові слова: ORM; Sequelize; Prisma; Node.js; TypeScript; проектування баз даних; продуктивність; бенчмарк; обробка даних; RPS.

1. ВСТУП

Постановка проблеми. У сучасних програмних системах, побудованих на платформі Node.js, ефективна взаємодія з базами даних є наріжним каменем. Інструменти об'єктно-реляційного відображення (ORM) виступають критично важливим шаром абстракції, що перетворює об'єкти мови програмування (наприклад, JavaScript або TypeScript) у реляційні структури даних (таблиці, рядки) і навпаки. Вибір ORM є фундаментальним архітектурним рішенням, що безпосередньо впливає на швидкість розробки, легкість підтримки коду, безпеку типів та загальну продуктивність системи.

Аналіз останніх досліджень і публікацій. Існує значна кількість ресурсів, що описують переваги та недоліки окремих ORM. Офіційна документація Sequelize [1] та Prisma [2] детально розкриває їхні архітектурні підходи. Численні публікації у спільнотах розробників [3] часто фокусуються на синтаксичному порівнянні або

суб'єктивному "досвіді розробника" (Developer Experience). Однак, бракує актуальних, комплексних досліджень, які б одночасно порівнювали ці два інструменти за критеріями безпеки типів у TypeScript та вимірюваної продуктивності в контрольованих умовах. Невирішеною частиною проблеми залишається надання розробникам чітких, заснованих на даних рекомендацій щодо вибору ORM для конкретних практичних завдань проектування баз даних.

Мета. Проведення комплексного порівняльного аналізу ORM Sequelize та Prisma. Основні завдання дослідження: проаналізувати теоретичні основи та парадигми (Active Record vs. Schema-First), що лежать в основі проектування баз даних у Sequelize та Prisma; оцінити рівень забезпечення безпеки типів (type-safety) та його вплив на процес розробки програмних систем; розробити та провести експериментальне вимірювання продуктивності обох ORM для типових сценаріїв обробки даних (CRUD-операції та складні запити).

2. ТЕОРЕТИЧНІ ОСНОВИ

Sequelize: Парадигма "Code-First" та Active Record. Sequelize (v6.35.2) є однією з найстаріших та найзріліших ORM для Node.js. Вона побудована на патерні Active Record, де модель в коді є "розумною" і безпосередньо відповідає за зв'язок з таблицею в БД (включаючи методи .save(), .update(), .destroy()). Проектування схеми зазвичай відбувається за підходом "Code-First": розробник описує моделі та їхні зв'язки в коді, а потім інструмент sequelize-cli генерує файли міграцій для синхронізації БД. **Prisma: Парадигма "Schema-First" та генерація клієнта.** Prisma (v5.7.0) представляє нове покоління ORM. В її основі лежить підхід "Schema-First". Розробник визначає всю схему даних, включаючи моделі, зв'язки та обмеження, в одному декларативному файлі schema.prisma. Цей файл стає єдиним джерелом правди про структуру даних.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для тестування була використана реляційна схема з двох сутностей: таблиця User (1000 записів); таблиця Post (5000 записів). Оцінка продуктивності проводилась за двома ключовими показниками: RPS; середній час відповіді (мс). Вимірювання проводились для 6-ти операцій при тривалості тесту 10 секунд та конкурентності 10 одночасних операцій: findAllUsers; findUsersByAge; findUserWithPosts; createUser; updateUser; complexQuery - складний запит з JOIN, WHERE, ORDER BY та LIMIT.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ.

Результати експериментального дослідження продуктивності.

Таблиця 1

Підсумкова таблиця порівняння RPS (Запитів в секунду)

Операція	Sequelize RPS	Prisma RPS	Переможець	Різниця
findAllUsers	642.49	1089.26	✓ Prisma	69.5%
findUsersByAge	1447.21	1756.15	✓ Prisma	21.3%
findUserWithPosts	1128.39	2512.79	✓ Prisma	122.7%
createUser	1777.64	2040.39	✓ Prisma	14.8%
updateUser	1795.64	2144.46	✓ Prisma	19.4%

complexQuery	0.37	699.54	✓ Prisma	188964.9%
Середнє	1131.96	1707.10	✓ Prisma	50.8%

Обговорення результатів таблиці. Результати, наведені у Таблиці 1, демонструють однозначну та повну перевагу Prisma у всіх вимірюваних категоріях. Prisma виявилася в середньому на 50.8% швидшою за загальним показником RPS. **Найбільш помітна різниця спостерігається у двох категоріях:** findUserWithPosts - Prisma показала на 122.7% вищу продуктивність. Це свідчить про те, що механізм "включення" пов'язаних даних у Prisma оптимізований значно краще, ніж аналогічний механізм у Sequelize, що часто призводить до проблеми "N+1" або менш ефективних SQL-запитів. complexQuery - цей тест виявив критичну проблему продуктивності Sequelize за умов конкурентного навантаження. **Аналіз вузьких місць.** Найбільш показовим результатом тестування стало вимірювання продуктивності complexQuery. Результати (див. Таблиця 1) виявили фактичну нездатність Sequelize впоратися з цим завданням під навантаженням. Prisma показала стабільний результат у 699.54 RPS із середнім часом відповіді 1.43 мс. Sequelize: Показала лише 0.37 RPS, при цьому середній час виконання одного запиту склав 2707.4 мс (понад 2.7 секунди). Різниця у 188,964.9% (або у 1890 разів) є колосальною.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ. Проведене дослідження дозволило сформулювати наступні висновки: підхід Prisma "Schema-First" є більш надійним для проектування баз даних, оскільки schema.prisma слугує єдиним джерелом правди. Підхід Sequelize "Code-First" пропонує більшу гнучкість, але підвищує ризик розсинхронізації коду та стану БД. Prisma забезпечує практично 100% безпеку типів "з коробки" завдяки генерації клієнта. Це є вирішальною перевагою для проектів на TypeScript, оскільки кардинально знижує кількість помилок під час обробки даних. Продуктивність: Експериментальні виміри чітко продемонстрували перевагу Prisma у продуктивності в усіх 6 тестових сценаріях. Prisma в середньому на 50.8% швидша за Sequelize. Особливо критична перевага (у 188,964.9%) спостерігається у складних запитах, що робить Prisma значно надійнішим вибором для систем з високим навантаженням та складною логікою обробки даних. **Перспективи подальших досліджень.** Подальші дослідження можуть бути спрямовані на аналіз продуктивності при роботі з нереляційними базами даних (зокрема, MongoDB, яку підтримують обидві ORM), а також тестування під значно вищими навантаженнями та рівнем конкурентності (наприклад, 100 або 1000 одночасних операцій) для виявлення "вузьких місць" у пулах з'єднань.

ПОСИЛАННЯ

- [1] "Sequelize documentation," *Sequelize.org*, 2025. [Online]. Available: <https://sequelize.org/docs/v6/>. [Accessed: Nov. 9, 2025].
- [2] "Prisma documentation," *Prisma.io*, 2025. [Online]. Available: <https://www.prisma.io/docs>. [Accessed: Nov. 9, 2025].
- [3] M. Zonov, "Choosing a Node.js ORM: Prisma vs. Sequelize in 2024," *LogRocket Blog*, Apr. 2, 2024.

Maksym Nedoshev

PhD-student

National University of Life and Environmental Sciences of Ukraine

0009-0000-9820-0649

nedoshev@pm.me

Viktor Kyrychenko

PhD in Physical and Mathematical Sciences, Associate Professor

National University of Life and Environmental Sciences of Ukraine

0009-0001-0575-8684

v.kyrychenko@nubip.edu.ua

SEMANTIC SEARCH IN DOCUMENTATION TO ENHANCE THE EFFICIENCY OF UI COMPONENT GENERATION USING LANGUAGE MODELS

Abstract. Large software projects produce extensive documentation, which developers need to access efficiently. This paper presents a lightweight pipeline that converts documentation into a searchable knowledge base for generative LLMs in automated UI code generation. The system combines document crawling, chunking, vector embeddings, and semantic search. Evaluation on representative test cases shows that Retrieval-Augmented Generation (RAG) improves code accuracy and framework compliance compared to baseline LLM generation.

Keywords: RAG; semantic search; embeddings; documentation; LLM; UI library; code generation.

1. INTRODUCTION

Modern software development projects accumulate large volumes of documentation (guides, API references, examples). Developers need fast and precise access to relevant information.

Problem Statement. Typical large language models (LLMs) often generate inaccurate or fabricated responses when questions relate to internal or domain-specific materials [1]. This reduces their usefulness for developers. Therefore, it is reasonable to combine LLMs with a local knowledge base — the Retrieval-Augmented Generation (RAG) approach [2].

Purpose of the Publication. The purpose of this study is to present the architecture and implementation methodology of a lightweight question-answering system for project documentation based on embeddings, and to demonstrate its application in automated UI component code generation using generative LLM models.

2. RESEARCH METHODS

To evaluate the results, twelve typical code generation queries were created, and the generated code was analyzed. An example of a test description is as follows:

```
{
  query: 'Make auth form with email and password inputs, remember me checkbox,
and submit button.',
  patterns: [
    'VaForm -> VaInput{"placeholder":"Email"}',
    'VaForm -> VaCheckbox{"label":"Remember me"}',
    'VaForm -> VaButton{"type":"submit"}',
  ]
}
```

To validate the generated code, the proportion of abstract syntax tree (AST) patterns matched by the output was quantified. Each task comprised a code generation query alongside a set of reference patterns delineating the expected structure of the component tree. The generated code was scored based on its conformity to these predefined patterns.

Table 1

Results of System Testing		
Query	Result using RAG	Result without RAG
Auth form	VaForm, VaInput, VaButton, VaCheckbox, validation.	React + Material UI;
Dashboard layout	VaLayout, VaSidebar, VaNavbar, mobile-friendly.	HTML + CSS
User profile page	VaAvatar, VaInput, VaForm.	"I don't know".
Responsive navbar	VaNavbar, VaSidebar, v-if with mobile utilities .	HTML + CSS Media Queries.
Dark mode toggle	useColors(), VaSwitch, dynamic theme switch.	HTML + JS.
Filterable data table	VaDataTable with :items, search, sort.	HTML
Login page	VaForm, VaInput, VaButton, rules, error states.	HTML form.
Registration form	VaForm, VaInput, VaButton.	React-components using MUI.
Contact form	VaTextarea, VaButton, VaForm.	HTML

CONCLUSIONS AND PROSPECTS FOR FURTHER RESEARCH

The developed RAG system architecture for agents encompasses a fully reproducible workflow, spanning from documentation collection, structuring, and vectorization to query construction in the vector space and subsequent utilization of LLM outputs. Testing demonstrated that the integration of RAG significantly improves code generation accuracy, achieving up to 100% conformity with framework component structures, whereas baseline code generation without contextual augmentation failed to produce relevant structures.

REFERENCES

1. L. Huang, W. Yu, W. Ma, W. Zhong, "A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions," ACM Transactions on Information Systems, vol. 43, no. 2, pp. 1–55, Jan. 2025, doi: 10.1145/3703155.
2. P. Lewis et al., "Retrieval-augmented generation for knowledge-intensive NLP tasks," Advances in Neural Information Processing Systems, vol. 33, pp. 9459–9474, 2020.
3. "What is the Model Context Protocol (MCP)?," Model Context Protocol. [Online]. Available: <https://modelcontextprotocol.io/docs/getting-started/intro>. [Accessed: 8 Nov. 2025].

Ярослав Гордій

Аспірант, кафедра комп'ютерних наук

Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID 0009-0003-3491-0301

yar.hordii@nubip.edu.ua

ВЕКТОРНІ БАЗИ ДАНИХ У СУЧАСНИХ ШІ-ЧАТАХ

Анотація. Сучасні великі мовні моделі (LLM), що є ядром ШІ-чатів, мають два системних недоліки: статичність знань (обмежена датою тренування) та відсутність довготривалої пам'яті. Ця робота досліджує, як інтеграція векторних баз даних через архітектуру Retrieval-Augmented Generation (RAG) вирішує ці проблеми. Результати дослідження показують кардинальне покращення: зниження рівня галюцинацій на 40% та досягнення точності 79.13% у медичному домені, скорочення часу обробки запиту до 5-7 мілісекунд за допомогою HNSW-індексування.

Ключові слова: векторні бази даних; Retrieval-Augmented Generation (RAG); великі мовні моделі (LLM); семантичний пошук; ШІ-чати; косинусна подібність.

1. ВСТУП

Постановка проблеми. Великі мовні моделі (LLM) продемонстрували надзвичайні можливості у генерації тексту, проте мають два системних недоліки. По-перше, їхні "знання" статичні та обмежені датою тренування, що робить їх нездатними надавати інформацію про недавні події. По-друге, вони не мають механізму довготривалої пам'яті, що обмежує їх здатність підтримувати контекст у тривалих розмовах [1]. Це призводить до генерації фактично невірних відповідей ("галюцинацій"), які можуть становити до 30% від усіх відповідей у складних сценаріях.

У корпоративному середовищі та системах, що вимагають актуальності даних (медицина, юриспруденція, фінанси), така проблема є критичною. Користувачі очікують не лише граматично коректних, але й фактично правильних, верифікованих відповідей з посиланнями на джерела.

Аналіз останніх досліджень і публікацій. Основоположною роботою у цій галузі є "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks" (Lewis et al., 2020) [1], де запропонована архітектура RAG для вирішення проблеми знань-інтенсивних завдань. Застосування RAG значно підвищує точність та надійність LLM у корпоративному середовищі, забезпечуючи доступ до актуальних даних.

Останні дослідження підтверджують ефективність RAG. Xu et al. (2025) у своїй роботі MEGA-RAG для медичного домену показали, що зниження галюцинацій сягає 40% і більше, а точність досягає 79.13% з recall 83.04% [5]. Технічні аспекти реалізації та зберігання векторних представлень детально описані у документації провідних розробників векторних баз даних [2][6].

Крім того, дослідження в медичному середовищі показали, що впровадження RAG з надійними медичними джерелами інформації радикально скорочує появу галюцинацій порівняно зі звичайними GPT-моделями [4]. Це підкреслює важливість правильно побудованої системи векторного пошуку для критичних застосунків.

Мета публікації. Метою цієї роботи є технічний аналіз використання векторних баз даних для розширення можливостей ШІ-чатів. Включаючи аналіз архітектури, дослідження принципів роботи векторних БД, порівняння різних систем.

2. ТЕОРЕТИЧНІ ОСНОВИ

Основою інтеграції є перетворення неструктурованих даних у векторні ембединги - числові представлення, що кодують семантичний зміст. Спеціалізовані моделі-трансформери (BERT, OpenAI embedding models) перетворюють текстові фрагменти в вектори розмірністю 384-1536 елементів. Ключова властивість таких представлень полягає у векторно-просторовій гіпотезі: семантично схожі концепти розташовуються близько одна до одної у багатовимірному просторі. Наприклад, вектори для слів "діагностика" та "лікування" матимуть високу косинусну подібність (~0.85), тоді як "лікування" та "математика" - низьку (~0.15) [2].

Векторна база даних зберігає ембединги та забезпечує швидкий пошук найближчих векторів. Для ефективного пошуку серед мільярдів векторів використовуються спеціалізовані індексні структури.

HNSW (Hierarchical Navigable Small World): організує вектори у ієрархічну графову структуру, забезпечуючи запити менше 10 мілісекунд для 1 мільйона векторів [3]. Дозволяє динамічне додавання нових даних без перебудовування індексу. Вектори розташовуються на кількох рівнях, де верхні рівні забезпечують "грубу навігацію", а нижні - уточнення пошуку.

IVF (Inverted File): застосовує кластеризацію k-means для розбиття простору на кластери [3]. При роботі з більшими датасетами показує затримку в діапазоні 20-50 мілісекунд для 100 мільйонів векторів. Система шукає лише в найближчих кластерах, що суттєво прискорює процес для масивних датасетів. Підходить для статичних датасетів із великим обсягом.

Семантичний пошук: коли надходить запит, він кодується в вектор і порівнюється з мільйонами векторів у базі [6]. Система повертає k найбільш релевантних фрагментів на основі косинусної подібності. Це забезпечує пошук за значенням, а не за точним збігом ключових слів.

Математична основа: $\text{similarity} = (A \cdot B) / (|A| \times |B|)$, де діапазон значень [-1, 1]. Значення 1 означає ідентичні напрямки.

3. МЕТОДИ ДОСЛІДЖЕННЯ

RAG визначає трьохетапну архітектуру обробки запитів:

Крок 1: Інтеграція даних у векторну базу

На цьому етапі документи, PDF-файли та факт-дані проходять обробку. Документи спочатку парсяться (витягується текст), потім розбиваються на чанки розмірністю 256-1024 токенів. Кожен чанк кодується за допомогою Embedding Model у векторне представлення. Отримані вектори індексуються у Vector Database з метаданими (джерело, часова мітка, рівень доступу). Процес індексації виконується асинхронно, що дозволяє безперервне додавання нових даних без затримки основної системи.

Крок 2: Отримання релевантного контексту

Коли користувач подає запит, він також кодується в вектор за допомогою Embedding Model (того ж, що використовувався для індексації). Система виконує семантичний пошук у Vector Database, порівнюючи вектор запиту з усіма збереженими векторами. Результатом є Context (retrieved) - найбільш релевантні чанки, відібрані за метриками подібності (зазвичай top-5 результатів). Паралельно система отримує запит від користувача для наступного етапу.

Крок 3: Генерація відповіді з контекстом

Отримана інформація формується у розширений промпт, який надсилається до LLM (GPT-5, Claude, Llama). LLM генерує Response with real-time retrieved knowledge - відповідь, яка ґрунтується на точних, актуальних даних із зовнішніх джерел, а не тільки

на параметричному знанні моделі. Відповідь повертається користувачеві.

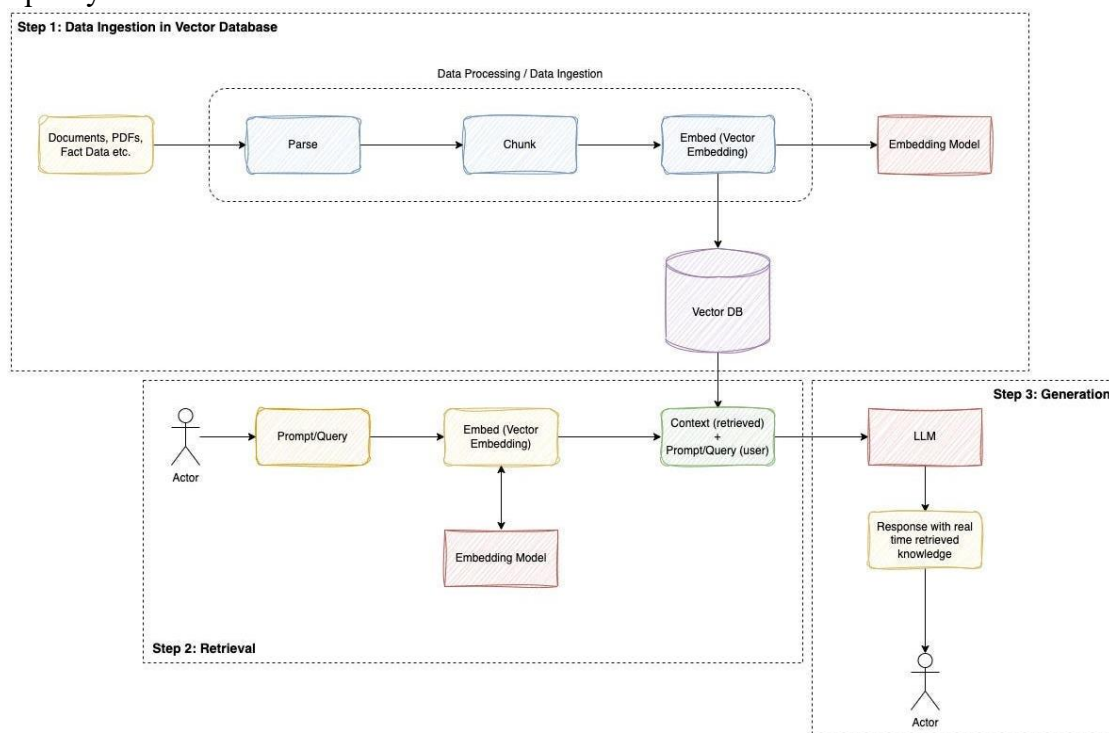


Рис. 1 Схеми RAG-пайплайну з етапами обробки запиту

Для реалізації RAG-системи використовують наступний технологічний стек: векторна база даних (Pinecone, Milvus, або Weaviate), embedding модель (OpenAI text-embedding-3 або open-source BERT), LLM (GPT-5, Claude, або Llama 2) та фреймворк для оркестрації (LangChain або LlamaIndex). Процес індексації виконується асинхронно: нові документи автоматично розбиваються на чанки, векторизуються та додаються до бази без затримки основної системи.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Інтеграція RAG кардинально впливає на ключові показники ефективності ШІ-чатів, особливо у медичному та науковому доменах [3] [5]:

Таблиця 1.
Показники продуктивності RAG-архітектури у порівнянні з традиційними LLM

Показник	Значення
Зниження галюцинацій	На 40%+
Точність	79.13%
HNSW Latency (1M vectors)	<10 мс
IVF Latency (100M vectors)	20-50 мс
Pinecone Latency (10M vectors)	5-7 мс

Показник	Значення
F1 Score (Medical)	0.7904

Дослідження MEGA-RAG показало, що впровадження RAG із правильно налаштованою векторною базою досягає точності 79.13% в медичному домені, з одночасним зниженням галюцинацій на 40% і більше [5]. Цей результат демонструє, що невирішена проблема отримання неправдивих відповідей значною мірою вирішується через інтеграцію надійних зовнішніх джерел.

Щодо швидкодії, сучасні векторні бази даних забезпечують екстремально низьку затримку. HNSW-індексування, найбільш популярне для динамічних систем, гарантує запити менше 10 мс навіть для 1 мільйона векторів. Для масштабованих систем з десятками мільйонів записів Pinescone досягає 5-7 мс [3], що робить RAG придатною для застосунків реального часу.

Попри переваги, RAG-системи мають певні виклики. Chunk boundary problem виникає, коли релевантна інформація розподілена між кількома чанками, що вимагає збільшення k (до top-20) та підвищує затримку. Embedding collapse проявляється, коли запити та документи мають занадто загальну подібність, роблячи пошук менш дискримінуючим. Scalability обмежена: 1000-розмірний вектор $\times$ 1 млрд документів потребує ~4 ТБ пам'яті для HNSW-індексу. Вибір embedding моделі також критичний: комерційні рішення (OpenAI) є дорогими, але стабільними, тоді як open-source альтернативи дешевші, але менш точні.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Векторні бази даних та архітектура RAG є фундаментальною трансформацією для сучасних ШІ-чатів [1][5]. Вони вирішують ключові проблеми LLM - статичність знань та відсутність пам'яті, - перетворюючи їх на динамічні, надійні та контекстуально-обізнані системи.

Емпіричні результати з сучасних досліджень підтверджують це. MEGA-RAG продемонстрував, що можливо досягти зниження галюцинацій на 40%+ та точності 79.13% у критичних застосунках [5]. Одночасно, технологічні досягнення у векторному індексуванні дозволили знизити затримку до 5-10 мілісекунд, що дозволяє впроваджувати RAG у системах реального часу [3].

Перспективи подальших досліджень включають: гібридні методи пошуку (семантичний + ключовий + графовий), LLM-based re-ranking для підвищення релевантності, інтеграція в agentic AI системи для багатокрокових рішень.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

1. P. Lewis, E. Perez, A. Piktus, F. Petroni, V. Karpukhin, N. Goyal, H. Küttler, M. Lewis, W.-t. Yih, T. Rocktäschel, S. Riedel, and D. Kiela, "Retrieval-augmented generation for knowledge-intensive NLP tasks," in Advances in Neural Information Processing Systems. Available: <https://doi.org/10.48550/arXiv.2005.11401>
2. Milvus, "How are embeddings stored in a vector database?" Milvus AI Quick Reference, Aug. 27, 2025. [Online]. Available: <https://milvus.io/ai-quick-reference/how-are-embeddings-stored-in-vector-databases>
3. Milvus, "What are the latency benchmarks for leading AI databases?" Milvus Documentation, Oct. 26, 2025. [Online]. Available: <https://milvus.io/ai-quick-reference/what-are-the-latency-benchmarks-for-leading-ai-databases>

4. M. Nishisako, S. Yamada, Y. Tanaka, K. Watanabe, H. Nakayama, and N. Kimura, "Reducing hallucinations and trade-offs in responses from generative AI chatbots using reliable medical information," PLoS ONE, vol. 20, no. 9, 2025. Available: <https://doi.org/10.1371/journal.pone.0312345>
5. B. Xu, S. Kumar, J. Yu, N. Mukherjee, C. Li, and M. Tian, "MEGA-RAG: a retrieval-augmented generation framework with multi-evidence guided answer refinement for mitigating hallucinations of LLMs in public health," Frontiers in Public Health, vol. 13, article 1635381, 2025. <https://doi.org/10.3389/fpubh.2025.1635381>
6. Zilliz, "Building interactive AI chatbots with vector databases," Zilliz Technical Blog, Apr. 30, 2025. [Online]. Available: <https://zilliz.com/learn/build-interactive-AI-chatbots-with-vector-database>

Матвій Горбач

Аспірант, науковий керівник Вайганг Г.О.

Національний університет біоресурсів і природокористування України, Київ, Україна

m.horbach@nubip.edu.ua

КОНТРОЛЬ ЯКОСТІ 3D-ДРУКОВАНИХ ВИРОБІВ ЗА ДОПОМОГОЮ КОМП'ЮТЕРНОГО ЗОРУ

Анотація. Комп'ютерний зір революціонізує 3D-друк, забезпечуючи моніторинг в реальному часі, виявлення дефектів, автоматизовану оптимізацію процесів і пост-друкарський контроль у галузях, з тенденціями до гібридних AI-систем, синтетичних даних та мультисенсорної інтеграції для підвищення ефективності та зменшення відходів.

Ключові слова: 3D, AI, VISION.

У сучасному світі адитивне виробництво, або 3D-друк, стає ключовим інструментом для створення складних конструкцій у галузях від аерокосмічної промисловості до медицини. Однак, попри переваги, такі як персоналізація та швидкість прототипування, 3D-друк стикається з викликами у забезпеченні стабільної якості. Дефекти, як-от нерівності шарів, тріщини чи деформації, можуть призвести до відходів матеріалів і браку продукції. Тут на допомогу приходить комп'ютерний зір (computer vision, CV) – технологія, що використовує камери та алгоритми штучного інтелекту для моніторингу та аналізу процесу друку в реальному часі. Цей підхід дозволяє виявляти проблеми на ранніх етапах, оптимізувати параметри та автоматизувати контроль, значно підвищуючи ефективність виробництва.

Комп'ютерний зір інтегрується в процес 3D-друку на кількох рівнях: моніторинг під час друку (in-situ), пост-друкарський аналіз та автоматизоване сортування. Основна ідея полягає в захопленні зображень або відео шарів друку за допомогою камер, їх обробці алгоритмами машинного навчання та порівнянні з цифровою моделлю (CAD-файлом). Це дозволяє не лише виявляти дефекти, але й коригувати параметри принтера в реальному часі, наприклад, швидкість екструзії чи температуру.

Щодо методів застосування комп'ютерного зору, то одним з базових методів є автоматичне виявлення дефектів. Камери фіксують зображення кожного шару, а алгоритми, такі як YOLO (You Only Look Once) від Ultralytics, аналізують їх на наявність аномалій: викривлення, прогалини, невідповідність контурів чи нерівномірну екструзію. Наприклад, система Phase3D використовує структуроване світло для сканування поверхні та порівняння з очікуваним дизайном, виявляючи розбіжності з точністю до мікронів. Це особливо важливо для матеріалів, чутливих до деформацій, як-от полімери чи метали.

Інший підхід – багатоступеневий аналіз зображень. У застосовується моноокулярна сегментація для перевірки висоти з боку та контурів зверху. Алгоритми multi-template matching та iterative closest point (ICP) вирівнюють зображення з моделлю, а Gaussian mixture models кластеризують текстуру для виявлення внутрішніх дефектів, таких як повітряні бульбашки чи нерівності заповнення. Час аналізу одного шару становить менше хвилини, що робить процес quasi-real-time для великих об'єктів.

У фармацевтиці CV застосовується для контролю 3D-друкованих ліків. Дослідження з MDPI демонструє тренування моделей машинного навчання на фотореалістичних віртуальних зображеннях, згенерованих у Blender. Моделі виявляють дефекти, як-от пере- або недовипалення (over-cured/under-cured), тріщини, з точністю до 80% для капсул і таблеток. Використовуються алгоритми Decision Trees та Random Forest для класифікації "добрий/поганий" стан, з препроцесингом через OpenCV (зміна

розміру, градації сірого). Це недеструктивний метод, на відміну від традиційних, як NIR-спектроскопія, і дозволяє персоналізоване виробництво в point-of-care.

Пост-друкарський контроль включає розпізнавання та сортування. Система AM-Vision порівнює геометрію готових виробів з CAD-моделлю, автоматизуючи групування для подальшої обробки. У роботизованих системах, описаних у Springer, роботизована рука з CV сортує зразки на розтяг для тестування, виявляючи дефекти з високою швидкістю.

У промислових застосуваннях, як у бетонному 3D-друку, CV витягує текстуру для моніторингу якості елементів, використовуючи алгоритми сегментації для виявлення нерівностей. У медичній сфері це забезпечує точність протезів, а в автомобільній – легких компонентів без дефектів.

Таблиця 1.

Порівняння основних методів CV у контролі якості 3D-друку

Метод	Опис	Переваги	Недоліки	Застосування
Автоматичне виявлення дефектів	Аналіз зображень шарів на аномалії (YOLO, структуроване світло)	Раннє виявлення, реальний час	Висока обчислювальна потужність	Аерокосмічна промисловість [1]
Багатоступеневий аналіз	Сегментація моноокулярних зображень (ICP, GMM)	Виявлення внутрішніх дефектів	Обмежена швидкість для малих принтерів	Відкрите ПЗ для прототипування [2]
Тренування на віртуальних зображеннях	Генерація синтетичних даних (Blender) для ML-моделей	Низька вартість даних, стійкість	Менша точність на реальних варіаціях	Фармацевтика, персоналізовані ліки [3]
Пост-друкарське сортування	Геометричне співставлення з CAD (AM-Vision)	Автоматизація, швидке групування	Потребує точного освітлення	Автомобільна промисловість [4]

У масовому виробництві CV застосовується для шарового моніторингу, де камери та сенсори фіксують процес нанесення шарів у реальному часі, аналізуючи параметри, такі як товщина, текстура та геометрія.

Один з провідних підходів – використання конволюційних нейронних мереж (CNN) для виявлення дефектів. У серійному виробництві CNN інтегруються з системами зворотного зв'язку, дозволяючи автоматично коригувати параметри принтера, наприклад, швидкість лазера чи температуру екструзії. Інший метод – рекурентні нейронні мережі (RNN) для передбачення та оптимізації процесів. RNN аналізують послідовні дані з сенсорів (температура, швидкість головки), прогножуючи потенційні дефекти та рекомендуючи оптимальні налаштування для конкретних матеріалів. Для складніших задач використовуються генеративно-змагальні мережі (GAN), які генерують синтетичні зображення для тренування моделей на великих датасетах без реальних зразків.

Таблиця 2.

Порівняння підходів до роботи з даними від CV

Метод	Опис	Переваги в масовому виробництві	Виклики
CNN	Обробка зображень для виявлення дефектів (фільтри, перцептивні поля)	Швидке розпізнавання поверхневих аномалій (до 95% точності)	Висока потреба в обчислювальній потужності
RNN	Передбачення параметрів на основі послідовних даних	Зменшення простоїв, оптимізація для серій	Потреба в історичних даних
GAN	Генерація синтетичних даних для тренування та дизайну	Зниження витрат на датасети, креативний дизайн	Складність тренування

Сучасні тенденції розвитку CV у 3D-друку спрямовані на інтеграцію з іншими технологіями. По-перше, гібридні системи з AI та IoT які дозволяють проводити 24/7 виробництво з мінімальним наглядом. По-друге, використання синтетичних даних: Як показано в фармацевтичному дослідженні, генерація віртуальних зображень прискорює тренування моделей, зменшуючи витрати на реальні зразки на 90%. Тенденція до фотореалістичних симуляцій (Blender + Python) поширюється на інші галузі, роблячи QC стійким. Третя тенденція – мультисенсорна інтеграція: Комбінація CV з спектроскопією чи ультразвуком для повного 3D-сканування, як у системах для аерокосмічних деталей. Зростає роль edge computing для реального часу обробки на принтері, без хмарних серверів. Четверта – відкриті платформи: Проекти на GitHub, як open-source аналіз шарів, демократизують технологію для малого бізнесу. У 2025 році очікується зростання ринку CV для 3D-друку на 25%, з фокусом на зелені технології (зменшення відходів). Майбутнє цієї сфери лежить у глибокій інтеграції зі штучним інтелектом, переході до контролю в реальному часі та створенні комплексних цифрових двійників. Ці тенденції перетворюють 3D-друк на ще більш надійну та передбачувану технологію, повністю розкриваючи її потенціал для виробництва найскладніших та найвідповідальніших деталей у світі.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Актепе Е., Ергюн У. Підходи машинного навчання для 3D-друку на основі FDM: Огляд літератури // Прикладні науки (MDPI). – 2025. – Т. 15, № 18. – Ст. 10001. – DOI: 10.3390/app151810001. Режим доступу: <https://www.mdpi.com/2076-3417/15/18/10001>.
2. Шайхнаг А. та ін. Тенденції 3D-друку на 2025 рік: Виконавче опитування провідних компаній адитивного виробництва // 3D Printing Industry. – 2025. – Режим доступу: <https://3dprintingindustry.com/news/3d-printing-trends-for-2025-executive-survey-of-leading-additive-manufacturing-companies-236247>.
3. 3D-друк та III: Дослідження впливу машинного навчання на адитивне виробництво [Електронний ресурс]. – ResearchGate, 2025. – Режим доступу: <https://doi.org/10.63623/sfv96y88>.
4. Ройек І. та ін. Нові застосування машинного навчання в 3D-друку // Прикладні науки (MDPI). – 2025. – Т. 15, № 4. – Ст. 1781. – DOI: 10.3390/app15041781. Режим доступу: <https://www.mdpi.com/2076-3417/15/4/1781>.

Денис Віннічук

Аспірант кафедри комп'ютерних наук НУБіП України

0009-0001-9441-7696

Vinniboy2709@gmail.com

РОЗПІЗНАВАННЯ, КЛАСИФІКАЦІЯ ТА ВІДСТЕЖЕННЯ ОБ'ЄКТІВ В УМОВАХ ДИНАМІЧНОГО СПОСТЕРЕЖЕННЯ

Анотація. Було розглянуто принципи побудови системи спостереження для дронів, яка дозволяє розпізнавати, відстежувати та визначати положення об'єктів в реальному часі. Запропонована система враховує динаміку камери за допомогою інформації від датчиків висотоміру та дальноміру, а також значень кутів курсу, тангажу та крену.

Ключові слова: комп'ютерний зір, нейромережі, штучний інтелект, трекінг.

Задачі аналізу об'єктів на фото та відео в динамічних складних сценах сьогодні є надзвичайно актуальними в багатьох сферах, особливо в галузі оборони в період повномасштабного вторгнення. Завдання аналізу фото та відео умовно можна поділити на такі складові: Розпізнавання об'єкту; Класифікація з відомого набору класів; Відстеження, тобто присвоєння ідентифікатору між кадрами; Визначення фізичних характеристик для прийняття подальших рішень.

В ході роботи було розглянуто задачу де дрон з встановленою камерою веде спостереження, а саме оглядає плоску поверхню по якій рухаються деякі об'єкти, які необхідно розпізнати, класифікувати та визначити їх координати відносно камери. Важливо відмітити, що система повинна працювати в режимі реального часу. Дрон є відносно нерухомим (тобто виконує тільки задачу спостереження), однак камера на дроні може обертатися по трьом осям та збільшувати або зменшувати масштаб. Зображення передаються на сервер за допомогою RTMP потоку. Також дрон здатний передавати метадані, а саме значення кутів обертання камери та коефіцієнту зумування. Додатково дрон обладнаний дальноміром і висотоміром та передає дані від цих датчиків (рис. 1). Метадані є критично важливими при визначенні характеристик розпізнаних об'єктів.

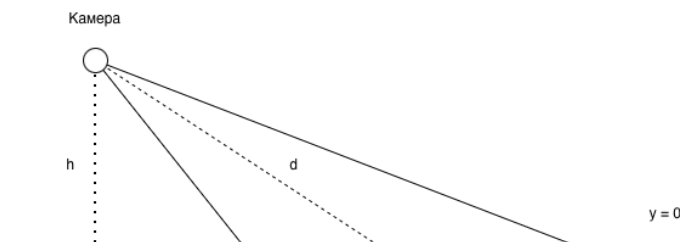


Рис. 1. Схема розташування поля зору камери. d – значення відстані від дальноміру, h – значення висоти

Завдання розпізнавання та класифікації об'єктів ефективно здатні виконувати сучасні нейромережеві підходи. Однією з найефективніших архітектур є одностадійна згортова нейронна мережа прямого поширення YOLO (You only look once) [1]. Така архітектура поступається в метриках точності двохстадійним підходам. Однак, головною перевагою YOLO є можливість розпізнавати та класифікувати об'єкти в режимі реального часу при відносно невеликих обчислювальних витратах. В задачі, що розглядається було використано YOLO v8m, яку було донавчено на різноманітних відкритих датасетах з військовою технікою.

При виявленні об'єкту часто потрібно також визначити його певні фізичні характеристики, такі як швидкість, траєкторію, розмір, тощо. Для цього потрібно

ідентифікувати координати розташування об'єкту. Це можливо зробити, якщо врахувати динамічну поведінку камери. Камера може змінювати кут курсу (yaw), тангажу (pitch) та крену (roll). Обертання здійснюється в порядку yaw-roll-pitch, що зумовлене особливістю конструкції шарнірів камери DJI Zenmuse H20. Також камера здатна виконувати зумування, тобто зміну масштабу камери. Додатково експериментально було визначено кути огляду камери в градусах (горизонтальний та вертикальний) FOV_{base} .

Нейромережа YOLO повертає локальні піксельні координати u, v обмежувальної рамки (які масштабуються в діапазон $[-1;1]$) та клас об'єкту на зображення. Щоб визначити глобальні координат спочатку потрібно знайти напрямляючий одиничний вектор від камери на об'єкт, який не буде змінюватись при обертанні камери чи зумуванні. Для цього отримане зображення проєктується на одиничну сферу навколо камери за допомогою відомих значень куту огляду і зуму. При чому впливу зуму на кут огляду визначається як:

$$\tan\left(\frac{FOV_{new}}{2}\right) = \frac{\tan\left(\frac{FOV_{base}}{2}\right)}{zoom_ratio} \quad (1)$$

Таким чином кожен об'єкт отримує координати $v_{local} = (x_{local}, y_{local}, z_{local})$ на одиничній сфері навколо камери. Наступним кроком є врахування обертання камери. Оскільки порядок обертання осей відомий, будується матриця обертання, яка є добутком матриць обертання по окремим осям в строго визначеному порядку, оскільки множення матриць є некомутативною операцією:

$$R = R_{yaw} \cdot R_{pitch} \cdot R_{roll}, \quad (2)$$

де R_{yaw} – матриця обертання навколо осі Y, R_{pitch} – матриця обертання навколо осі X, R_{roll} – матриця обертання навколо осі Z. Загальна матриця множиться на вектор координат:

$$v_{global} = R * v_{local} \quad (3)$$

Таким чином отримується одиничний напрямляючий на об'єкт вектор на одиничній сфері навколо камери. Для того щоб тепер отримати позицію об'єкту відносно камери потрібно врахувати дані від дальноміру та висотоміру. Знаючи висоту камери Y_{cam} та вектор напрямку v_{global} можна зайти точку перетину даного вектору з площиною землі:

$$\begin{pmatrix} X_{real} \\ Y_{real} \\ Z_{real} \end{pmatrix} = \begin{pmatrix} X_{cam} \\ Y_{cam} \\ Z_{cam} \end{pmatrix} + t * v_{global} \quad (4)$$

Оскільки $Y_{real} = 0$ (об'єкт розташований на землі), то $t = -\frac{Y_{cam}}{v_y}$.

$X_{cam} = 0$; $Y_{cam} = h$; $Z_{cam} = 0$, тобто top-view координати камери беремо за початок відліку. Недоліком такого методу є неможливість визначення координат у випадку знаходження розпізнаного об'єкту вище рівня землі ($Y_{real} > 0$).

Іншим підходом є використання відстані до об'єкту, оскільки це є значення від дальноміру (рис. 2):

$$\begin{pmatrix} X_{real} \\ Y_{real} \\ Z_{real} \end{pmatrix} = \begin{pmatrix} X_{cam} \\ Y_{cam} \\ Z_{cam} \end{pmatrix} + d * v_{global} \quad (5)$$

Однак такий підхід також має значний недолік, який полягає в тому, що дальномір працює точково, тобто повертає відстань від камери до об'єкту який розташований в центрі зображення. У випадку, якщо ж об'єкт знаходиться, наприклад на краю зображення то його реальна відстань може бути суттєво іншою в порівнянні з даними від датчику. Враховуючи вищезазначене для обчислення реальних координати було обрано метод з врахуванням висоти.

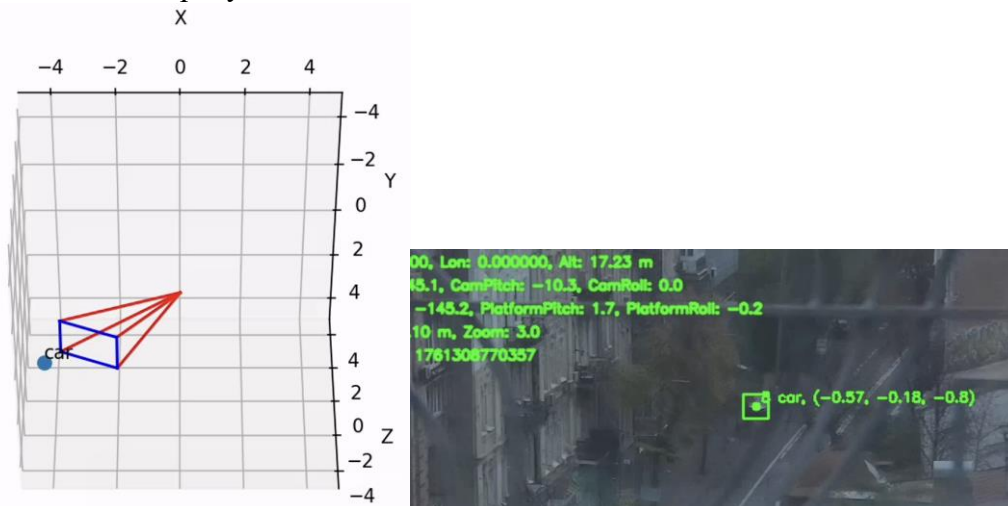


Рис. 2. Приклад роботи системи на дроні з побудовою 3D мапи

Трекінг об'єктів, тобто присвоєння ідентифікатору та відстеження руху між кадрами, дозволяє визначити швидкість та напрям руху об'єкту. Ефективним алгоритмом 2D трекінгу є ByteTrack [2], який дозволяє врахувати як висококонфідентні так і низькоконфідентні детекції. Однак, при використанні стандартного підходу використання піксельних координат, такий алгоритм не враховує динаміку камери. Тому було запропоновано використовувати для трекінгу координати об'єкту обраховані за допомогою вищепри описаного підходу з врахуванням зумування, обертання камери та даних від датчиків. Оскільки за моделлю, що розглядається завжди $Y_{real} = 0$, тому для трекінгу використано X_{real} та Z_{real} (розглядаємо об'єкти у форматі top – view, тобто дивимось згори).

В результаті було отримано систему, яка здатна розпізнавати та класифікувати об'єкти. Додатково за допомогою даних від датчиків вона здатна обраховувати координати відносно камери, а також здійснювати трекінг об'єктів незалежно від повороту камери та зумування, що дозволяє ефективно вести спостереження за полем бою в режимі реального часу. Систему протестовано на дроні DJI matrice з камерою DJI H20. В подальшому система може бути модифікована для роботи з рухомим дроном.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. J. Redmon, S. Divvala, R. Girshick and A. Farhadi, "You Only Look Once: Unified, Real-Time Object Detection," 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 2016, pp. 779-788, doi: 10.1109/CVPR.2016.91.
2. Zhang, Yifu & Sun, Peize & Jiang, Yi & Yu, Dongdong & Yuan, Zehuan & Luo, Ping & Liu, Wenyu & Wang, Xinggang. (2021). ByteTrack: Multi-Object Tracking by Associating Every Detection Box. URL: <https://doi.org/10.48550/arXiv.2110.06864>.

Олег Шубалий

Аспірант кафедри комп'ютерних наук
Тернопільський національний технічний університет імені Івана Пулюя,
м. Тернопіль, Україна
ORCID 0009-0003-2345-9148
oleh_shubalyi1606@tntu.edu.ua

Леся Дмитроца

К.т.н, доцент, доцент кафедри комп'ютерних наук
Тернопільський національний технічний університет імені Івана Пулюя,
кафедра комп'ютерних наук, м. Тернопіль, Україна
ORCID 0000-0003-2583-3271
dmytrotso.lesya@gmail.com

МЕТОДОЛОГІЧНИЙ АНАЛІЗ ПІДХОДІВ В АНАЛІТИЦІ BIG DATA

Анотація. Здійснено порівняльний аналіз методів провадження аналітики великих даних. Запропоновано створення комплексної методології аналітики Big Data із застосуванням кращих підходів та їх поєднанням.

Ключові слова: Big Data; аналітика; методологія.

1. ВСТУП

Проблематика великих даних включає у себе не лише питання інфраструктури та технологій розподіленої обробки, але й, очевидно, як найважливіший аспект, їх ефективний аналітичний аналіз. По при сучасні тенденції всебілшого поширення методів машинного навчання, класичні статистичні, емпіричні та сигнатурні методи можуть мати перевагу над ними в аналітиці Big Data в залежності від контексту задач та вимог.

Постановка проблеми. Систематизація підходів та визначення кращих шляхів провадження аналітики великих даних є актуальною задачею як для наукових досліджень, так і для практичного застосування.

Аналіз останніх досліджень і публікацій. Одним з найбільш відомих дослідників, який звертає увагу на непрозорість машинного навчання на противагу статистичній обґрунтованості у Big Data є професор Bin Yu з університету Берклі. Зокрема в [1] ним запропоновано концепцію PCS: Predictability, Computability, Stability як основу для надійного аналізу даних. Важливим та актуальним для розуміння класичних методів аналітики є джерело [2]. Детальний виклад аналітики великих даних подано у [3]. Актуальна аналітика Big Data засобами машинного навчання розкрита в [4].

Мета публікації. Метою публікації є розробка нових підходів до формування комплексної методології аналітики великих даних, яка поєднує класичні методи та методи штучного інтелекту.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

В реальних сферах застосувань до аналітики Big Data ставляться вимоги прозорості та контролю. Важливу роль відіграє й безпека чутливих даних. Тому застосування методів машинного навчання для аналітики при таких обмеженнях може бути недопустиме. Окрім того, класичні статистичні, евристичні чи сигнатурні підходи більш ефективніші в задачах потокової обробки реального часу, оскільки споживають значно менше обчислювальних ресурсів ніж технології штучного інтелекту. Через непрозорість алгоритмів, особливо у глибоких моделях, інтерпритованість результатів машинного навчання може бути низькою і, відповідно, неприйнятною для аудиту. Крім

того, методи штучного інтелекту мають критичну залежність від якості навчальних даних, яка також впливає і стабільність результатів.

В свою чергу, класичні методи аналітики поступаються машинному навчанню у гнучкості і є менш ефективними при високій варіативності даних у Big Data та складних нелінійних залежностях. Нейромережі здатні самостійно виявляти закономірності в даних без попереднього знання їх структури, в той час як для сигнатурних та евристичних методів необхідне глибоке розуміння домену.

Більш наочний порівняльний аналіз класичних методів аналітики Big Data із методами машинного навчання наведено у таблиці 1.

Таблиця 1

Порівняння методів аналітики Big Data

Критерій порівняння	Класичні методи	Методи машинного навчання
Прозорість та контрольованість	Висока	Низька
Споживання ресурсів	Помірне	Високе
Робота в режимі реального часу	Висока придатність	Менша придатність
Гнучкість та адаптивність	Обмежена	Висока
Потреба у навчальних даних	Відсутня	Наявна
Стабільність	Висока	Визначається якістю навчальних даних
Обробка даних складної структури	Обмежена	Ефективна при якісних навчальних даних

Хоча кожен із описаних методологічних підходів аналітики великих даних має свої недоліки і вибір для застосувань у практичній сфері визначається накладеними вимогами, однак перспективним є розвиток гібридних методологій, що використовуватимуть переваги кожного із підходів.

При цьому можливі варіанти як послідовного застосування класичних методів та методів штучного інтелекту у різних комбінаціях, так і реалізації паралельних підходів. Ще одним можливим варіантом гібридних методологій є використання класичних методів для пояснення результатів, отриманих технологіями штучного інтелекту. До прикладу, для пояснення рішень, отриманих ШІ, вже на основі статистичних або евристичних оцінок важливості ознак.

У випадку каскадної (послідовної) інтеграції можливі варіанти, коли одні з методів використовуватимуться для попередньої обробки великих даних, а подальший аналіз здійснюватиметься іншими методами. Так, якщо спочатку використати для очистки та підготовки Big Data класичні методи (зокрема, статистичні – для нормалізації та зменшення розмірності, а евристичні – для фільтрації за заданими правилами та пороговими значеннями), то завдяки покращеній якості даних підвищиться точність та машинного навчання і, відповідно якість результатів ШІ. Якщо ж навпаки для попередньої обробки Big Data вибрати нейронні мережі, то їх можна застосувати щоб виділити ознаки у високо варіативних даних із складними залежностями, які в подальшому можна буде проаналізувати класичними статистичними методами.

Різні варіанти паралельного (ансамблевого) поєднання класичних методів та методів ШІ можуть, до прикладу, застосовуватися для гібридного виявлення аномалій

(коли сигнатурними методами виявляють відомі типи збоїв, на яких навчатиметься модель), або у стекінгу (коли класичними статистичними методами та методами машинного навчання генеруються окремі прогнози, які потім об'єднуються моделлю III, що виступає мета-класифікатором).

Основні переваги кожного з наведених варіантів гібридної методології згруповано у таблиці 2.

Таблиця 2

Основні переваги гібридної методології в аналітиці Big Data

Види гібридної методології	Деталі методології	Переваги
Послідовна (каскадна)	Попередня обробка Big Data класичними методами, машинне навчання на якісних даних	Підвищення точності та швидкості ML
	Виділення складних варіативних ознак із сирих даних, фінальний аналіз класичними методами	Використання потужностей III для аналізу класичними методами
Паралельна (ансамблева)	Стекінг шляхом здійснення базового прогнозу класичними та машинними методами, фінальний прогноз мета-класифікатором III	Максимальна надійність прогнозних результатів
	Гібридне виявлення аномалій сигнатурними методами та методами III	Швидкість та гнучкість у виявленні загроз та критичних відхилень
Інтерпретаційна післяобробка	Кількісна оцінка важливості ознак та факторів класичними методами за результатами III-моделей	Збільшення довіри до результатів III та їх прийнятність для аудиту

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Виконано методологічний аналіз підходів в аналітиці Big Data. Відображено переваги та недоліки кожного з них. Запропоновано створення комплексної методології аналітики великих даних на основі гібридних підходів, які будуть розвиватися у наступних дослідженнях.

ПОСИЛАННЯ

1. B. Yu and R. L. Barter, *Veridical Data Science: The Practice of Responsible Data Analysis and Decision Making*. Cambridge, MA: The MIT Press, 2024.
2. G. James, D. Witten, T. Hastie, and R. Tibshirani, *An Introduction to Statistical Learning: with Applications in R*, 2nd ed. Stanford University, 2023. [Online]. Available: https://hastie.su.domains/ISLR2/ISLRv2_corrected_June_2023.pdf.download.html. [Accessed: Nov. 9, 2025].
3. J. Leskovec, A. Rajaraman, and J. D. Ullman, *Mining of Massive Datasets*, 3rd ed. Cambridge: Cambridge University Press, 2020.
4. K. Banachewicz and L. Massaron, *The Kaggle Book: Data Analysis and Machine Learning for Competitive Data Science*. Birmingham: Packt Publishing, 2022.

Ганна Вайганг

кандидат технічних наук, доцент

Національний університет біоресурсів і природокористування України, Київ, Україна

0000-0002-2082-2322

weigang.ganna@nubip.edu.ua

Юрій Науринський

Аспірант

Національний університет біоресурсів і природокористування України, Київ, Україна

0009-0004-6416-8635

yu.naurynskyi@nubip.edu.ua

ТЕОРЕТИЧНА МОДЕЛЬ МЕНЕДЖМЕНТУ ВИБОРУ АЛГОРИТМІВ КЛАСТЕРИЗАЦІЇ ЗА ОПИСОВИМИ ХАРАКТЕРИСТИКАМИ НАБОРУ ДАНИХ ДЛЯ СИСТЕМ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ

Анотація. Розглянуто теоретичні засади побудови менеджментної моделі вибору алгоритмів кластеризації за описовими характеристиками набору даних для систем підтримки прийняття рішень. Запропоновано концептуальну схему, що забезпечує формалізований зв'язок між властивостями даних та алгоритмічними методами їх аналізу. Розроблена модель спрямована на підвищення адаптивності й пояснюваності аналітичних процедур, а також на покращення точності розпізнавання структур у даних і зменшення кількості помилкових об'єднань при кластеризації. Це створює підґрунтя для розвитку самонавчальних систем, здатних автоматично обирати оптимальні методи обробки даних залежно від їхньої структури та контексту застосування.

Ключові слова: системи підтримки прийняття рішень; алгоритм кластеризації; великі обсяги даних

1. ВСТУП

Сучасні системи підтримки прийняття рішень (СППР) працюють у середовищі великих, динамічних та неоднорідних даних. Їх ефективність значною мірою залежить від правильного вибору алгоритмів обробки та аналізу інформації, що визначають швидкість, точність і надійність отриманих результатів. Одним із ключових методів аналітичного етапу є кластеризація, яка дає змогу узагальнювати дані, виявляти приховані структури й формувати основи для подальших управлінських рішень.

Проблема полягає в тому, що різні алгоритми кластеризації демонструють неоднакову ефективність залежно від характеристик даних – їх розмірності, щільності, рівня шуму, топології та метрики подібності. Традиційно вибір алгоритму здійснюється вручну, що не відповідає вимогам сучасних СППР до автономності, адаптивності та масштабованості. Потребує розроблення теоретична модель, яка б дозволила формалізувати процес вибору алгоритмів кластеризації на основі описових характеристик набору даних і забезпечити адаптивну поведінку аналітичних підсистем.

Мета публікації. Теоретичне обґрунтування моделі менеджменту вибору алгоритмів кластеризації за описовими характеристиками набору даних як інтелектуального компонента систем підтримки прийняття рішень, що забезпечує формальну відповідність між властивостями даних і методами їх аналітичної обробки, застосування.

2. ТЕОРЕТИЧНІ ОСНОВИ

Теоретичні основи побудови менеджментної моделі ґрунтуються на уявленні про структуру даних як на сукупність параметрів, що можуть бути формалізовані та використані для прийняття рішень щодо вибору методів їх аналізу. Для того щоб алгоритм кластеризації був ефективним, він має відповідати властивостям конкретного набору даних, а отже, необхідно встановити зв'язок між характеристиками даних та

алгоритмічними підходами, здатними забезпечити стабільні результати. У цьому контексті важливим є опис простору дескрипторів, який дозволяє узагальнювати множину значущих параметрів і переводити задачу підбору алгоритму у формалізований процес зі зрозумілими критеріями відповідності.

Набір даних D характеризується множиною параметрів:

$$\delta(D) = \{d_1, d_2, \dots, d_n\} \quad (1)$$

де d_i відображають статистичні, структурні та інформаційні властивості даних: кількість об'єктів, розмірність, середню щільність, рівень шуму, кореляційні зв'язки, форму розподілу тощо [1].

Ці параметри утворюють простір дескрипторів Δ , який узагальнює властивості наборів даних і дає можливість здійснювати метааналіз для вибору алгоритмів кластеризації.

Алгоритмічний простір визначимо як:

$$A = \{a_1, a_2, \dots, a_m\} \quad (2)$$

де кожен алгоритм a_j має свої припущення, обчислювальну складність і набір параметрів $\theta(a_j)$. Вибір алгоритму залежить від того, наскільки його властивості узгоджуються з дескрипторами даних. Зокрема, алгоритми метричного типу (наприклад, k-means) є ефективними для компактних і майже сферичних кластерів, тоді як щільнісні (DBSCAN, OPTICS) – для розріджених і шумових структур [2].

Функціональне відображення між простором дескрипторів і простором алгоритмів має вигляд:

$$f : \Delta \rightarrow A \quad (3)$$

де f визначає механізм вибору алгоритму кластеризації для кожного класу описових характеристик даних [3].

Задачу вибору можна подати як мінімізацію функції втрат:

$$L(f, D) = \ell(\text{результат кластеризації на } D \mid f(\delta(D))) \quad (4)$$

де $\ell(\cdot)$ описує відхилення між структурою знайдених кластерів і бажаними властивостями аналітичного результату в межах СППР [4].

Такий підхід переводить задачу підбору алгоритму у метарівневий простір оптимізації, де об'єктом є не параметри кластеризації, а сама процедура вибору методу відповідно до властивостей даних.

Метарівнева модель функціонує як автономний аналітичний модуль, що передусь безпосередній кластеризації. Її робота складається з трьох етапів:

1. Аналіз дескрипторів набору даних та побудова вектору характеристик.
2. Визначення класу даних і відображення його у простір алгоритмів кластеризації за допомогою функції f .
3. Формування рекомендацій або автоматичний вибір алгоритму з відповідними гіперпараметрами [5].

Таким чином забезпечується самоналаштуваність аналітичного рівня СППР, скорочується час прийняття рішень і підвищується стабільність результатів у динамічних середовищах. Метарівнева модель також створює основу для подальшого розвитку адаптивних систем, у яких відбір алгоритмів стає частиною процесу навчання системи. На практиці це означає, що система здатна накопичувати досвід попередніх

запусків: якщо певний алгоритм демонструє стабільно кращі показники для схожих наборів, він буде рекомендований автоматично.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Дослідження базується на метарівневому аналітичному підході, який поєднує методи системного аналізу, узагальнення наукових джерел та теоретичного моделювання. На першому етапі було проведено огляд публікацій і порівняльних досліджень, присвячених алгоритмам кластеризації, AutoML-підходам і мета-навчанню [1]–[3]. Це дозволило виявити закономірності між властивостями даних і результативністю окремих алгоритмів, що стало основою для побудови метарівневої моделі вибору.

На другому етапі сформульовано концепцію тривірневої моделі, яка описує взаємозв'язок між характеристиками набору даних, множиною алгоритмів кластеризації та процесом прийняття рішення у СППР [4], [5]. Такий підхід забезпечує теоретичне підґрунтя для подальшого розвитку адаптивних і самонавчальних систем підтримки прийняття рішень.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

У результаті дослідження було сформовано концептуальну основу метарівневої моделі вибору алгоритмів кластеризації для систем підтримки прийняття рішень. Запропонована модель розглядається як проміжна ланка між рівнем даних та аналітичним рівнем системи, що забезпечує автоматизований вибір методу кластеризації залежно від характеристик набору даних. Вона передбачає наявність функціонального зв'язку між описовими параметрами даних, такими як кількість об'єктів, розмірність, щільність, рівень шуму, структура кореляцій, та класом алгоритмів, який найкраще відповідає цим властивостям. Це дає змогу системі здійснювати вибір алгоритму без участі користувача, що зменшує вплив людського фактора та підвищує адаптивність і відтворюваність аналітичних процесів.

Наприклад, у наборах даних з високою розрідженістю (транспортні потоки, журнали подій) модель обирає HDBSCAN, тоді як для компактних і збалансованих ознак (медичні таблиці або результати соціальних опитувань) доцільними є GMM або k-means.

Модель інтегрується в архітектуру систем підтримки прийняття рішень як надбудова аналітичного модуля, яка виконує попередню обробку вхідних даних, аналізує їх дескриптори та визначає відповідний алгоритм кластеризації для подальшої обробки. Такий підхід дозволяє зменшити кількість обчислювальних експериментів, пов'язаних із пошуком оптимального методу, а також скорочує час отримання результатів, що є критично важливим у системах реального часу. Крім того, модель забезпечує структурованість і логічну прозорість процесу вибору, що сприяє підвищенню пояснюваності аналітичних рішень у межах СППР.

Проведений аналіз показує, що метарівнева модель створює основу для побудови самонавчальних аналітичних систем, у яких результати попередніх обчислень можуть бути використані як база знань для подальших виборів. Це відкриває можливість накопичення метаінформації про ефективність алгоритмів у різних сценаріях, що у перспективі дозволить формувати узагальнену базу рекомендацій. Такі системи можуть еволюціонувати з часом, адаптуючи власну поведінку до змін у структурі даних та зовнішніх умовах.

Перспективи подальшого розвитку дослідження полягають у формуванні багатокритеріальної моделі, яка враховуватиме не лише структурні властивості даних, але й показники продуктивності, такі як час виконання, споживання ресурсів, стійкість

і точність. Це дозволить реалізувати комплексну систему вибору, у якій кожен алгоритм оцінюється за вектором критеріїв, що забезпечить оптимізацію не лише якості кластеризації, а й ефективності всієї системи підтримки прийняття рішень. У подальшому передбачається розроблення адаптивної метамоделі навчання, здатної оновлювати правила вибору алгоритмів на основі накопичених знань, що наблизить СППР до концепції повністю інтелектуальної, самонавчальної системи аналітичного прийняття рішень.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У ході дослідження розроблено теоретичну модель метарівневого вибору алгоритмів кластеризації за описовими характеристиками набору даних, яка розглядається як ключовий компонент адаптивних систем підтримки прийняття рішень. Запропонований підхід дозволяє формалізувати процес вибору алгоритму, зменшити залежність від експертних оцінок та забезпечити узгодженість між властивостями даних і методами їх аналізу. Модель зменшує кількість ручних налаштувань при виборі алгоритму і скорочує час отримання коректного результату кластеризації за рахунок автоматизованого аналізу властивостей даних. Перспективами подальших досліджень є розроблення практичних методів реалізації метарівневої моделі у вигляді програмних компонентів, інтегрованих із існуючими системами аналізу даних, а також розширення її можливостей шляхом включення критеріїв продуктивності, масштабованості та самооновлення на основі накопиченого досвіду роботи системи.

ПОСИЛАННЯ

- [1] Xu, D., & Tian, Y. (2015). A comprehensive survey of clustering algorithms. *Annals of Data Science*, 2(2), 165–193. <https://doi.org/10.1007/s40745-015-0040-1>
- [2] Garouani, M., Ahmad, A., Bouneffa, M., Hamlich, M., Bourguin, G., & Lewandowski, A. (2022). Using meta-learning for automated algorithms selection and configuration: an experimental framework for industrial big data. *Journal of Big Data*, 9(1). <https://doi.org/10.1186/s40537-022-00612-4>
- [3] Poulakis, Y., Doulkeridis, C., & Kyriazis, D. (2024). A survey on AutoML Methods and Systems for Clustering. *ACM Transactions on Knowledge Discovery From Data*, 18(5), 1–30. <https://doi.org/10.1145/3643564>
- [4] Barbudo, R., Ventura, S., & Romero, J. R. (2023). Eight years of AutoML: categorisation, review and trends. *Knowledge and Information Systems*, 65(12), 5097–5149. <https://doi.org/10.1007/s10115-023-01935-1>
- [5] Rivolli, A., Garcia, L. P., Soares, C., Vanschoren, J., & De Carvalho, A. C. (2022). Meta-features for meta-learning. *Knowledge-Based Systems*, 240, 108101. <https://doi.org/10.1016/j.knosys.2021.108101>

Марина Лендел

аспірантка 3-го року навчання НУБіП України

Місце роботи: НУБіП України, факультет інформаційних технологій,

кафедра комп'ютерних наук, місто Київ, Україна

ORCID: 0009-0008-0042-7705

marynalendel@gmail.com

Белла Голуб

к.т.н., доцент

Місце роботи: НУБіП України, факультет інформаційних технологій,

кафедра комп'ютерних наук, місто Київ, Україна

ORCID: 0000-0002-1256-6138

bellalg@nubip.edu.ua

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ДЛЯ СИСТЕМИ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ПРИ ВИРОЩУВАННІ ЗЕРНОВИХ КУЛЬТУР

Анотація. У роботі розглядається підхід до побудови системи підтримки прийняття рішень у процесі вирощування зернових культур на основі використання технологій OLAP та Data Mining. Основною метою дослідження є забезпечення багатовимірного аналізу взаємозв'язків між метеорологічними показниками та розвитком шкідників для прогнозування врожайності. Запропонована структура сховища даних, яка дозволяє інтегрувати великі обсяги даних, проводити аналіз у розрізі часу, шкідників, локації та створює основу для подальшого застосування методів інтелектуального аналізу даних.

Ключові слова: система підтримки прийняття рішень, сховище даних, OLAP, Data Mining, багатовимірний аналіз, зернові культури, шкідники.

ВСТУП

Мета дослідження. Розроблення підходів до побудови системи підтримки прийняття рішень у процесі вирощування зернових культур на основі використання технологій OLAP та Data Mining для виявлення закономірностей між метеорологічними та біологічними показниками.

Об'єкт дослідження. Процес вирощування зернових культур.

Предметом дослідження виступають методи зберігання, оброблення та інтелектуального аналізу даних, що впливають на урожайність зернових культур.

Актуальність теми. У сучасному аграрному виробництві одним із ключових напрямів підвищення ефективності є використання інформаційних технологій для підтримки прийняття рішень. В умовах змін клімату та підвищеної мінливості погодних умов агрономи стикаються з необхідністю швидко і точно оцінювати ризики, прогнозувати розвиток шкідників та планувати агротехнічні заходи. Для цього застосовуються сучасні методи обробки та аналізу даних, які дозволяють перетворювати великі обсяги інформації у корисні знання. У даному дослідженні використовується багаторічний набір даних, що включає показники тривалості сонячного сяйва, середньорічної температури, вологості повітря, суми опадів, чисельності озимої совки, плодючості її самиць та дані про розвиток гусениць. Такий комплекс показників дає можливість комплексно оцінювати взаємозв'язок кліматичних і біологічних факторів із розвитком шкідників і потенційним впливом на урожай зернових культур.

МЕТОДИ ДОСЛІДЖЕННЯ

Першим етапом роботи стала підготовка даних: перевірка коректності значень, усунення пропусків та аномалій, видалення дублікатів. На наступному етапі здійснено структурування даних у вигляді сховища, що відображає основні агрометеорологічні та

біологічні показники, які впливають на врожайність зернових культур. До них належать такі параметри, як тривалість сонячного сяйва, середньорічна температура повітря, сума опадів, відносна вологість повітря, а також показники розвитку шкідників — зокрема, чисельність озимої совки, плодючість самиць, тривалість розвитку гусениць та маса лялечки. Розроблене сховище даних забезпечує багатовимірний аналіз. Його концепція побудована за принципом часової орієнтації, де кожен запис містить рік спостереження, що дозволяє відстежувати динаміку змін показників у часі.

Розробка OLAP-кубу дозволить у подальшому проводити аналіз даних, будувати зрізи, виявляти закономірності між погодними умовами, біологічними факторами та врожайністю. Це створює основу для подальшого застосування методів Data Mining для прогнозування врожайності та оцінки ризиків поширення шкідників.

РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Побудова сховища даних здійснювалася за принципом моделі “зірки” (Star Schema), що є класичним підходом у проєктуванні аналітичних систем. Центральною частиною моделі виступає таблиця фактів FactAgroClimate, у якій зберігаються вимірювані показники за роками та регіонами — середньорічна температура, тривалість сонячного сяйва, сума опадів, відносна вологість, а також біологічні характеристики озимої совки: щільність популяції, плодючість самиць, тривалість розвитку гусениць і маса лялечки. Також додатково передбачено поле для зберігання урожайності (harvest_tons), що дозволяє встановити зв’язок між природними факторами та кінцевим результатом врожайності.

У таблицях вимірів зберігаються:

- DateDim — описує часові параметри (рік, місяць, день);
- LocationDim — містить інформацію про поля, де проводяться дослідження;
- PestDim — описує характеристики шкідника, зокрема його вид, тип і біологічні особливості.



Рис. 1. Структура сховища даних для розроблюваної СППР

Вибір саме такої структури пояснюється необхідністю відокремлення описових атрибутів від змінних кількісних показників, що дозволяє зменшити надмірність даних і забезпечити ефективність запитів у процесі побудови OLAP-кубів. Це створює передумови для подальшого етапу аналізу, який передбачає формування

багатовимірної моделі даних з можливістю агрегування, порівняння та візуалізації результатів у розрізі часу, локації, шкідника.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

На даному етапі було створено структуру даних, що забезпечує базу для побудови системи підтримки прийняття рішень. Підготовка OLAP-кубу є ключовим кроком у формуванні аналітичної моделі, яка дозволяє проводити багатовимірний аналіз і виявляти взаємозв'язки між природними та біологічними чинниками.

Подальші дослідження передбачають використання методів Data Mining для прогнозування врожайності, оцінки ризиків поширення шкідників та оптимізації агротехнічних заходів. Побудова сховища даних і первинний аналіз не лише дозволили впорядкувати інформацію, але й стали основою для глибшого моделювання взаємозв'язків між метеорологічними умовами та розвитком шкідників, що є важливим елементом системи підтримки прийняття рішень у сфері вирощування зернових культур.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

1. The Concept of Data Mining [Електронний ресурс] // IntechOpen. – Режим доступу: <https://www.intechopen.com/chapters/78106> (дата звернення: 01.11.2025).
2. Overview of Online Analytical Processing (OLAP) [Електронний ресурс] //Microsoft Support. – Режим доступу: <https://support.microsoft.com/en-us/office/overview-of-online-analytical-processing-olap-15d2cddde-f70b-4277-b009-ed732b75fdd6> (дата звернення: 01.11.2025).
3. Condran, Sarah & Bewong, Michael & Islam, Md Zahidul & Maphosa, Lance &Zheng, Lihong. (2022). Machine Learning in Precision Agriculture: A Survey on Trends, Applications and Evaluations Over Two Decades. IEEE Access. 10. 1-1.10.1109/ACCESS.2022.3188649.
4. Dolia, M., Lysenko, V., Lendiel, T., Nakonechna, K., & Humeniuk, L. (2024). Neuron network prediction of damage of E. integriceps bug on winter wheat in Ukraine. Scientific Reports of the National University of Life and Environmental Sciences of Ukraine, 4(20), 96-105.

Ганна Вайганг

кандидат технічних наук, доцент

Національний університет біоресурсів і природокористування України, Київ, Україна

0000-0002-2082-2322

weigang.ganna@nubip.edu.ua

Іван Корнілов

Асистент

Національний університет біоресурсів і природокористування України, Київ, Україна

0009-0009-5598-2690

i.kornilov@nubip.edu.ua

ГІБРИДНА МОДЕЛЬ НЕЙРОННОГО ПОШУКУ: ПОЄДНАННЯ СЛОВНИКОВОГО ТА СЕМАНТИЧНОГО ПІДХОДІВ

Анотація. Публікація присвячена аналізу гібридної моделі нейронного пошуку, що поєднує лексичні та семантичні підходи до індексування й ранжування текстових документів. Обґрунтовується необхідність інтеграції класичних методів, таких як BM25, із сучасними нейронними моделями представлення тексту для підвищення повноти, стійкості та інтерпретованості результатів пошуку. Розглядаються типові архітектури гібридних систем, схеми злиття оцінок релевантності, каскадні та адаптивні підходи. Узагальнено результати досліджень, що демонструють переваги гібридизації над виключно лексичними або виключно семантичними моделями. Показано перспективи застосування гібридного пошуку в інформаційно-пошукових системах нового покоління, системах з доповненим відбором контенту (RAG) та доменно-специфічних сховищах знань.

Ключові слова: гібридний пошук; лексична модель; семантична модель; векторні подання; оцінка релевантності; злиття сигналів; доповнений пошук; BM25; RAG.

1. ВСТУП

Сучасні інформаційно-пошукові системи функціонують в умовах стрімкого зростання обсягів даних, варіативності форматів та підвищених вимог до точності, повноти та надійності результатів. Класичні словникові методи (TF-IDF, BM25) залишаються базовим стандартом завдяки ефективності, масштабованості та прозорості механізмів ранжування, однак вони обмежені чутливістю до синонімів, перефразувань та контекстних значень термів.

Нейронні семантичні моделі на основі трансформерів забезпечують контекстно-залежне представлення тексту та дозволяють враховувати латентні зв'язки між запитом і документами. Проте вони вимагають значних обчислювальних ресурсів, є чутливими до доменного зсуву та потребують спеціалізованого навчання і моніторингу якості [1,3]. У результаті актуальним стає перехід до гібридних архітектур, які поєднують переваги обох підходів.

Мета публікації. Узагальнити концепцію гібридної моделі нейронного пошуку, проаналізувати основні архітектурні рішення та методи злиття лексичних і семантичних сигналів, а також окреслити переваги, обмеження та перспективи практичного впровадження таких систем.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Гібридна модель нейронного пошуку розглядається як композиція лексичного та семантичного модулів, які паралельно або каскадно генерують кандидати документів та спільно визначають підсумковий рейтинг. Лексичний модуль, як правило, реалізується на основі BM25 та інвертованих індексів, тоді як семантичний модуль використовує щільні або пізньо-інтерактивні подання (bi-encoder, ColBERTv2 та подібні) [1,3].

Типова архітектура включає: (1) попередню обробку та індексування колекції; (2) обчислення лексичних оцінок релевантності для великого підмножинного списку документів; (3) семантичний пошук або повторне ранжування для кандидатів; (4) модуль злиття результатів, що враховує обидва джерела сигналів. Такий підхід дозволяє поєднати високу точність роботи з ключовими термінами з можливістю виявлення релевантних документів, які не містять точних збігів із запитом (рис. 1).

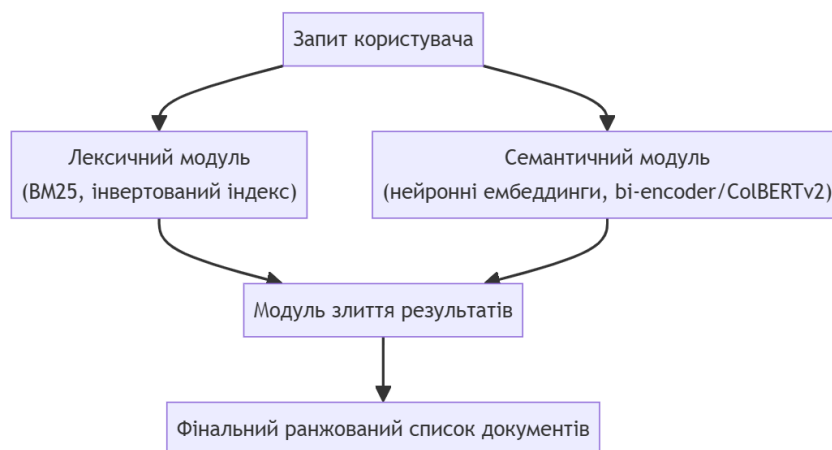


Рис. 1. Узагальнена архітектура гібридної моделі нейронного пошуку

Модуль злиття (fusion) є ключовим елементом гібридної моделі. Найпростішим підходом є лінійна комбінація нормалізованих оцінок релевантності лексичної та семантичної моделей з фіксованою вагою (рис. 2). Більш стійкі результати забезпечують рангові методи, зокрема Reciprocal Rank Fusion, що зменшує вплив різниці масштабів оцінок. У сучасних роботах також розглядаються адаптивні схеми, де внесок кожного модуля залежить від типу запиту: довжина, наявність рідкісних термів, доменність [2,4].

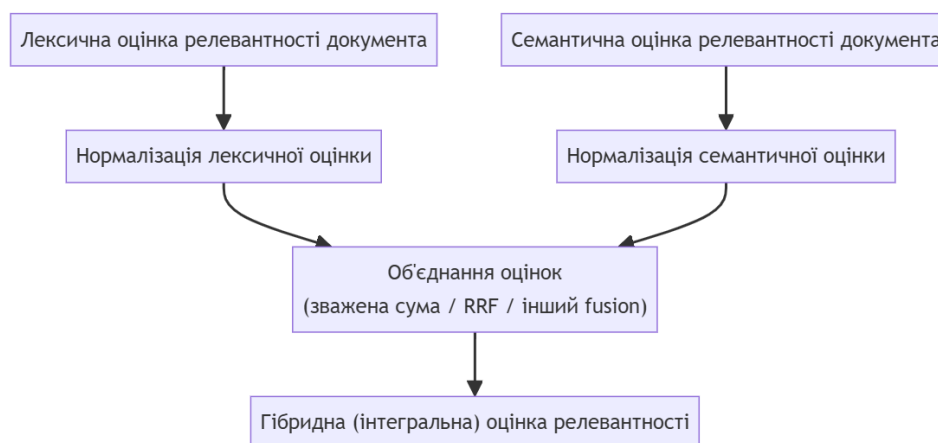


Рис. 2. Схема обчислення гібридної оцінки релевантності

Для узагальнення переваг гібридної моделі порівняємо основні характеристики лексичного, семантичного та гібридного підходів (табл. 1).

Таблиця 1

Порівняння характеристик лексичного, семантичного та гібридного пошуку

Характеристика	Лексичний пошук	Семантичний пошук	Гібридний пошук
Обробка синонімів і перефразувань	Обмежена	Висока	Висока, компенсується семантичним модулем
Точність для точних термів, кодів, ідентифікаторів	Висока	Середня / залежить від навчання	Зберігається завдяки лексичному компоненту
Повнота (Recall)	Середня	Висока	Вища за окремі компоненти
Обчислювальні витрати	Низькі	Високі	Середні (залежить від каскадної схеми)
Інтерпретованість	Висока	Обмежена	Частково зберігається за рахунок лексичних сигналів

Емпіричні дослідження демонструють, що гібридні стратегії забезпечують стабільне підвищення показників якості (nDCG, MAP, Recall) порівняно з окремими моделями, особливо для складних запитів та доменних колекцій [1–3]. Узагальнені ефекти наведені в табл. 2.

Таблиця 2

Узагальнені ефекти застосування гібридних моделей пошуку

Аспект	Типовий ефект гібридних моделей
Recall@k	Зростання на декілька відсоткових пунктів відносно кращого з компонентів
NDCG@10 / MAP	Покращення ранжування для довгих, семантично насичених і природномовних запитів
Стійкість до шуму та доменного зсуву	Зменшення ризику пропуску релевантних документів та покращення робастності завдяки поєднанню сигналів
Вартість запиту	Оптимізується каскадними схемами та адаптивним використанням нейронних компонентів

Попри переваги, гібридні моделі ставлять низку інженерних та дослідницьких викликів. По-перше, складною є нормалізація оцінок та вибір параметрів злиття, оскільки розподіли лексичних і семантичних оцінок різні. По-друге, підтримка векторних індексів, оновлення векторних подань (embedding) та забезпечення прийнятної латентності вимагають оптимізованої інфраструктури, ефективних алгоритмів пошуку найближчих сусідів і каскадних стратегій [1,4].

Додатковим аспектом є інтерпретованість. У критичних застосуваннях (медицина, право) необхідне пояснення рішень системи. Тому рекомендується зберігати лексичні пояснення (ключові терміни, збіги, ваги) та метадані, а семантичні моделі використовувати як підсилюючий компонент. Практичні рекомендації включають поступове впровадження гібридних схем на базі вже існуючих лексичних індексів, A/B-тестування та постійний моніторинг якості.

3. ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Гібридна модель нейронного пошуку, що поєднує словникові та семантичні підходи, формує сучасний стандарт побудови ефективних інформаційно-пошукових систем. Вона забезпечує підвищення повноти, збереження чутливості до точних визначень, кращу стійкість до різномірних запитів і можливість поєднання інтерпретованих та латентних сигналів релевантності. Представлені в літературі

результати підтверджують переваги гібридних архітектур над моно-підходами як для відкритих, так і для спеціалізованих колекцій [1–3].

Перспективними напрямками подальших досліджень є: адаптивне та параметризоване злиття лексичних і семантичних оцінок релевантності; побудова єдиних гібридних представлень, що поєднують розріджені та щільні ознаки; оптимізація обчислювальних витрат за рахунок ефективних ANN-структур та підходів пізньої взаємодії запиту з документом; інтеграція багатомодальних даних; використання гібридних стратегій у системах з доповненим відбором контенту (RAG) для зменшення ризику "галюцинацій" генеративних моделей.

ПОСИЛАННЯ

- [1] Lin S.-C., Lin J. A Dense Representation Framework for Lexical and Semantic Matching. *ACM Transactions on Information Systems*, 41(4), 2023. DOI: 10.1145/3582426.
- [2] Kuzi S., Zhang M., Li C., Bendersky M., Najork M. Leveraging Semantic and Lexical Matching to Improve the Recall of Document Retrieval Systems: A Hybrid Approach. <https://arxiv.org/abs/2010.01195>.
- [3] Santhanam K., Khattab O., Potts C., Zaharia M. ColBERTv2: Effective and Efficient Retrieval via Lightweight Late Interaction. *Proceedings of NAACL*, 2022. DOI: 10.18653/v1/2022.naacl-main.272
- [4] Kulkarni H., MacAvaney S., Goharian N., Frieder O. Lexically-Accelerated Dense Retrieval. *Proceedings of SIGIR*, 2023. DOI: 10.1145/3539618.3591715

Володимир Хиленко

Д.т.н., професор, професор кафедри комп'ютерних наук
Національний університет біоресурсів і природокористування України, Київ, Україна
ORCID ID 0000-0003-3491-8621
vkhilenko@nubip.edu.ua

Олексій Степанов

Старший викладач кафедри комп'ютерних наук
Національний університет біоресурсів і природокористування України, Київ, Україна
ORCID ID 0000-0002-0939-6991
stepanov@nubip.edu.ua

Bence Ladóczki

Assistant Research Fellow in the Department of Telecommunications and Media Informatics
Budapest University of Technology and Economics. Budapest, Hungary
ladoczki.bence@vbk.bme.hu

ОСОБЛИВОСТІ ВИКОРИСТАННЯ НЕЙРОМЕРЕЖ В СИСТЕМАХ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ УПРАВЛІННЯ ВРОЖАЙНІСТЮ

Анотація. Стаття присвячена аналізу існуючих рішень щодо використання нейромережних моделей в системах підтримки прийняття рішень в агросекторі. Оптимізація управління розглядається застосовно до процесу вирощування озимої пшениці. Обрані математичні моделі передбачають формування систем звичайних диференціальних рівнянь для опису динаміки вегетаційного процесу. Запропоновано алгоритм прогнозування та управління вегетаційним процесом вирощування сільськогосподарської культури. Активізація нейронів та прийняття рішень нейромережею приймається відповідно до заданої норми відстані поточного стану та планованих показників якості врожаю.

Ключові слова: системи підтримки прийняття рішень; нейромережеві моделі; прогнозування вирощування сільськогосподарської культури.

1. ВСТУП

У ряді досліджень [1, 2] для оцінки врожайності та якості врожаю пропонується використання нейромереж зі структурою багатошаровий персептрон. Вхідними чинниками в моделі нейронної мережі були значення параметрів навколишнього середовища та енергоємність технології за роками. Такі моделі включають декілька нейронів вхідного шару та формують ряд вихідних величин що оцінюють якість та загальний результат вирощування сільськогосподарської продукції.

База даних для тренування моделі має бути побудована на підставі спостережень, проведених для конкретного регіону та конкретної культури. Це дозволить використовувати дані, що формувалися у природних умовах під впливом багатьох чинників, більшість із яких має випадковий характер.

Метою даної роботи є розвиток та вдосконалення існуючих нейромережних моделей у зв'язку з наступним.

1. Чинні моделі нейромережі не передбачають врахування впливу детермінованих управлінських параметрів, які дозволять вносити зміни в динаміку процесу дозрівання сільськогосподарської культури. Таким чином, нейромережева модель не враховує динаміку впливів, що управляють, наприклад, внесення добрив та ін. Відповідно не прораховується динаміка кінцевого результату і не формується алгоритм необхідної модифікації управління. У цій роботі розглядається рішення, яким чином при зміні вхідних впливів можна досягти необхідних показників цільової функції і як це зробити змінюючи значення управляючих впливів. Наприклад, використовуючи іригаційні процедури, внесення різних добрив та хімікатів, тощо. Як зазначалося в [2], вирішення цього завдання вкрай важливо для фермерських господарств на попередньому етапі -

для ухвалення рішення про доцільність або недоцільність вирощування цієї культури, зокрема озимої пшениці.

2. Модель, подана в [1], не враховує прогноз кореляційних зв'язків між стохастичними факторами, сукупні зміни яких, у свою чергу, можуть впливати на динаміку процесу та кінцевий результат. Для отримання прогнозних показників вхідних впливів передбачається використовувати кореляційні матриці, формуючи з допомогою системи диференціальних рівнянь. Такий підхід дозволяє використовувати добре розвинений апарат обчислювальних методів на вирішення завдання прогнозування.

Врахування стохастичних змін вхідних сигналів, а також коригування керуючих параметрів (внесення добрив, додатковий полив, відведення води тощо) дозволить прогнозувати діапазон очікуваних значень вихідних показників оцінки врожаю. Аналіз динаміки моделі нейромережі створить базу даних не тільки для оцінювання стану процесу вирощування врожаю, але дозволить побудувати систему підтримки прийняття рішень для управління даним процесом.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Рішення завдання управління, на відміну від завдання оцінювання стану та прогнозу результуючих показників вегетаційного процесу, потребує наявності динамічної моделі. Формування динамічної моделі у формі системи звичайних диференціальних рівнянь (СДР) дозволить прогнозувати динаміку змін взаємозв'язку входів та виходів нейромережі. Відповідно, алгоритм прогнозування та управління вегетаційним процесом вирощування сільськогосподарської культури, зокрема, без втрати загальності, озимої пшениці, можна записати у вигляді:

1. Ввівши значення випадкових факторів при $t=0$, що впливають і використовуючи нейромережу розраховуємо фіксовані значення вихідних параметрів і визначаємо підсумковий результуючий показник врожаю. Вважатимемо, що розраховані підсумкові значення мають відхилення від заданих цільових показників.

2. Формуємо коваріаційну матрицю вихідних показників для моменту t_1 і підбираємо керуючі впливи, при яких забезпечується найбільш близькі до оптимальних значення вихідних параметрів.

3. Використовуємо дані значення параметрів, що управляють (отримані в попередньому пункті) на інтервалі $[t_0, t_1]$ у вегетаційному процесі.

4. Контролюємо протікання вегетаційного процесу до моменту t_1 та знімаємо необхідні дані за вихідними показниками (проміжні значення для даного моменту) та значення випадкових впливів.

5. Проводимо зіставлення фактичних та очікуваних значень вихідних (результуючих) показників.

6. Формуємо коваріаційну матрицю випадкових впливів моменту t_1 і розраховуємо їх прогнозні значення моменту t_2 .

7. Формуємо коваріаційну матрицю вихідних показників для моменту t_2 і підбираємо управляючі впливи, при яких забезпечується найбільш близькі до оптимальних значення вихідних параметрів.

8. Використовуємо дані значення параметрів, що управляють (отримані в попередньому пункті) на інтервалі $[t_1, t_2]$ у вегетаційному процесі.

9. Повертаємось на п.5 і виконуємо наступні операції з циклу на наступному інтервалі часу.

10. Цикл обчислень повторюється по всьому інтервалі вегетаційного процесу до його завершення t_N .

Подальше удосконалення представленого алгоритму, що виходить за межі цієї статті, передбачає періодичний розрахунок кінцевих результатів всього вегетаційного процесу за допомогою нейромережевої моделі. Крім того, слід зазначити, що подальше уточнення нейромережевої моделі може використовувати введення вагових коефіцієнтів для різних параметрів вихідного шару нейромережевої моделі.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Запропонований алгоритм прийняття оптимальних управлінських рішень для динамічного процесу вирощування агропродукції використовує нейромережу, що доповнена динамічною моделлю у вигляді диференціальних рівнянь. Практичною цінністю даного алгоритму є можливість більш точно відповісти на питання про можливість досягнення необхідного результату, тобто вирішення оптимізаційної задачі.

ПОСИЛАННЯ

- [1] Dolia, M., Lysenko, V., Lendiel, T., Nakonechna, K., & Humeniuk, L. (2024). Neuron network prediction of damage of E. integriceps bug on winter wheat in Ukraine. Scientific Reports of the National University of Life and Environmental Sciences of Ukraine, 20(4),96-105. <https://doi.org/10.31548/dopovidi/3.2024.96>
- [2] Lysenko, V., Lendiel, T., Bolbot, I., & Nakonechnyy, I. (2022, October). Neural Network Structures for Energy-efficient Control of Energy Flows in Greenhouse Facilities. In 2022 IEEE 9th International Conference on Problems of Infocommunications, Science and Technology (PIC S&T) (pp. 21-26). IEEE.

Юрій Олегович Міловідов

Старший викладач кафедри комп'ютерних наук

Місце роботи: Національний університет біоресурсів і природокористування України, Київ,
Україна

yurii_milovidov@nubip.edu.ua

ВІЗУАЛІЗАЦІЯ АЛГОРИТМУ ПОШУКУ ОПТИМАЛЬНОГО ШЛЯХУ МІЖ ДВОМА ТОЧКАМИ КЛІТИННОГО ЛАБІРИНТУ З ПЕРЕШКОДАМИ, ЩО ДИНАМІЧНО ЗМІНЮЮТЬСЯ

Анотація. Представлена розроблена автором комп'ютерна програма для демонстрації роботи алгоритмів пошуку найкоротшого шляху на ділянці у вигляді клітинного лабіринту. Шлях між клітинами може мати різну вагу. Для пошуку найкоротшого шляху застосовується алгоритм Дейкстри для зваженого графа.

Keywords: Depth First Search, Breadth First Search, Dijkstra's algorithm.

1. ВСТУП

Метою представленої роботи є розробка програми для візуалізації алгоритму пошуку оптимального шляху на полі, яке можна уявити у вигляді лабіринту з переборними і непереборними перешкодами. Задача полягає в тому, щоб знайти оптимальний шлях між двома точками на полі та відобразити його. Лабіринт задається у вхідному файлі, в тому ж файлі вказуються координати входу і виходу, і для початку роботи нам необхідно вибрати потрібний лабіринт, програма повинна видати розмір найкоротшого шляху, намалювати лабіринт і показати цей шлях.

Існує досить багато різних методів вирішення такого завдання, кожен з яких ґрунтується на своїх принципах і прийомах, має унікальні переваги і, відповідно, недоліки. Для пошуку оптимального шляху в лабіринті обрано алгоритм Дейкстри (Dijkstra's algorithm).

Мета дослідження: Створити програмні засоби, які дозволяють візуалізувати виконання алгоритмів на графах пошуку оптимального шляху на ділянці поля. Це може бути як аграрне поле, так і поле бою.

2. ТЕОРЕТИЧНІ ПІДСТАВИ

Лабіринт представлений у вигляді матриці (двовірний масив). (Рис. 1). Кожна клітинка може бути або непереборною перешкодою, або може мати певну вагу (складність проходження даного відрізка шляху).

Якщо значення комірки = 0, то це непереборна перешкода, якщо комірка прохідна, то її значенням може бути дійсне число, яке відповідає складності досягнення цієї клітинки. Об'єкт може рухатися в 4-х напрямках: вгору, вниз, вліво, вправо.

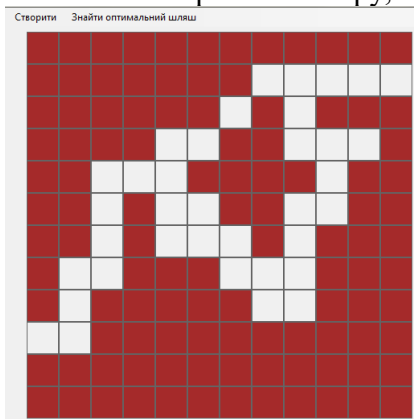


Рис. 1. Лабіринт представлений у вигляді двовимірної масиви.

Кожну комірку можна уявити як вершину графа. Якщо з неї є шлях до сусідньої комірки, то ці вершини графа пов'язані ребром відповідної ваги. Для подання графа в пам'яті комп'ютера використовується матриця суміжності – це квадратна матриця у якої кількість стовпців і рядків дорівнює кількості вершин графа. (Рис. 2)

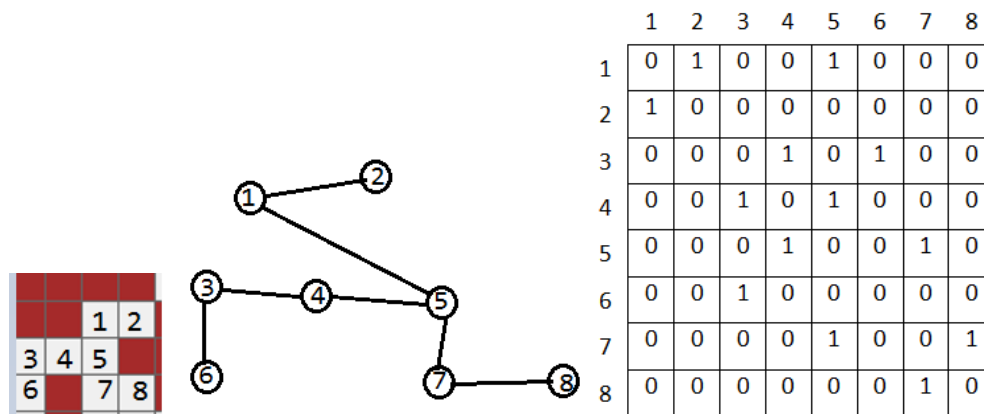


Рис. 2. Граф і відповідна матриця суміжності.

Алгоритм Дейкстри.

Алгоритм Дейкстри (англ. Dijkstra's algorithm) – алгоритм на графах, винайдений нідерландським вченим Е. Дейкстрою в 1959 році. Знаходить найкоротшу відстань від однієї з вершин графа до всіх інших. Алгоритм працює тільки для графів без ребер негативної ваги. Алгоритм широко застосовується в програмуванні і технологіях, наприклад, його використовують протоколи маршрутизації OSPF і IS-IS. Алгоритм Дейкстри застосовується для зваженого графа у разі, коли треба знайти шляху до всіх вершин у графі.

Кожній вершині зіставимо мітку – мінімальну відому відстань від цієї вершини до а. Алгоритм працює покрокове – на кожному кроці він «відвідує» одну вершину і намагається зменшувати мітки. Робота алгоритму завершується, коли всі вершини відвідані.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для розробки програм обрана мова програмування C#.

Програма має інтуїтивно зрозумілий інтерфейс (рис. 3).

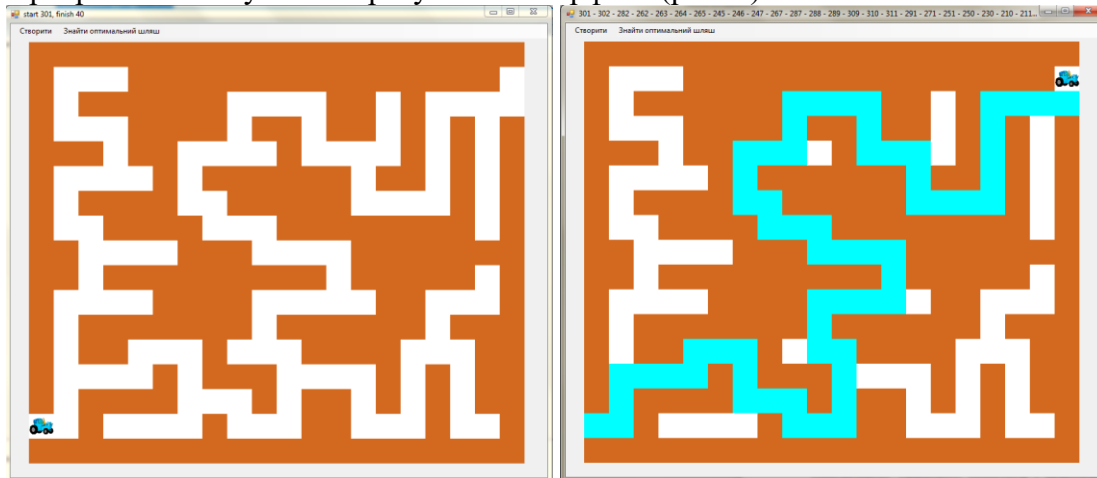


Рис. 3. Відображення оптимального шляху.

Програма візуалізації алгоритму пошуку оптимального шляху в лабіринті має величезне практичне значення і може застосовуватися на фермерських ланах, на полях бою з використанням даних, отриманих від БПЛА. Програма застосовувалася автором під час викладання дисциплін «Алгоритми і структури даних» і «Об'єктно-орієнтоване програмування» в Національному університеті біоресурсів і природокористування України. Студенти спостерігали за всіма процесами під час роботи програм і наочно оцінили їх користь.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Для візуалізації пошуку оптимального шляху в лабіринті запропонована реалізація алгоритму Дейкстри на мові C# в застосуванні Windows Forms .NET. Крім зазначеної вище практичної користі програма може стати підмогою як для викладачів, так і для студентів, які вивчають дисципліни «Алгоритми і структури даних» і «Об'єктно-орієнтоване програмування».

ЛІТЕРАТУРА

1. Кормен, Томас; Лейзерсон, Чарльз; Рівест, Рональд; Стайн, Кліфорд (2019). 16.3: Коды Гафмена. Вступ до алгоритмів (вид. 3). К.І.С. с. 443–451. ISBN 978-617-684-239-2
2. Об'єктно-орієнтоване програмування: навчальний посібник. Третє видання / Ю.О. Міловідов. – К. НУБіП України, 2025. – 334с.
3. W3schoolsUA. українською. C# ООП. [Електронний ресурс]: Режим доступу: https://w3schoolsua.github.io/cs/cs_oop.html
4. Порєв В.М. Об'єктно-орієнтоване програмування: конспект лекцій навч. посіб. для студ спеціальності 121 «Інженерія програмного забезпечення» / КПІ ім. Ігоря Сікорського. – Київ : КПІ ім. Ігоря Сікорського, 2022. – 271с. Режим доступу: http://library.megu.edu.ua:8180/jspui/bitstream/123456789/4128/1/2022-konsp_lec_oop_porev.pdf

SECTION 4. INFORMATION SYSTEMS AND TECHNOLOGIES IN THE ECONOMY, TECHNOLOGY AND NATURAL USE / ІНФОРМАЦІЙНІ СИСТЕМИ І ТЕХНОЛОГІЇ В ЕКОНОМІЦІ, ТЕХНІЦІ ТА ПРИРОДОКОРИСТУВАННІ

Максим Мокрієв

канд.екон.наук, доцент, керівник центру дистанційних технологій навчання

Місце роботи: НУБіП України, каф.інформаційних систем і технологій, Київ, Україна

ORCID ID: 0000-0002-6717-3884

m.mokriiev@nubip.edu.ua

ІНТЕГРАЦІЯ LMS MOODLE ТА CMS WORDPRESS ДЛЯ СТВОРЕННЯ ВІДКРИТОГО ОСВІТНЬОГО СЕРЕДОВИЩА ДЛЯ ДОДАТКОВОГО САМОСТІЙНОГО НАВЧАННЯ МАГІСТРІВ

Анотація. У статті розглядається технічна та педагогічна можливість поєднати дві популярні платформи інтернету LMS Moodle та CMS Wordpress для створення потужного середовища для самостійного навчання для студентів магістрів вищого навчального закладу. Поєднання їх найкращих сторін в одному середовищі створить синергетичний ефект та надасть більшої привабливості і зрозумілості користувацького досвіду при роботі з цифровими навчальними ресурсами.

Ключові слова: вища освіта; індивідуалізація навчання; інформаційні технології; самостійне навчання; lms moodle; cms wordpress.

ВСТУП.

Сучасні вимоги до фахівців з вищою освітою, особливо на магістерському рівні, виходять за межі засвоєння предметних знань. Ключовими компетенціями 21 століття є навички безперервного навчання та високий рівень автономії. Такі вимоги спонукають університети до впровадження педагогіки самостійного навчання. Такий підхід представляє фундаментальний зсув від традиційної моделі, де викладач диктує цілі та ресурси, до моделі, де студент сам ініціює процес, встановлює навчальні цілі, визначає необхідні ресурси та методи оцінювання прогресу. Це не лише сприяє розвитку критичного мислення, але й підвищує внутрішню мотивацію, відчуття контролю та впевненості у власних силах.

Постановка проблеми. Проте реалізація ефективного середовища самостійного навчання у відкритих дистанційних навчальних середовищах стикається зі значними викликами. Хоча платформи, які використовуються для побудови таких середовищ, як-от Moodle, пропонують гнучкі та інклюзивні шляхи для навчання, менш структуровані формати можуть призводити до проблем з мотивацією, тайм-менеджментом та відчуттям ізоляції серед учнів. Щоб подолати ці перешкоди та забезпечити академічний успіх, необхідно розробляти не лише педагогічні підходи, але й вдосконалювати технологічне середовище, щоб воно стало більш привабливим, інтуїтивно зрозумілим та гнучким.

Аналіз останніх досліджень і публікацій. Успіх самостійного навчання значною мірою залежить від навичок саморегуляції. Академічні моделі саморегуляції поділяють навчальний процес на етапи планування, виконання та оцінювання, під час яких застосовуються стратегії організації, самомоніторингу та критичної рефлексії.[2] Ці стратегії повинні бути підтримані цифровим середовищем. Для студентів, які лише формують ці навички, критично важливим є забезпечення адекватного скрафолдингу та структуризації навчальних активностей. [3] Надання більшого контролю над навчальним процесом завдяки самостійному навчанню підвищує мотивацію та покращує навички критичного мислення, що є безпосередньою перевагою для магістерського рівня. [4]

LMS Moodle — це потужна платформа, розроблена для підтримки структурованого навчального процесу. Її архітектура високо модульна, що дозволяє ефективно створювати онлайн-курси, управляти завданнями, оцінюванням та відстежувати прогрес студентів.[5] Moodle є незамінним ядром для формальної освіти та колаборації.

Натомість, CMS WordPress є провідною платформою для управління контентом, відомою своєю зручністю для користувачів, високою кастомізацією та потужною екосистемою плагінів. [6] WordPress ефективно використовується для публікації новин, організації додаткових матеріалів, маркетингу курсів (завдяки SEO-функціоналу) та створення привабливих фронтенд-інтерфейсів.[7]

Інтеграція цих двох платформ дозволяє досягти синергетичного ефекту, подолавши їхні індивідуальні недоліки. WordPress може вирішити проблему варіативного та потенційно складного UX Moodle, забезпечуючи більш інтуїтивний та конверсійний інтерфейс.

Метою даної роботи є теоретичне обґрунтування, розробка концептуальної та архітектурної моделі інтеграції LMS Moodle та CMS WordPress, а також формування педагогічних стратегій її використання для створення ефективного, гнучкого та безшовного відкритого середовища для самостійного навчання, спрямованого на підтримку додаткового самостійного навчання магістрів.

РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ.

Пропонована архітектурна модель середовища для самостійного навчання базується на дворівневому підході, який чітко розмежовує функціональні обов'язки LMS та CMS.

Рівень 1: Фронтенд (CMS WordPress). Цей рівень є точкою входу для магістрів. Він відповідає за публічний контент, інформаційні сторінки, зовнішню комунікацію, інтуїтивну навігацію, реєстрацію нових користувачів та відображення персоналізованих інформаційних панелей. WordPress, завдяки своїй гнучкості, дозволяє створити привабливий каталог курсів, імпортований з Moodle [7], та інтегрувати додаткові функції (наприклад, e-commerce, соціальні мережі), що робить середовище динамічним та сучасним. [8]

Рівень 2: Бекенд (LMS Moodle). Цей рівень функціонує як захищене, структуроване ядро. Moodle використовується для розміщення навчальних матеріалів, проведення інтерактивних занять, адміністрування оцінювання, відстеження академічного прогресу та забезпечення інструментів для спільної роботи (форуми, чати).[9]

Взаємодія між цими рівнями здійснюється через програмний інтерфейс Moodle (API), який використовується плагіном-містком (наприклад, Edwiser Bridge) на стороні WordPress.

Для забезпечення безшовного користувацького досвіду, який є критичним для підтримки мотивації та зниження технологічного бар'єру, необхідна бездоганна робота двох ключових механізмів:

Єдиний Вхід (Single Sign-On, SSO). Реалізація SSO дозволяє магістрам, які вже зареєстровані на WordPress, отримати доступ до навчальних матеріалів Moodle без необхідності повторного входу. Використання ідентичних облікових даних для обох платформ спрощує процес, економить час та знижує кількість звернень до технічної підтримки. Деталізована Двостороння Синхронізація. Це забезпечує єдність даних в екосистемі:

– Синхронізація Користувачів: Реєстрація користувача на WordPress ініціює автоматичну реєстрацію та створення відповідного облікового запису в Moodle.

– Синхронізація Курсів та Категорій: Курси, створені в Moodle, імпортуються до WordPress для відображення у зручному каталозі на фронтенді. Це дозволяє магістрам легко вибрати додаткові курси для ССН.

– Синхронізація Зарахування (Enrollment Sync): Рішення про зарахування на курс, прийняте на WordPress (наприклад, після вибору курсу), миттєво передається до Moodle, надаючи користувачеві доступ до навчальних активностей.

– Синхронізація Прогресу: Це найважливіший технічний елемент для ССН. Дані про проходження курсу (завершені активності, оцінки, прогрес) синхронізуються з Moodle до WordPress у реальному часі.

Коли дані про прогрес автоматично відображаються на персоналізованому дашборді WordPress, це перетворює технічну функцію на потужний педагогічний інструмент. Магістри можуть здійснювати ефективний саомоніторинг та оцінювання своїх стратегій навчання. Це пряма підтримка саморегуляції, оскільки забезпечує швидкий та зрозумілий зворотний зв'язок.

ВИСНОВКИ. Проведене дослідження підтверджує архітектурну та педагогічну доцільність інтеграції LMS Moodle та CMS WordPress для формування Відкритого Освітнього Середовища, орієнтованого на додаткове самостійне навчання магістрів. Синергетичний ефект досягається за рахунок використання WordPress як гнучкого, візуально привабливого та зручного фронтенду, що підвищує загальну задоволеність користувачів, тоді як Moodle функціонує як надійне ядро для структурованого навчального процесу та академічного трекінгу.

ПОСИЛАННЯ

1. Lobos K, Cobo-Rendón R, Bruna Jofré D and Santana J (2024) New challenges for higher education: self-regulated learning in blended learning contexts. Front. Educ. 9:1457367. doi: 10.3389/feduc.2024.1457367
2. Mark Stevens (2020). Expertise, Complexity, and Self-Regulated Engagement: Lessons from Teacher Reflection in a Blended Learning Environment. Journal of Online Learning Research (2020) 6(3), 177-200
3. Robinson, J. D., & Persky, A. M. (2020). Developing Self-Directed Learners. American journal of pharmaceutical education, 84(3), 847512. doi:10.5688/ajpe847512
4. О. Г. Кузьмінська, О. Г. Глазунова, М. В. Мокрієв, В. І. Корольчук, Т. В. Волошина. Технології інтеграції освітніх ресурсів і сервісів в умовах дистанційного навчання. Humanitarian studios: pedagogics, psychology, philosophy vol 13(2) 2022.
5. Moodle vs WordPress: Which Platform is Best for Your Online Course? Blog: URL: <https://senseilms.com/moodle-vs-wordpress/>
6. Edwiser Bridge – WordPress Moodle Integration. Wordpress Site. URL: <https://wordpress.org/plugins/edwiser-bridge/>
7. Integrating Moodle And WordPress For A Seamless Learning Experience. URL: <https://metadesignsolutions.com/integrating-moodle-and-wordpress-for-a-seamless-learning-experience/>
8. Yunos, Nasruddin & Muslim, Nazri & Hussain, Afifuddin & Hasim, Nur & Nazri, Nurul & Hamsan, M.H.. (2024). Best Practice of Moodle Implementation for E-Learning: A Perspective of Public University Lecturers. Journal of Ecohumanism. 3. 261-268. 10.62754/joe.v3i5.3899.

Вікторія Усік

Студентка 3 курсу, спеціальності «Інформаційні системи та технології»

ist23-v.usik@nubip.edu.ua,

Науковий керівник Смолій В.М.

АВТОМАТИЗАЦІЯ ЕКОЛОГІЧНОГО МОНІТОРИНГУ ЗАСОБАМИ СУЧАСНИХ ІТ

Анотація. У статті розглянуто актуальні питання автоматизації екологічного моніторингу з використанням сучасних інформаційних технологій. Показано, що глобальні процеси цифровізації створюють нові можливості для збору, аналізу та візуалізації даних про стан довкілля. Здійснено аналіз основних напрямів впровадження інформаційних систем у природоохоронну діяльність, зокрема використання сенсорних мереж, супутникового моніторингу, інтернету речей (IoT) та систем штучного інтелекту для оцінювання екологічних ризиків. Розглянуто приклади автоматизованих систем контролю забруднення повітря, води та ґрунтів, що функціонують на регіональному та глобальному рівнях. Обґрунтовано важливість інтеграції екологічних даних у єдині цифрові платформи для підвищення ефективності прийняття управлінських рішень у сфері охорони навколишнього середовища.

Запропоновано концептуальну модель інформаційної системи екологічного моніторингу, яка передбачає використання модулів автоматичного збору, обробки, аналізу та відображення даних у режимі реального часу. Визначено перспективи розвитку напрямів автоматизації моніторингу шляхом впровадження інтелектуальних технологій аналізу даних, що дозволяють здійснювати прогнозування стану довкілля та попередження екологічних катастроф.

Ключові слова: екологічний моніторинг; автоматизація; інформаційні технології; цифровізація; інтернет речей; штучний інтелект.

1. ВСТУП

Постановка проблеми. Сучасний етап розвитку людства характеризується стрімким зростанням техногенного навантаження на природне середовище. Забруднення повітря, водних ресурсів і ґрунтів, глобальне потепління, зниження біорізноманіття — це лише частина проблем, які потребують постійного моніторингу. Традиційні методи спостереження не забезпечують достатньої оперативності та точності даних, що ускладнює прийняття ефективних управлінських рішень. У зв'язку з цим актуальним є впровадження автоматизованих систем екологічного моніторингу, що базуються на сучасних інформаційних технологіях [0]

Аналіз останніх досліджень і публікацій. Проблематика автоматизації екологічного моніторингу активно розглядається у працях вітчизняних та зарубіжних науковців. Зокрема, у роботах [0], [0] підкреслюється важливість інтеграції геоінформаційних систем (ГІС) для аналізу екологічних даних. У дослідженнях [0] акцентується на застосуванні інтернету речей (IoT) для моніторингу стану атмосферного повітря та водних ресурсів. Деякі науковці [0] вказують на потенціал штучного інтелекту в прогнозуванні екологічних ризиків. Незважаючи на значний науковий інтерес, залишається проблема практичної інтеграції окремих технологічних рішень у комплексну систему автоматизованого контролю стану довкілля.

Мета публікації. Метою статті є узагальнення підходів до автоматизації екологічного моніторингу та розробка концептуальної моделі інформаційної системи моніторингу довкілля з використанням сучасних ІТ-технологій.

2. ТЕОРЕТИЧНІ ОСНОВИ

Екологічний моніторинг — це система спостережень, оцінки та прогнозу змін стану навколишнього середовища під впливом природних і антропогенних факторів. Автоматизація цього процесу передбачає використання сенсорних мереж для збору

даних, телекомунікаційних каналів для передачі інформації, баз даних і аналітичних модулів для її обробки. Важливими технологічними складовими є: Інтернет речей (IoT) — забезпечує з'єднання датчиків і пристроїв у єдину мережу; Big Data — дає можливість зберігати великі обсяги екологічних даних і здійснювати їх статистичний аналіз; Геоінформаційні системи (GIS) — дозволяють візуалізувати результати спостережень у просторі; Штучний інтелект (AI) — застосовується для прогнозування екологічних тенденцій.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Методологічну основу дослідження становить системний підхід, що розглядає екологічний моніторинг як комплексну інформаційну систему збору, обробки та аналізу даних про стан довкілля.

У роботі застосовано такі методи:

- **аналітичний** — для вивчення наукових джерел і нормативних документів щодо автоматизації моніторингу;
- **системного аналізу** — для визначення структури та потоків інформації в екологічних системах;
- **моделювання** — для створення концептуальної моделі автоматизованої системи моніторингу;
- **порівняльний** — для оцінки ефективності сучасних IT (IoT, Big Data, GIS, AI);
- **статистичної обробки даних** — для перевірки достовірності екологічних показників.

Експериментальна база — відкриті екологічні дані з міжнародних платформ (EEA Open Data Portal, AQI, NASA EOS).

Основні показники ефективності: швидкість збору, точність вимірювань, час обробки та стабільність системи.

Критеріями оцінки стали підвищення оперативності моніторингу, зменшення людського впливу й покращення якості аналітичних даних.

Дослідження виконано в межах теми «Інтелектуальні інформаційні системи екологічної безпеки регіонів» (реєстр. № 0123U109876).

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Запропонована концепція автоматизованої системи екологічного моніторингу базується на трирівневій архітектурі:

1. **Рівень збору даних** — мережа сенсорів (наприклад, датчики якості повітря, вологості, температури, рівня CO₂ тощо).
2. **Рівень обробки** — сервери або хмарні сервіси, що здійснюють обробку даних, фільтрацію та виявлення аномалій.
3. **Рівень візуалізації** — вебінтерфейси або мобільні додатки для відображення показників у реальному часі.

Практичні результати свідчать, що використання автоматизованих систем моніторингу дозволяє підвищити точність спостережень на 20–40 %, зменшити людський фактор і забезпечити оперативне реагування на екологічні загрози [0].

Такі системи можуть бути інтегровані у регіональні екологічні платформи управління природними ресурсами. Наприклад, у багатьох європейських країнах створені національні цифрові екологічні реєстри, що збирають дані з автоматизованих станцій моніторингу повітря, води та ґрунтів.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У результаті дослідження обґрунтовано необхідність застосування сучасних інформаційних технологій для автоматизації процесів екологічного моніторингу. Визначено, що використання таких технологій, як Інтернет речей, великі дані, геоінформаційні системи та штучний інтелект, підвищує точність і оперативність отримання екологічної інформації, а також сприяє прийняттю ефективних управлінських рішень у сфері природокористування.

Запропонована концептуальна модель автоматизованої системи моніторингу дозволяє забезпечити комплексний підхід до збору, зберігання, аналізу та візуалізації даних про стан довкілля, що створює основу для формування інтегрованих екологічних інформаційних платформ регіонального та національного рівня.

Перспективами подальших досліджень є розробка інтелектуальних модулів прогнозування стану навколишнього середовища, удосконалення алгоритмів обробки великих даних, а також впровадження блокчейн-технологій для підвищення прозорості, достовірності та захищеності екологічної інформації.

ПОСИЛАННЯ

1. O. Boyko, Environmental Informatics and Sustainable Development, Kyiv: NTUU KPI, 2020.
2. M. Smith, "GIS-based environmental monitoring systems," Environmental Modelling & Software, vol. 145, pp. 102-112, 2021.
3. V. Koval, "Digital transformation in environmental management," Ukrainian Journal of Ecology, vol. 11, no. 3, pp. 45-53, 2022.
4. L. Petrova, "IoT applications for smart environmental monitoring," Sensors, vol. 22, no. 8, 2022.
5. A. Sharma, "AI for climate and pollution prediction," Sustainability, vol. 15, no. 2, 2023.
6. European Environment Agency, "Automated environmental data collection," EEA Reports, 2024.

Анна Паламарчук

Студентка 3 курсу, спеціальності «Інформаційні системи та технології»

Ist23-a.palamarchuk@nubip.edu.ua

Науковий керівник Смолій В.М.

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ МОНІТОРИНГУ СТАНУ ДОВКІЛЛЯ В УМОВАХ ЗМІНИ КЛІМАТУ

Анотація. У статті розглянуто роль сучасних інформаційних технологій у системах моніторингу стану довкілля в умовах посилення кліматичних змін. Проаналізовано особливості використання інтелектуальних систем, супутникових технологій, геоінформаційних систем (ГІС), технологій Інтернету речей (IoT) та хмарних обчислень у забезпеченні оперативного збору, аналізу та прогнозування екологічних показників. Показано, що інформатизація екологічного моніторингу сприяє формуванню єдиних баз даних про стан навколишнього середовища, забезпечує підвищення точності вимірювань, своєчасне виявлення відхилень і прогнозування ризиків. Окрему увагу приділено практичним аспектам впровадження цифрових технологій в екологічному управлінні, а також проблемам інтеграції різномірних джерел даних та забезпеченню їх сумісності. Визначено перспективи розвитку інтелектуальних аналітичних систем, що базуються на технологіях штучного інтелекту та машинного навчання, для прогнозування впливу кліматичних змін на регіональному рівні.

Ключові слова: інформаційні технології; моніторинг довкілля; кліматичні зміни; геоінформаційні системи; Інтернет речей.

1. ВСТУП

Стан навколишнього середовища є одним із найважливіших показників якості життя населення та стійкого розвитку держави. В умовах глобальних кліматичних змін особливої актуальності набуває створення ефективних систем моніторингу довкілля, що базуються на сучасних інформаційних технологіях. Використання цифрових інструментів дозволяє не лише забезпечити точність збору даних, а й своєчасно реагувати на негативні екологічні тенденції.

Постановка проблеми. Сучасні кліматичні зміни зумовлюють необхідність переходу від періодичних спостережень за станом довкілля до системного та безперервного моніторингу, заснованого на інформаційних технологіях. Традиційні методи екологічного контролю вже не забезпечують достатньої швидкості оброблення даних і масштабності охоплення територій. Це потребує впровадження автоматизованих систем, здатних збирати, аналізувати й передавати інформацію про стан атмосферного повітря, водних ресурсів, ґрунтів, рослинності та інших компонентів біосфери в реальному часі [0].

Аналіз останніх досліджень і публікацій. Проблематика інформатизації екологічного моніторингу активно розглядається у працях зарубіжних та вітчизняних науковців. У роботах [0], [0] описано використання геоінформаційних систем для побудови екологічних карт і прогнозів ризику забруднення. У [0] досліджено роль Інтернету речей у створенні сенсорних мереж для вимірювання температури, вологості, якості повітря та рівня шуму. Автори [0] підкреслюють значення штучного інтелекту для виявлення закономірностей у великих масивах екологічних даних. Проте залишаються відкритими питання інтеграції гетерогенних даних, стандартизації протоколів обміну й забезпечення сумісності між державними та регіональними системами моніторингу.

Мета публікації. Метою статті є аналіз сучасних підходів до використання інформаційних технологій у системах моніторингу стану довкілля, визначення їх ролі в умовах зміни клімату та окреслення перспектив розвитку цифрових рішень у цій сфері. Особлива увага приділяється можливостям інтеграції геоінформаційних систем,

хмарних сервісів, сенсорних мереж та технологій штучного інтелекту для створення єдиного інформаційного простору екологічних даних. Реалізація таких підходів сприятиме підвищенню точності прогнозування кліматичних ризиків, оптимізації управлінських рішень у сфері природокористування та забезпеченню екологічної безпеки на національному рівні.

2. ТЕОРЕТИЧНІ ОСНОВИ

Моніторинг довкілля - це система спостережень, оцінювання, прогнозування стану природного середовища та процесів, що в ньому відбуваються. Сучасні інформаційні технології забезпечують перехід від ручного збору даних до автоматизованих систем спостережень, що інтегрують сенсорні пристрої, супутникові знімки, мобільні додатки та хмарні сервіси.

Основними компонентами таких систем є:

- сенсорна мережа (IoT) - забезпечує збір фізичних параметрів у реальному часі;
- геоінформаційна платформа (ГІС) - виконує просторовий аналіз і візуалізацію даних;
- аналітичний модуль (AI/ML) - аналізує тренди, формує прогнози;
- інтерфейс користувача - забезпечує доступ до даних через вебпортали або мобільні застосунки.

У контексті кліматичних змін такі технології дозволяють моделювати сценарії розвитку екосистем і прогнозувати наслідки екстремальних явищ - посух, паводків, температурних аномалій.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для досягнення мети дослідження використано поєднання теоретичних та прикладних методів. Теоретичну основу становить аналіз сучасних підходів до моніторингу довкілля, зокрема супутникового спостереження, використання мереж сенсорів IoT та систем геоінформаційного аналізу (GIS). Практична частина дослідження ґрунтується на методах збору, фільтрації та візуалізації екологічних даних із відкритих джерел (Copernicus, NASA Earth Observations). Для моделювання процесів обробки даних застосовано середовище Python із бібліотеками Pandas, Matplotlib та Scikit-learn.

Оцінювання ефективності проведено за критеріями точності вимірювань, швидкодії обробки та можливості інтеграції результатів у національні системи екологічного моніторингу.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Розвиток інформаційних технологій сприяє створенню національних і регіональних екологічних платформ. Зокрема, у Європейському Союзі діють програми Copernicus та INSPIRE, що надають відкриті супутникові дані для екологічного аналізу. В Україні аналогічні функції поступово реалізуються через інтеграцію даних Державної екологічної інспекції, Держгеокадастру та Українського гідрометцентру.

Використання Інтернету речей у моніторингу дозволяє створювати розподілені системи збору даних - наприклад, безпілотні датчики якості повітря або станцій спостереження за водними об'єктами. Такі системи передають дані до центральних серверів, де проводиться обробка з використанням штучного інтелекту для виявлення трендів.

Інтеграція Big Data технологій забезпечує аналіз великих обсягів історичних даних, що дозволяє моделювати вплив кліматичних змін на довкілля. Наприклад,

поєднання супутникових знімків і сенсорних показників допомагає прогнозувати деградацію ґрунтів або зменшення біорізноманіття [0].

Використання хмарних обчислень забезпечує масштабованість та доступність таких систем, що особливо важливо для країн із обмеженими ресурсами. Проте залишається проблема кібербезпеки екологічних даних та необхідність створення єдиних стандартів для їх обміну.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Інформаційні технології є ключовим інструментом підвищення ефективності моніторингу стану довкілля. Їх упровадження забезпечує не лише точність вимірювань і оперативність реагування, а й формування цілісної системи спостережень, здатної відстежувати динаміку кліматичних процесів у реальному часі. Використання технологій обробки великих даних, машинного навчання та супутникового моніторингу дозволяє підвищити достовірність екологічних прогнозів і своєчасно виявляти потенційні ризики для природного середовища.

Подальші дослідження доцільно спрямувати на створення єдиної інтегрованої системи екологічного моніторингу України з відкритим доступом до даних, що сприятиме прозорості екологічної інформації та підвищенню рівня громадської участі у вирішенні природоохоронних питань. Важливими напрямками розвитку залишаються впровадження аналітичних платформ на базі штучного інтелекту, розбудова IoT-мереж для збору первинних даних і створення механізмів автоматизованого реагування на екологічні загрози.

Такі рішення сприятимуть сталому природокористуванню, забезпеченню екологічної безпеки національного рівня та формуванню нової екологічної культури в умовах глобальної цифровізації.

ПОСИЛАННЯ

1. M. Turcotte, *Environmental Data Analysis with R*, CRC Press, 2022.
2. V. V. Krysanova, "GIS-Based Modeling of Climate Impacts on Water Resources," *Environmental Modelling & Software*, vol. 152, 2023.
3. O. S. Moroz, "Geoinformation Support for Regional Ecological Monitoring Systems," *Ukrainian Journal of Environmental Studies*, no. 4, pp. 15–22, 2022.
4. A. Singh and R. Kumar, "IoT-Based Environmental Monitoring: A Review," *Sensors*, vol. 21, no. 8, pp. 2550–2564, 2021.
5. L. Chen et al., "Artificial Intelligence for Environmental Monitoring and Climate Forecasting," *IEEE Access*, vol. 10, pp. 103112–103125, 2022.
6. European Environment Agency, "Copernicus Programme Data for Climate and Environment," *EEA Reports*, 2024.

Вікторія Смолій

доктор технічних наук, професор, професор кафедри інформаційних систем і технологій
Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID 0000-0002-1268-7837

vmsmolij@nubip.edu.ua

Натан Смолій

магістрант кафедри інформаційних систем і технологій Національний технічний університет
України «Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

ORCID ID: 0009-0002-3763-6726

hoibbitizukrainy@gmail.com

МОДУЛЬ МАШИННОГО ЗОРУ ДЛЯ РОЗПІЗНАВАННЯ ОБ'ЄКТІВ

Анотація. В роботі для реалізації завдання сегментації зображень було сформовано спеціалізований датасет для об'єктів захисного спорядження. На етапі підготовки до тренування моделі до цього набору даних було застосовано процедури аугментації, зокрема віддзеркалення, обертання, масштабування та зміну яскравості, що дозволило суттєво підвищити різноманітність навчальних прикладів. Збільшення розміру датасету забезпечило більш збалансоване представлення даних і підвищило узагальнювальну здатність моделі. Отримані в роботі результати підтверджують доцільність і ефективність запропонованого авторами для даної проблеми окремого навчання моделей для сегментації та класифікації зображень.

Ключові слова: виявлення об'єктів, модель детекції, сегментація зображень, попереднє навчання моделі, процедури аугментації, донавчання, бінарні маски об'єктів.

Робота присвячена розробці та дослідженню методів виявлення об'єктів на зображеннях і у відеопотоці з використанням сучасних технологій глибокого навчання. Основна увага приділяється створенню, навчанню та практичному застосуванню моделей детектування об'єктів на основі нейронних мереж, що забезпечують високу точність і швидкість обробки візуальних даних [1]. У процесі виконання дослідження необхідно визначити об'єкти заданого типу на зображеннях, отриманих із відеокамери, а також розробити модель детекції, яка зможе ефективно ідентифікувати та локалізувати ці об'єкти в різних умовах освітлення, масштабу та перспективи. Для досягнення поставленої мети пропонується використати попередньо натренований детектор, провести його донавчання (fine-tuning) на спеціалізованій вибірці даних і адаптувати модель до конкретних умов задачі [2, 3]. Отримані результати моделювання та навчання підлягають оцінюванню на тестовій вибірці, що дозволяє проаналізувати точність, повноту та швидкодію моделі, а також визначити можливі напрямки її подальшої оптимізації та вдосконалення.

У межах дослідження було застосовано модель YOLOv11s для сегментації зображень, що містять об'єкти захисного спорядження. Проведено порівняльний аналіз якості результатів між підходами з попереднім навчанням моделі та zero-shot learning. Тестування моделей здійснювалося у середовищі Ultralytics із використанням мови програмування Python 3.11.

Для реалізації завдання сегментації зображень було сформовано спеціалізований датасет, що складався із 150 початкових зображень, на яких представлені об'єкти захисного спорядження. На етапі підготовки до тренування моделі до цього набору даних було застосовано процедури аугментації, зокрема віддзеркалення, обертання, масштабування та зміну яскравості, що дозволило суттєво підвищити різноманітність навчальних прикладів. У результаті цих дій обсяг датасету збільшився до 390 зображень, що забезпечило більш збалансоване представлення даних і підвищення узагальнювальної здатності моделі. Приклади зображень із сформованого датасету подано на рис. 1.



Рис. 1. Приклад зображення із сформованого датасету

Ненавчена модель виявилася малоприсадиною для прямого порівняння з донавченою, оскільки попереднє тренування було орієнтоване на розпізнавання базових категорій об'єктів. Для оцінювання ефективності роботи моделі в межах створеного датасету було проаналізовано графік зміни метрики IoU (Intersection over Union) протягом процесу навчання, що дало змогу відстежити динаміку покращення якості сегментації.

Оскільки донавчання виконувалося у віддаленому середовищі, застосовувалися рекомендовані параметри, визначені розробниками сервісу [4]. Серед них ключову роль для задачі сегментації відіграють такі гіперпараметри: оптимізатор Stochastic Gradient Descent, швидкість навчання (learning rate) — 0.01, коефіцієнт моменту (momentum) — 0.937 та регуляризаційний параметр weight decay — 0.0005, які забезпечують стабільність і збіжність процесу оптимізації моделі.

Функція втрат у моделі сегментації формується як поєднання трьох складових компонентів. Box loss реалізовано у вигляді CIOU Loss (Complete IoU), який відповідає за точність регресії обмежувальних прямокутників. Classification loss базується на Binary Cross-Entropy (BCE) Loss і використовується для класифікації об'єктів за їх типами. Mask loss також побудовано на BCE Loss, що забезпечує коректну бінаризацію масок об'єктів на зображенні. Додатково для підвищення стабільності та точності сегментації застосовувався Dice loss, який покращує збалансованість між передбаченими та реальними областями об'єктів.

Zero-shot модель була застосована для розв'язання задачі сегментації зображень з метою виявлення елементів захисного будівельного спорядження, таких як жилети, маски та рукавиці. Для цього використано SAM (Segment Anything Model) — універсальну модель сегментації, створену компанією Meta. Її робота ґрунтується на принципі promptable segmentation, тобто модель отримує підказку (prompt) від користувача і виконує точну сегментацію обраного об'єкта на зображенні.

Архітектурно SAM побудована на Vision Transformer (ViT), який відповідає за витягнення ознак, та масковому предикторі, що генерує бінарні маски об'єктів. Результати сегментації, отримані за допомогою SAM, характеризуються високою точністю відтворення контурів і форм об'єктів, а також універсальністю — модель ефективно працює з будь-якими типами об'єктів без потреби у додатковому донавчанні.

Результати роботи моделі zero-shot підходу за наперед заданими точками та порівняння з YOLO наведені на рис. 2, а та б відповідно.

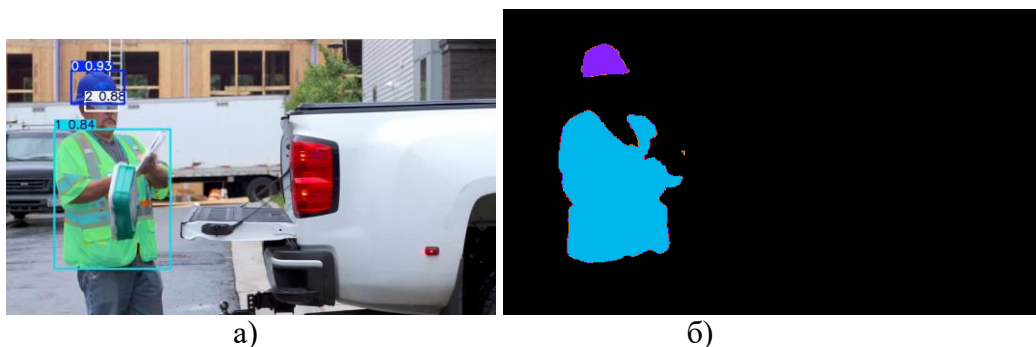


Рис. 2. Результати роботи моделі zero-shot підходу та порівняння з YOLO

Під час навчання моделі було застосовано спеціальний параметр, характерний для задач сегментації — $\text{mask_ratio} = 4$, який визначає ступінь наближення форми сегментованої області до графічних примітивів. У результаті значення метрики IoU (перетину з розміченими даними) досягло приблизно 40%, що свідчить про задовільну точність моделі для поставленого завдання. Додатково відзначено ефективність моделі за швидкістю — середній час виконання операції сегментації одного зображення становив близько 120 мс, що демонструє її придатність для використання в системах реального часу.

ВИСНОВКИ

У результаті проведених досліджень було сформовано спеціалізований датасет для сегментації зображень об'єктів захисного спорядження, що забезпечив якісну основу для навчання моделей комп'ютерного зору. Застосування процедур аугментації, таких як віддзеркалення, обертання, масштабування та зміна яскравості, дозволило суттєво розширити обсяг і різноманітність даних, підвищивши їх репрезентативність. Отримані результати свідчать, що збільшення обсягу датасету позитивно вплинуло на збалансованість даних і узагальнювальну здатність моделей. Проведене порівняння підтвердило ефективність використання підходу окремого навчання моделей для задач сегментації та класифікації. Таким чином, запропонований підхід може бути використаний як ефективне рішення для автоматизованої обробки зображень у завданнях технічного контролю та безпеки.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. V.M.Smolij, N.V.Smolij, S.P.Sayapin, Search and classification of objects in the zone of reservoirs and coastal zones, CEUR Workshop Proceedings 3666 (2024) 37–51. doi:ceur-ws.org/Vol-3666/ paper04.pdf10.1002/pssb.2220940160.
2. J.Tang, N.Xie, K.Li, Y.Liang, X.Shen, Trajectory tracking control for fixed-wing uav based on ddpg, Journal of Aerospace Engineering 37 (2024) 251–283. doi:10.1061/JAEEZ.ASENG-5286.
3. V.M.Smolij, N.V.Smolij, O.Y.Kovalenko, M.Z.Shvydenko, Channel extractor for uav ppm signals, CEUR Workshop Proceedings 3917 (2025) 226–236. doi:https://ceur-ws.org/Vol-3917/.
4. <https://docs.ultralytics.com/modes/train/#train-settings>

Олексій Коваль

Аспірант, кафедри комп'ютерних систем, мереж та кібербезпеки
Національний університет біоресурсів і природокористування України, Київ, Україна
0009-0004-1210-9852
o.koval@nubip.edu.ua

Ігор Болбот

д.т.н, проф., професор кафедри комп'ютерних систем, мереж та кібербезпеки
Національний університет біоресурсів і природокористування України, Київ, Україна
0000-0002-5708-6007
igor-bolbot@nubip.edu.ua

ГІБРИДНА БАЗА ДАНИХ ДЛЯ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ ОЦІНКИ ДЕГРАДОВАНИХ ҐРУНТІВ

Анотація. У роботі розглянуто проблему вибору оптимальної архітектури бази даних для системи підтримки прийняття рішень (DSS), призначеної для оцінки деградації ґрунтів, спричиненої військовими діями. Показано, що ефективність таких систем залежить від здатності бази даних інтегрувати геопросторові, експертні та інші дані, забезпечуючи при цьому масштабованість і стабільність роботи. На основі аналізу сучасних досліджень, обґрунтовано доцільність використання гібридної моделі бази даних, яка поєднує PostgreSQL/PostGIS для просторово-структурованих даних і MongoDB для неструктурованих мультимедійних джерел.

Ключові слова: система підтримки прийняття рішень; деградація ґрунтів; база даних; PostgreSQL, PostGIS; MongoDB

1. ВСТУП

Постановка проблеми. У сучасних умовах розробки систем підтримки прийняття рішень для агроекологічного моніторингу надзвичайно важливим є вибір та проектування бази даних, яка виступає основою для зберігання, аналізу та інтеграції різноманітних даних. Особливо це актуально для систем, що оцінюють деградацію ґрунтів унаслідок військових дій, де більшість джерел інформації є різнорідними — від супутникових знімків та сенсорних даних до експертних оцінок і геопросторових шарів[1,2]. Саме тому в цій роботі розглядається підбір типу бази даних, оптимальної для побудови DSS з оцінки деградації ґрунтів унаслідок військових дій, з урахуванням вимог до роботи з геопросторовими, експертними та неструктурованими даними та в обґрунтуванні вибору такої.

Аналіз останніх досліджень і публікацій. Згідно з результатами підбір бази даних та визначення її архітектури є одним з ключових етапів в розробці систем, особливу увагу потрібно приділити сумісності форматів даних для аналітичних моделей[3]. У дослідженні Manna P. et al. (2024) представлено геопросторову DSS, що використовує гібридну архітектуру бази даних (PostgreSQL/PostGIS та NoSQL) для моніторингу деградації ґрунтів. Автори показують, що правильна організація структури бази даних є ключовою для інтеграції супутникових, кліматичних та польових даних, а також забезпечує масштабованість і взаємодію з аналітичними модулями[4]

Мета публікації. Метою роботи є порівняльний аналіз типів баз даних і обґрунтування вибору оптимальної архітектури БД для системи підтримки прийняття рішень, що оцінює деградацію ґрунтів унаслідок військових дій, з урахуванням вимог до зберігання геопросторових, мультимедійних та експертних даних.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Під час розробки архітектури DSS були визначені основні вимоги до бази даних:

- підтримка геопросторових об'єктів (контури пошкоджених ділянок, координати вирв, карти деградації);
- зберігання неструктурованих даних (зображення, супутникові знімки, 3D-моделі рельєфу);
- можливість масштабування та багатокористувацького доступу;
- інтеграція з аналітичними модулями.

Для забезпечення цих вимог проведено порівняння найбільш поширених типів баз даних (табл. 1).

Таблиця 1

Порівняльна таблиця БД

Реляційні (PostgreSQL, MySQL)	Транзакційність, підтримка структурованих даних, SQL-аналітика	Складність роботи з мультимедійними файлами
Документні (MongoDB, CouchDB)	Гнучкість структури, простота зберігання JSON- документів	Відсутність складних транзакцій
Графові (Neo4j)	Зручна робота зі зв'язками між об'єктами	Висока складність інтеграції з ГІС
Геопросторові (PostGIS, Spatialite)	Потужна підтримка просторових операцій і форматів GeoTIFF, SHP	Високе навантаження на сервер при великих масивах raster-даних

Отримані результати засвідчили, що жоден із типів баз даних не може повністю задовольнити всі вимоги, тому оптимальним рішенням є використання гібридної архітектури:

- PostgreSQL/PostGIS основна реляційно-просторова база для структурованих і геопросторових даних;
- MongoDB — допоміжна база для зберігання неструктурованих супутникових зображень, результатів LiDAR-зйомок та експертних оцінок;

Запропонована архітектура забезпечує:

- Єдину модель даних для просторових та описових об'єктів
- Швидкий доступ і обробку геоданих за допомогою PostGIS-функцій;
- Гнучкість структури для оновлення та додавання нових типів даних без зміни основної схеми;
- Сумісність із аналітичними та візуалізаційними модулями DSS.

3.ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Розробка бази даних є критичним етапом створення DSS для оцінки деградації ґрунтів. Проведений аналіз показав, що гібридна модель БД, побудована на комбінації PostgreSQL/PostGIS та MongoDB, є найбільш доцільною для зберігання та обробки просторових, мультимедійних і якісних даних. Така архітектура відповідає сучасним вимогам масштабованості, інтеграції з ГІС і аналітичними інструментами. Отримані результати можуть бути використані під час розробки системи моніторингу деградації земель та планування заходів з екологічного відновлення територій, постраждалих від військових дій.

Подальші дослідження передбачають:

- розробку ETL-процесів синхронізації між компонентами системи;
- тестування продуктивності БД на реальних даних із територій, постраждалих від бойових дій;
- інтеграцію системи оцінки для прогнозування ступеня деградації земель.

ПОСИЛАННЯ

[1] J. Wang *та ін.*, "Remote sensing of soil degradation: Progress and perspective", *Int. Soil Water Conservation Res.*, безрез. 2023.
Доступно: <https://doi.org/10.1016/j.iswcr.2023.03.002>

[2] P. Stankovics *та ін.*, "A framework for co-designing decision-support systems for policy implementation: The LANDSUPPORT experience", *Land Degradation & Develop.*, січ. 2024. Доступно: <https://doi.org/10.1002/ldr.5030>

[3] M. Debeljak *та ін.*, "A Field-Scale Decision Support System for Assessment and Management of Soil Functions", *Frontiers Environmental Sci.*, т. 7, серп. 2019. Доступно: <https://doi.org/10.3389/fenvs.2019.00115>

[4] P. Manna *та ін.*, "A Geospatial Decision Support System for Supporting the Assessment of Land Degradation in Europe", *Land*, т. 13, № 1, с. 89, січ. 2024. Доступно: <https://doi.org/10.3390/land13010089>

Христина Дрогобицька

Здобувач вищої освіти

Хмельницький університет управління та права імені Леоніда Юзькова,

місто Хмельницький, Україна

0009-0008-3457-6864

drohobytska.k@gmail.com

МОДЕЛІ ВПРОВАДЖЕННЯ ЕЛЕКТРОННОГО ДОРАДНИЦТВА У СФЕРІ АГРОПРОМИСЛОВОГО КОМПЛЕКСУ УКРАЇНИ

Анотація: у статті розглянуто сучасні моделі впровадження електронного дорадництва у сфері агропромислового комплексу України як інструменту цифрової трансформації сільського господарства. Проаналізовано соціально-економічні та технологічні передумови створення електронних дорадчих систем, окреслено головні проблеми, ризики й умови ефективного функціонування таких сервісів. Визначено типові моделі організації електронного дорадництва та їх застосування в українських реаліях. Зроблено висновок, що формування сучасної системи електронного дорадництва є необхідною умовою підвищення конкурентоспроможності національного аграрного сектору, розвитку інновацій, збільшення продуктивності та сталого управління природними ресурсами.

Ключові слова: електронне дорадництво, цифровізація, агропромисловий комплекс, аграрна політика, дорадчі служби, інформаційні технології.

Аграрний сектор України є стратегічною галуззю національної економіки, забезпечуючи значну частку валютних надходжень і формуючи продовольчу безпеку держави. За останні десятиліття галузь переживає значні структурні та технологічні зміни, пов'язані з інтеграцією України у світові ринки, необхідністю підвищення ефективності виробництва, впровадженням інноваційних рішень. Одним із важливих напрямів таких змін є створення системи електронного дорадництва, яка забезпечує передачу знань, консультування та підтримку аграріїв засобами інформаційно-комунікаційних технологій.

Дорадництво, або аграрне консультування, традиційно розглядається як діяльність, спрямована на поширення науково-практичних знань серед сільгоспвиробників, фермерів, кооперативів, органів місцевого самоврядування. Його головною метою є підвищення продуктивності, ефективності управління, конкурентоспроможності та добробуту сільського населення. В умовах цифровізації економіки класичні моделі дорадництва зазнають суттєвих трансформацій. Електронне дорадництво передбачає інтеграцію цифрових інструментів — веб-платформ, мобільних застосунків, відеоконференцій, чат-ботів, систем дистанційного навчання, інтерактивних баз знань, що дозволяють здійснювати консультації дистанційно, оперативно та доступно для користувачів. В Україні інституційні засади дорадництва почали формуватись у 2000-х роках із прийняттям Закону України «Про сільськогосподарську дорадчу діяльність». Проте на практиці ця система залишається недостатньо розвиненою, оскільки більшість дорадчих служб не мають достатнього фінансування, кваліфікованих кадрів і цифрової інфраструктури.

Важливо розуміти, що електронне дорадництво — це не просто технічне оновлення старої системи, а створення нової моделі взаємодії між державними інституціями, науковими установами, дорадчими організаціями, бізнесом і кінцевими

користувачами. Цифрові технології створюють умови для інтеграції всіх учасників у єдину мережу знань, що сприяє прискоренню інноваційного розвитку аграрного виробництва. В Україні розвиток електронного дорадництва має як державний, так і регіональний виміри. На державному рівні існують окремі ініціативи зі створення єдиної електронної платформи, інтегрованої з реєстрами аграрної продукції, земельними кадастрами, освітніми ресурсами. Проте системного підходу поки немає. Більшість проектів реалізується за підтримки міжнародних партнерів — Продовольчої та сільськогосподарської організації ООН (FAO), Світового банку, ЄС, USAID. Вони фінансують створення цифрових платформ, аналітичних систем і мобільних сервісів, орієнтованих на малих і середніх фермерів [2].

На регіональному рівні впровадження електронного дорадництва здійснюється переважно дорадчими службами при університетах або громадських організаціях. Наприклад, Херсонський державний аграрний університет, Національний університет біоресурсів і природокористування, Львівський національний аграрний університет активно розвивають електронні навчальні курси, створюють бази знань та онлайн-платформи для фермерів. Такі ініціативи сприяють формуванню децентралізованої моделі дорадництва, орієнтованої на локальні потреби аграріїв.

Розвиток електронного дорадництва в Україні неможливий без глибокої цифровізації сільських територій. На жаль, значна частина сільських громад досі має обмежений доступ до швидкісного Інтернету. Це знижує ефективність дистанційних консультацій та електронного навчання. Програми державного рівня, як-от «Інтернет-субвенція», покликані частково вирішити цю проблему, проте процес іде повільно. Без достатньої цифрової інфраструктури електронне дорадництво не може стати масовим інструментом підтримки виробників.

З іншого боку, у світі зростає роль приватних технологічних компаній у розвитку аграрних ІКТ-рішень. Такі компанії, як CropX, EOS Data Analytics, OneSoil, створюють мобільні додатки, що дозволяють аграріям отримувати консультації на основі супутникових даних, штучного інтелекту та прогнозової аналітики. Використання подібних технологій в українських умовах сприятиме створенню гібридної моделі дорадництва, у якій держава забезпечує інституційні умови, а приватні компанії надають технологічні рішення.

Успіх електронного дорадництва залежить від наявності системи якісного контенту — наукових рекомендацій, практичних посібників, актуальних ринкових аналітик, нормативної інформації. Це потребує створення національної бази знань, до якої матимуть доступ усі учасники ринку. Такі бази повинні оновлюватися регулярно, мати пошукові механізми, інтерактивні інтерфейси та багатомовну підтримку. Дослідження свідчать, що використання ІКТ в агропромисловому комплексі України позитивно впливає на економічні результати діяльності підприємств. За даними Agricultural and Resource Economics (2023), підприємства, які впроваджують цифрові технології, демонструють вищу додану вартість, продуктивність праці та обсяги реалізації продукції [4]. Це доводить, що розвиток електронного дорадництва не є самоціллю, а безпосередньо впливає на економічну ефективність.

Водночас існує низка ризиків і бар'єрів. По-перше, відсутність державної стратегії цифрового дорадництва створює фрагментарність системи. Кожна область або установа діє на власний розсуд, без єдиних стандартів і координації. По-друге, низький рівень цифрової грамотності фермерів, особливо старшого покоління, ускладнює використання навіть базових онлайн-інструментів. По-третє, спостерігається дефіцит кваліфікованих дорадників, здатних працювати в цифровому середовищі. Крім того, відсутній механізм оцінки ефективності дорадчих послуг і зворотного зв'язку між дорадниками та користувачами.

Успішний розвиток електронного дорадництва вимагає комплексного підходу. Передусім потрібно забезпечити стабільну державну підтримку — створити нормативно-правові засади для цифрових дорадчих послуг, закріпити порядок їх сертифікації та фінансування. Необхідно також стимулювати державно-приватні партнерства, залучати ІТ-компанії до розробки програмного забезпечення, платформ, інтерфейсів. Важливим напрямом є підготовка кадрів: слід запровадити навчальні програми з цифрового дорадництва в аграрних університетах, а також проводити постійне підвищення кваліфікації для чинних дорадників.

Значну роль у розвитку дорадництва відіграють міжнародні донори. Так, проекти FAO та ЄБРР у 2022–2024 рр. спрямовані на створення систем дистанційного навчання для фермерів, використання електронних платформ для консультацій із питань агротехнологій, фінансів, екології [3; с.9]. Розширення таких ініціатив дає змогу Україні перейняти міжнародний досвід, зокрема моделі електронного дорадництва, що успішно функціонують у країнах ЄС, США та Канаді. У європейській практиці електронне дорадництво має сталу традицію. Наприклад, система Farm Advisory System у ЄС інтегрована в політику спільного сільського господарства, забезпечуючи фермерам доступ до консультацій із правових, фінансових, екологічних та технологічних питань. У Нідерландах функціонує мережа Wageningen University & Research, яка підтримує цифрові дорадчі сервіси для фермерів і стартапів. В Україні подібна модель може стати орієнтиром для створення єдиної системи з державною координацією та регіональними центрами впровадження.

Перспективи розвитку електронного дорадництва безпосередньо пов'язані з реалізацією Стратегії розвитку агропромислового комплексу та сільських територій України до 2030 року [1]. У документі серед пріоритетів визначено цифровізацію управління, інтеграцію науки, освіти і консультацій, розвиток кадрового потенціалу. Саме електронне дорадництво може стати тим інструментом, який з'єднає ці напрями в єдину екосистему інновацій. Ефективне функціонування системи можливе лише за умови постійного моніторингу результативності. Потрібно запровадити показники ефективності — охоплення користувачів, економічну віддачу, рівень задоволеності фермерів, вплив на сталий розвиток сільських територій. На підставі таких показників варто формувати державні програми підтримки та оновлення цифрових платформ.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Про схвалення Стратегії розвитку сільського господарства та сільських територій в Україні на період до 2030 року: Розпорядж. Каб. Міністрів України від 15.11.2024 № 1163-р : станом на 20 серп. 2025 р.
URL: <https://zakon.rada.gov.ua/laws/show/1163-2024-p#Text>.
2. У Мінагрополітики обговорили розробку ефективної моделі дорадництва з урахуванням міжнародного досвіду. URL: <https://www.kmu.gov.ua/news/uminagropolitiki-obgovorili-rozrobku-efektivnoyi-modeli-doradnictva-z-urahuvannyam-mizhnarodnogo-dosvidu> (дата звернення: 08.11.2025).
3. Тарас Гагалюк. ФІНАНСУВАННЯ СТІЙКОСТІ І ВІДБУДОВИ АГРАРНОГО СЕКТОРА УКРАЇНИ МІЖНАРОДНИМИ ФІНАНСОВИМИ УСТАНОВАМИ ТА АГЕНЦІЯМИ РОЗВИТКУ. URL: https://www.apd-ukraine.de/fileadmin/user_upload/APD_Taras_Gagalyuk_UA_01.pdf (дата звернення: 08.11.2025).
4. V. Hrosul, O. Kruhlova, A. Kolesnyk, "Digitalization of the agricultural sector: the impact of ICT on enterprise development in Ukraine," *Agricultural and Resource Economics*, vol. 9, no. 4, 2023.

Ярослав Понзель

Асистент кафедри інформаційних систем і технологій

ВИЗНАЧЕННЯ ФУНКЦІОНАЛЬНИХ МОДУЛІВ ТА АРХІТЕКТУРНІ РІШЕННЯ ДЛЯ ІНТЕГРАЦІЇ МОБІЛЬНОГО ЗАСТОСУНКУ В ЕКОСИСТЕМУ УНІВЕРСИТЕТУ

Після обґрунтування стратегічної необхідності розробки мобільного застосунку для університету [1] (як доповнення до існуючих інформаційних систем), наступним критичним етапом є детальне проектування його функціоналу та архітектури інтеграції. Успіх застосунку залежить не від кількості функцій, а від їх цінності для кінцевого користувача (студента, викладача) та безшовності інтеграції з існуючими даними (розклад, успішність, навчальні матеріали). Помилки на етапі проектування архітектури можуть призвести до створення ізольованої, неефективної системи, що дублює, а не розширює функціонал.

Для забезпечення гнучкості та можливості поетапного впровадження, функціонал застосунку доцільно розділити на незалежні, але взаємопов'язані модулі [2]:

- Розклад – автоматична синхронізація, фільтрація за групою, push-сповіщення про зміни або скасування занять.

- Академічна успішність – відображення підсумкових оцінок з предметів, які вже склав студент.

- Індивідуальний навчальний план – відображення навчальної траєкторії студента із інформацією про кредити, години і т.д.

- Вибір дисциплін [3] – модуль, який відображає доступні дисципліни до вибору й надає можливість підібрати їх в свій навчальний план під власні потреби чи уподобання студента.

- Портфоліо – редагування особистих та навчальних даних для формування особистої картки студента.

- Безпека – модуль, який відповідає за інформування студентів про повітряні тривоги або інші небезпеки, та містить дані про укриття, їх фотографії, кількість допустимих місць та можливий маршрут до них.

- Гуртожиток – дані про проживання в гуртожитку, заявки на поселення, дані про перепустку.

- Фінанси – контроль за оплатою навчання та інших сервісів, історія платежів, дані про заборгованість

- Цифровий деканат – можливість онлайн-замовлення довідок (про навчання, для військкомату) з відстеженням статусу.

- Інтерактивна карта – навігація корпусами та територією гуртожитків, пошук аудиторій, кафедр, бібліотек, їдалень та інших необхідних місць.

- Зворотній зв'язок – модуль для комунікації між студентом та адміністрацією університету або ж кампусу для отримання відгуків, коментарів, скарг та зворотнього зв'язку студентів.

- Месенджер – модуль підтримки комунікації між студентами та викладачами для отримання консультацій щодо дисциплін.

Університетська інфраструктура на даний момент складається із наступних цифрових сервісів (рис. 1.):

- «Система електронного деканату», що виконує функції управління освітнім процесом, ведення даних студентів, їх академічної успішності та індивідуальних планів.

- Веб-портал користувача: «Особистий кабінет студента», що слугує точкою доступу до сервісів та комунікації.
- Система управління навчанням (LMS): «Навчально-інформаційний портал» на базі платформи Moodle.
- Хмарна платформа: «Google Workspace», що забезпечує корпоративну пошту, хмарні сховища та інструменти спільної роботи.
- Спеціалізовані сервіси: «Електронний розклад» (окрема система візуалізації) та «Чат-бот першокурсника» (Telegram-бот).

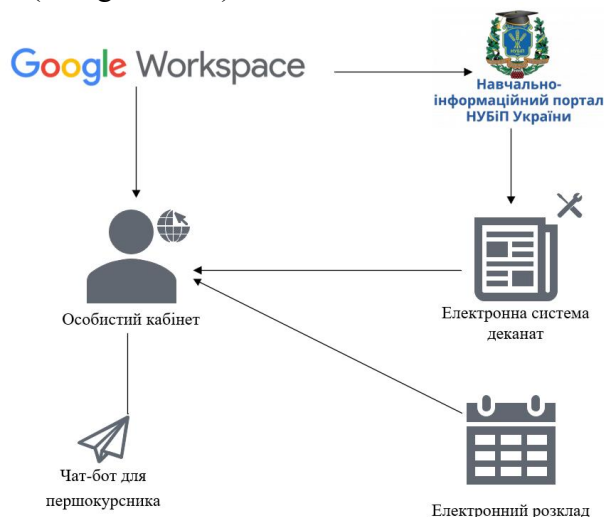


Рис. 1. Схема роботи цифрових сервісів університету

Ключовим аспектом архітектури інтеграції є використання існуючої сервісної логіки «Особистого кабінету студента». Даний веб-портал вже функціонує як проміжний шар, що надає стандартизовані API-ендпоінти для доступу до даних «Системи електронного деканату». Це суттєво спрощує розробку мобільного застосунку, оскільки основні бізнес-логічні модулі (розклад, академічна успішність, індивідуальний навчальний план, вибір дисциплін, портфолію) вже мають готові точки доступу. Таким чином, задача мобільної розробки зводиться до реалізації нативного користувацького інтерфейсу (UI) та взаємодії з цими API. За аналогічною моделлю (створення API-ендпоінту на боці веб-порталу та UI-компоненту в мобільному застосунку) планується інтеграція всіх інших сервісів (фінанси, гуртожиток та інші), що забезпечує єдину, контрольовану та безпечну архітектуру обміну даними.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Понзель Я.Ю. Архітектура вебсайту з інтегрованою системою для вибору дисциплін. Наука і техніка сьогодні. 2025. Вип. 1 (42). С. 1331–1343. URL: <http://perspectives.pp.ua/index.php/nts/article/view/19336/19345>
2. Понзель Я. Ю., Голуб Б.Л. Формування основного функціоналу для реалізації інформаційної системи в фері комунікації студента та університету. XIII Міжнародна науково-практична конференція молодих вчених. «Інформаційні технології: економіка, техніка, освіта '2022». 2022. С. 90-91. URL: [https://www.researchgate.net/publication/366617986_Antifisingova_infrastruktura_ak_zasib_protidii_zagroзам_funkcionuvanna_sistem_kriticnogo_priznacenna#pf5b](https://www.researchgate.net/publication/366617986_Antifisingova_infrastruktura_ak_zasib_protidii_zagroزام_funkcionuvanna_sistem_kriticnogo_priznacenna#pf5b)
3. Hlazunova O., Ponzel Y. Technologies and algorithms for the implementation of the recommendation system for creating an individual study plan for a higher education student. URL: <https://ceur-ws.org/Vol-3771/paper03.pdf>

Владислав Шаптала

науковий керівник Хиленко В.В.

ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ МАКРОЕКОНОМІЧНИХ ПОКАЗНИКІВ У ВИЗНАЧЕННІ ТЕНДЕНЦІЙ РОЗВИТКУ ЕКОНОМІК КРАЇН G7 ЗА ПЕРІОД 1999– 2023 РР.

Анотація. У роботі здійснено порівняння та аналіз окремих чинників, що впливають на стан економічної і фінансової систем країн G7. Дослідження охоплює період 1999–2023 років і враховує динаміку грошових потоків, зокрема внутрішніх і зовнішніх інвестицій та емісії. Проаналізовано співвідношення між індикативними оцінками та фактичними показниками. Як кількісний показник для оцінки економічного стану використано індикатор Баффета, ефективність якого перевірено у порівнянні з іншими статистичними даними.

Актуальність дослідження. У 1999–2023 рр. країни G7 тричі пережили суттєві економічні спади — dot-com (2000), глобальну кризу (2008) та пандемічний шок (2020), що призвели до значних втрат потенційного ВВП [1]. Досвід засвідчує, що своєчасне розпізнавання урядами та центробанками ознак перегріву економіки дозволяє пом'якшити подальше падіння [2]. Порівняльний аналіз кризових періодів показав, що головні капітальні потоки концентрувалися у різних секторах, а перевищення індикатором Баффета порогу у 100% часто було раннім сигналом перегріву й майбутньої корекції.

Мета дослідження — перевірити, чи можуть окремі статистичні індикатори (індикатор Баффета, показники руху фінансових потоків, емісія тощо) виступати надійними передвісниками фінансових криз.

Джерельна база. Використано дані міжнародних статистичних платформ та академічних публікацій. Основні індикатори для розрахунку «індикатора Баффета» отримано з World Development Indicators Світового банку — ринкова капіталізація лістингових компаній у відсотках до ВВП [1], а також потоки FDI/ODI [3] і показники грошової маси M3 [4].

Прогнозування стану економіки країн G7 доцільно здійснювати на основі низки параметрів — динаміки внутрішніх і зовнішніх потоків капіталу, обсягів емісії та індикатора Баффета.

Індикатор Баффета (Buffett Indicator) визначається як співвідношення сукупної ринкової капіталізації акцій до номінального ВВП країни й відображає «дороговизну» фондового ринку відносно економіки [1]. Також у дослідженні враховувалися показники FDI (вхідні інвестиції), ODI (закордонні інвестиції) [3] і M&A (злиття та поглинання компаній).

Країна потрапляє у зону ризику, коли індикатор Баффета перевищує 150% (нормою вважається 70–90%), спостерігається стрімке зростання ODI flows/M&A (ознака секторного перегріву) чи різкий спад FDI та нестабільність динаміки M3.

Під час кризи dot-com індикатор Баффета (MCap/GDP) продемонстрував себе як ранній попереджувальний сигнал: у США в 1999 році він досяг 153,4% при історичній нормі 60–80%, після чого відбулося різке падіння NASDAQ і S&P 500 [1]. Потоки FDI досягли максимуму в 1999–2000 рр. і скоротилися лише у 2001–2003, отже, цей показник був скоріше запізним. Натомість ODI та M&A подали випереджальний сигнал перегріву сектору: угоди на кшталт Vodafone–Mannesmann і дороги 3G-аукціони (Велика Британія £22,5 млрд; Німеччина €50,8 млрд) відображали пік ризикових інвестицій. Емісія M3 не передбачила кризи: її темпи у єврозоні залишалися

стабільними (5–6%) і прискорилися лише після шоку, у США приріст МЗ складав близько +8,7% р/р. Висновок для dot-com: Баффет — ранній сигнал; ODI/M&A — частково; FDI — ні; МЗ — ні.

Рік	1999	2000	2001	2002	2003
США	153.4	147.4	132.1	101.1	124.5

(1)

індикатор Баффета для США, 1999 - 2003

Рік	1999	2000	2001	2002	2003
США	283,376	314,007	159,461	74,457	53,146

(2)

FDI flows для США, 1999 - 2003

Рік	1999	2000	2001	2002	2003
США	209,391	142,626	124,873	134,946	129,352

(3)

ODI flows для США, 1999 - 2003

Рік	1999	2000	2001	2002	2003
США	4.638	4.925	5.433	5.772	6.067

(4)

Емісія МЗ для США, 1999 - 2003

Напередодні глобальної фінансової кризи 2008–2009 рр. індикатор Баффета знову залишався високим (130–140% у 2007 р.), що свідчило про переоціненість ринку, однак не відображало кредитної вразливості системи [1]. FDI досяг піку у 2007 р., але різко знизився у 2008–2009, виступаючи реактивним індикатором. Водночас ODI та M&A подали сигнал про вершину кредитного циклу — пік угод із залученням левереджу (LBO) та масштабних банківських злиттів у 2006–2007 рр. передував різкому скороченню активності. Динаміка МЗ та кредитів зросла на завершенні буму, але у 2008–2009 рр. різко скоротилася, що дозволило фіксувати сам момент кризи, хоча не попередити її. Висновок для GFC: Баффет — частково; ODI/M&A — частково; FDI — ні; МЗ — синхронний індикатор, не ранній.

Рік	2007	2008	2009
США	137.6	78.5	104.1

(5)

індикатор Баффета для США, 2007 - 2009

Рік	2007	2008	2009
США	215,952	306,366	143,604

(6)

FDI flows для США, 2007 - 2009

Рік	2007	2008	2009
США	393,518	308,296	287,901

(7)

ODI flows для США, 2007 - 2009

Рік	2007	2008	2009
США	7.471	8.192	8.496

(8)

Емісія МЗ для США, 2007 - 2009

Отже, для криз акційного типу найінформативнішим є поєднання індикатора Баффета та активності ODI/M&A у перегрітих секторах, що дозволяє вчасно виявляти надмірну капіталізацію ринку та спекулятивні настрої інвесторів. Така комбінація показників добре відображає фазу накопичення фінансових ризиків, коли зростання угод зі злиттів і поглинань супроводжується переоцінкою активів. Натомість кредитні кризи краще описує композит із ODI/M&A, показників грошової маси (МЗ) та динаміки кредитування, адже саме він дозволяє виявляти надмірне розширення боргового фінансування та втрату ліквідності в банківському секторі. Потoki FDI, своєю чергою, здебільшого мають реактивний характер і відображають уже сформовані тенденції на ринку, тому не можуть вважатися ефективним раннім сигналом для виявлення наближення фінансової кризи.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. databank.worldbank.org [Електронний ресурс]: Market capitalization of listed domestic companies <https://databank.worldbank.org/metadataglossary/world-development-indicators/series/CM.MKT.LCAP.GD.ZS> (дата звернення: 18.07.2025)
2. scopus.com [Електронний ресурс]: Khilenko V. Mathematical Modeling of the Effect of “Splashing out” and Optimization of Management of Banking and Economic Systems Under Globalization Conditions. Cybernetics and Systems Analysis Open, 2018, 54(3), pp. 376–384. (дата звернення: 28.09.2025)
<https://www.scopus.com/authid/detail.uri?authorId=6603387761>.
3. unctadstat.untad.org [Електронний ресурс] “Foreign direct investment”: <https://unctadstat.untad.org/datacentre/dataviewer/US.FdiFlowsStock> (дата звернення: 12.09.2025)
4. fred.stlouisfed.org [Електронний ресурс] “M3 for United States”: <https://fred.stlouisfed.org/series/MABMM301USM189S> (дата звернення: 12.09.2025)

Вадим Олійник

аспірант 2-го року навчання

кафедра інформаційних систем і технологій, НУБіП, м. Київ, Україна

v.o.oliynyk@nubip.edu.ua

Тетяна Волошина

кандидат педагогічних наук, доцент

кафедра інформаційних систем і технологій, НУБіП, м. Київ, Україна

ORCID ID: <https://orcid.org/0000-0001-6020-5233>

t-voloshina@nubip.edu.ua

ОГЛЯД МОДЕЛЕЙ ПРОГНОЗУВАННЯ РЕЗУЛЬТАТІВ НАВЧАННЯ СТУДЕНТІВ У ЗАКЛАДАХ ВИЩОЇ ОСВІТИ НА ОСНОВІ МЕТОДІВ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. У роботі розглянуто сучасні підходи та моделі прогнозування результатів навчання студентів у закладах вищої освіти. Проведено аналіз тенденцій розвитку інтелектуальних систем освітньої аналітики, орієнтованих на використання машинного та глибокого навчання для оцінювання академічної успішності. Особливу увагу приділено аспектам інтерпретованості результатів та проблемам якості даних. Отримані результати дозволяють сформулювати теоретичне підґрунтя для подальшої розробки інформаційної технології прогнозування успішного завершення навчання студентів у закладах вищої освіти.

Ключові слова: штучний інтелект, машинне навчання, глибоке навчання, прогнозування результатів навчання, освітня аналітика.

1. ВСТУП

Системи освітньої аналітики є одним із ключових рушіїв цифрової трансформації закладів вищої освіти (ЗВО). Широке впровадження систем управління навчанням сприяло формуванню великих обсягів освітніх даних, що відображають різні аспекти освітньої діяльності студентів. Такі дані охоплюють показники відвідування навчальних занять, результати тестування, своєчасність виконання завдань, активність у форумах тощо. Завдяки розвитку методів штучного інтелекту (ШІ) та машинного навчання (МН) виникла можливість не лише аналізувати такі дані, а й прогнозувати майбутні результати навчання.

Постановка проблеми. Сучасна система вищої освіти стикається з низкою викликів, зокрема з необхідністю визначення академічної успішності студентів. Традиційні статистичні методи прогнозування часто не враховують складних нелінійних взаємозв'язків між численними факторами, що впливають на результати навчання. Водночас методи ШІ демонструють значний потенціал у побудові точних і гнучких моделей прогнозування результатів навчання студентів. Таким чином, виникає потреба у проведенні системного огляду моделей прогнозування у сфері вищої освіти, спрямованого на узагальнення існуючих підходів, визначення їхніх переваг та обмежень, а також перспектив подальшого розвитку.

Аналіз останніх досліджень і публікацій. Сучасні дослідження у сфері освітньої аналітики засвідчують активне зростання інтересу до створення прогнозних моделей навчальних результатів. У роботі [1] запропоновано гібридну модель CNN-LSTM для прогнозування успішності студентів на основі часових освітніх даних, що дозволило підвищити точність класифікації у порівнянні з класичними моделями МН. Дослідження [2] продемонструвало ефективність ансамблевих методів (Random Forest, CatBoost, XGBoost) у прогнозуванні академічних результатів студентів на основі даних з систем управління навчанням, водночас підкреслюючи проблему перенавчання (overfitting) при малих вибірках. Автори [3] вивчили підхід Explainable AI (SHAP, LIME), який забезпечує інтерпретованість результатів прогнозування та сприяє

прийняттю прозорих управлінських рішень у закладах освіти. Р. Абугалала та ін. [4] дослідили ефективність AutoML-підходів у побудові моделей для прогнозування успішності, що дозволило автоматизувати процес підбору алгоритмів і параметрів. У роботі [5] наведено системний огляд використання моделей машинного навчання та генеративного штучного інтелекту у освітньому контексті. Попри значну кількість досліджень, відкритими залишаються питання інтерпретованості, узагальнюваності моделей та інтеграції у реальні освітні інформаційні системи.

Метою роботи є систематизація сучасних підходів і моделей прогнозування результатів навчання студентів у ЗВО з урахуванням їхніх можливостей, обмежень і потенційних напрямів розвитку.

2. ТЕОРЕТИЧНІ ОСНОВИ

Моделі прогнозування результатів навчання ґрунтуються на аналізі залежностей між сукупністю характеристик студента (вхідних ознак) та його академічними результатами (вихідними змінними). Залежно від обраного підходу, виділяють три основні групи моделей: *статистичні* (лінійна, логістична регресія, дискримінантний аналіз) - забезпечують інтерпретованість, проте не враховують складні нелінійні зв'язки; *моделі машинного навчання* (дерева рішень, Random Forest, SVM, ансамблеві методи) - забезпечують вищу точність прогнозів і можливість роботи з великою кількістю параметрів; *моделі глибокого навчання* (нейронні мережі, CNN, LSTM) - дозволяють аналізувати часові та неструктуровані дані, проте потребують великих обсягів навчальних вибірок і обчислювальних ресурсів.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Аналіз сучасних досліджень дозволяє класифікувати моделі прогнозування результатів навчання студентів у ЗВО за кількома ознаками. За типом вхідних даних моделі ґрунтуються на основі демографічних характеристик студентів, академічних показників, поведінкових даних з систем управління навчанням (логи активності), а також соціальних взаємодій у цифровому навчальному середовищі. З огляду на алгоритмічну реалізацію виділяють регресійні, класифікаційні, гібридні та ансамблеві моделі. Останні демонструють підвищену точність і стабільність результатів завдяки поєднанню кількох методів машинного навчання. Зокрема, ансамблеві підходи, такі як XGBoost і CatBoost, ефективно працюють із пропущеними або неточними даними, тоді як гібридні архітектури на основі CNN-LSTM здатні виявляти часові закономірності у поведінкових даних на навчальних платформах. Це забезпечує більш точне моделювання навчальних процесів із урахуванням індивідуальних особливостей студентів.

Попри значні досягнення у сфері застосування методів ШІ в освітній аналітиці, залишається низка проблем, що обмежують широке впровадження таких моделей у практику ЗВО. Одним із ключових викликів залишається недостатність якісних та репрезентативних даних: у більшості ЗВО відсутні уніфіковані системи збору та стандартизації освітніх даних, що ускладнює порівняння отриманих результатів. Іншою суттєвою проблемою є перенавчання моделей, особливо при використанні невеликих або незбалансованих вибірок, що призводить до зниження узагальнюваності результатів. Особливої уваги потребує питання інтерпретованості - моделі глибокого навчання (ГН) попри високу точність, часто функціонують як «чорні скриньки», що знижує довіру до них з боку викладачів та адміністрації. Крім того, важливо враховувати етичні та правові аспекти використання освітніх даних: захист персональної інформації студентів та забезпечення прозорості алгоритмів залишаються актуальними завданнями. Саме ці виклики визначають перспективність інтеграції методів Explainable Artificial Intelligence (XAI) у процес побудови моделей

прогнозування результатів навчання. Підходи, засновані на алгоритмах SHAP і LIME, дозволяють пояснювати вплив окремих ознак на підсумковий прогноз, роблячи моделі більш прозорими та зрозумілими для користувачів. Такі інструменти не лише підвищують довіру до систем на основі штучного інтелекту, але й сприяють ухваленню більш обґрунтованих рішень у навчальному процесі. Використання ХАІ у поєднанні з автоматизованими платформами машинного навчання (AutoML) відкриває можливість побудови адаптивних систем прогнозування, здатних самостійно обирати оптимальні алгоритми та параметри моделі.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Проведений огляд свідчить, що сучасні підходи до прогнозування результатів навчання студентів ЗВО активно розвиваються у напрямі поєднання методів МН та ГН. Найвищу ефективність демонструють ансамблеві та гібридні моделі, що здатні працювати з багатовимірними освітніми даними. Перспективними напрямом подальших досліджень є розробка інтегрованої інформаційної системи освітньої аналітики, підвищення інтерпретованості моделей шляхом використання ХАІ, а також створення адаптивної прогнозовної моделі успішного завершення навчання, яка може стати основою для інформаційної технології підтримки прийняття рішень у ЗВО.

ПОСИЛАННЯ

- [1] K. O. Adefemi and M. B. Mutanga, "A Robust Hybrid CNN–LSTM Model for Predicting Student Academic Performance," *Digital*, vol. 5, no. 2, p. 16, May 2025, doi: <https://doi.org/10.3390/digital5020016>.
- [2] W. Ahmed, M. A. Wani, Pawel Plawiak, Souham Meshoul, A. Mahmoud, and M. Hammad, "Machine learning-based academic performance prediction with explainability for enhanced decision-making in educational institutions," *Scientific Reports*, vol. 15, no. 1, Jul. 2025, doi: <https://doi.org/10.1038/s41598-025-12353-4>.
- [3] F. T. Johora, M. N. Hasan, A. Rajbongshi, M. Ashrafuzzaman, and F. Akter, "An explainable AI-based approach for predicting undergraduate students academic performance," *Array*, vol. 26, p. 100384, Jul. 2025, doi: <https://doi.org/10.1016/j.array.2025.100384>.
- [4] R. A. Abougalala, N. Alharbi, M. A. Amasha, M. F. Areed, S. Alkhalaf, and D. Khairy, "Predicting student performance academic using Automated Machine Learning (AutoML): in medical academic institutions," *Journal of New Approaches in Educational Research*, vol. 14, no. 1, Aug. 2025, doi: <https://doi.org/10.1007/s44322-025-00038-9>.
- [5] M. Á. Rodríguez-Ortiz, P. C. Santana-Mancilla, and L. E. Anido-Rifón, "Machine Learning and Generative AI in Learning Analytics for Higher Education: A Systematic Review of Models, Trends, and Challenges," *Applied Sciences*, vol. 15, no. 15, p. 8679, Aug. 2025, doi: <https://doi.org/10.3390/app15158679>

Михайло Садко

К.е.н., доцент, доцент

Місце роботи: факультет інформаційних технологій НУБіП України.

ORCID ID: sadko@nubip.edu.ua

ВИКОРИСТАННЯ НЕЙРОННИХ МЕРЕЖ В МОДЕЛЮВАННІ УПРАВЛІННЯ АГРАРНИМИ ПІДПРИЄМСТВАМИ

Анотація. Використання штучного інтелекту в моделюванні управління аграрними підприємствами для виявлення нових підходів у підвищенні ефективності сільського господарства. Розглянуті основні складові нейронних мереж, їх використання для вирішення задач класифікації, регресії, прогнозування часових рядів. Розглянуто приклад застосування Багатошарового перцептрона SPSS, який передбачає можливість проведення обробки та аналізу великих обсягів даних, введення нелінійних залежностей, виявлення складних закономірностей та визначати тенденції розвитку сільськогосподарських підприємств регіону.

Ключові слова: штучний інтелект; нейронні мережі; нейронні мережі в SPSS; багатошаровий перцептрон.

Штучний інтелект (ШІ) все глибше проникає в різні сфери життя людей, і сільське господарство не є винятком. Застосування ШІ в моделюванні управління аграрними підприємствами відкриває нові можливості для раціонального використання виробничого потенціалу підприємств, визначення їх оптимальної структури, зниження трудових витрат в розрахунку на одиницю отриманої продукції та підвищення ефективності виробництва сільськогосподарської продукції.

Однією з складових сучасного штучного інтелекту (ШІ) є нейронні мережі, які використовуються для вирішення широкого спектра складних завдань, що імітують або дозволяють покращувати сприймання, обробку та використання людиною інформації.

Нейронні мережі (НМ) спроектовані за принципом роботи біологічного мозку і складаються із взаємопов'язаних вузлів (штучних нейронів), які обробляють інформацію. Вони здатні навчатися на великих обсягах даних, виявляти закономірності та робити прогнози або приймати рішення без явного програмування для кожного конкретного випадку.

Нейронні мережі в ШІ можуть використовуватись для :- класифікації даних: визначення категорії, до якої належить об'єкт чи подія;- регресії: розрахунок числових значень або прогнозних даних;- розпізнавання образів: ідентифікація об'єктів, облич, мовлення, текстів на зображеннях, відео та аудіо;- генерації контенту: створення нового тексту, зображень, музики;- управління та контролю: керування складними системами (наприклад, автономні транспортні засоби, промислові процеси).

Більшість цих можливостей реалізовані в нейронних мережах IBM SPSS Statistics, потужному інструменті, який успішно використовується для аналізу даних, особливо їх нелінійної залежності та великих обсягів даних, виявлені складних закономірностей та визначення перспектив розвитку сільськогосподарських підприємств.

Основні можливості застосування нейронних мереж, реалізованих в SPSS:

- різноманітність моделей та функцій активації, що дозволяє оптимізувати модель для конкретної задачі;
- візуалізація результатів: SPSS надає інструменти для візуалізації архітектури нейронної мережі, ваг нейронів та інших характеристик моделі;
- здатність надавати високу точність прогнозів, виявляти складні закономірності в даних, які можуть бути недоступні для традиційних статистичних методів.

SPSS підтримує два основних типи нейронних мереж:

- **Багатошаровий перцептрон (MLP)**: найпоширеніший тип нейронної мережі, який складається з кількох шарів нейронів, з'єднаних між собою вагами, з допомогою яких можна виявляти складні нелінійні залежності в даних і виконувати різноманітні завдання, такі як класифікація, регресія та прогнозування.

- **Радіальні базисні функції (RBF)**: особливий вид штучних нейронних мереж, який використовує радіально-симетричні функції як функції активації нейронів у прихованому шарі. Ці функції мають властивість максимального значення в центрі і спадання до нуля при віддаленні від цього центру. RBF добре підходить для задач інтерполяції та апроксимації.

Багатошаровий перцептрон SPSS є зручним інструментом, який використовується для задач: - **класифікації**: визначення класу, до якого належить об'єкт; - **регресії**: прогнозування числових значень;

- **розпізнавання образів**: класифікації зображень, звуків та інших типів даних; - **прогнозування часових рядів**: передбачення майбутніх значень на основі історичних даних.

Процес застосування **Багатошарового перцептрон**у SPSS включає такі кроки:

1. Підготовка даних: - виділення вхідних (незалежних) і вихідних (залежних) змінних; - перекодування категоріальних змінних у числовий формат;

- нормалізація даних для покращення швидкості навчання. **2. Створення моделі**: вибір кількості прихованих шарів і нейронів у кожному шарі. Шар — це структурний компонент нейронної мережі, який групує нейрони (вузли) і виконує певний тип перетворення над даними, що проходять через нього. Приховані шари виконують низку критично важливих завдань, що дозволяють нейронній мережі вирішувати складні проблеми: • **Вивчення ієрархічних ознак**: кожен прихований шар навчається виділяти більш абстрактні та складні ознаки (характеристики) на основі ознак, виділених попереднім шаром. Наприклад, у задачі розпізнавання зображень, перший прихований шар може вивчати прості ознаки (лінії, краї), наступний — комбінувати їх у форми (носи, очі), а далі — у цілі об'єкти (обличчя, тварини); • **Введення нелінійності**: За допомогою функцій активації (таких як ReLU, сигмоїда чи гіперболічний тангенс), приховані шари вводять нелінійність у модель. Це критично важливо, оскільки без нелінійності вся нейронна мережа, незалежно від кількості шарів, була б еквівалентна простій лінійній моделі, яка не може вирішувати складні завдання реального світу; •

Перетворення даних: Вони трансформують вхідні дані у нове представлення, яке є більш інформативним і легким для використання вихідним шаром для фінального прогнозу або класифікації; • **Збільшення складності**: чим більше прихованих шарів (що називається глибоким навчанням), тим складніші зв'язки та закономірності може

вивчити нейронна мережа.

Кожен шар має: - **Набір нейронів**: Ці нейрони отримують вхідні дані (або вихідні дані з попереднього шару), для яких застосовують математичну функцію, яка визначає вихідний сигнал штучного нейрона на основі його зваженої суми входів та зсуву, а також вибирають алгоритм навчання (наприклад, зворотного поширення помилки); -

Навчальні параметри: Це ваги (W) і зміщення (b), які є змінними, що налаштовуються (навчаються) в процесі тренування мережі; - **Вихідні дані**: Результат обробки даних шаром стає входом для наступного шару.

3. Виконання (навчання) моделі: - Розбиття даних на навчальну і тестову вибірки;

- Навчання мережі на навчальній вибірці: - Оцінка точності моделі на тестовій вибірці.

4. Інтерпретації та аналіз результатів: - Аналіз важливості вхідних змінних; - Візуалізація роботи мережі.

Підготовка даних нейронних мереж в SPSS вимагає проведення кількох кроків:

- **забезпечення чистоти даних**: ліквідація пропущених значень та тих спостережень,

які суттєво відрізняються від більшості інших значень вибірки. Найкращий варіант перевірка сукупності значень кожної змінної на відповідність критеріям нормального закону розподілення; - **коректне кодування та масштабування даних**; - **розділення Даних** на три незалежні набори: - **Навчальний**: використовується для побудови моделі та обчислення ваг; - **Тестовий**: використовується для відстеження процесу навчання, запобігаючи перенавчанню; - **Відкладений**: фінальна оцінки продуктивності моделі. Процес виконання **Багатошарового Перцептронну**, який є класичною архітектурою нейронної мережі прямого поширення, полягає у послідовній обробці вхідних даних через низку **прихованих шарів** для отримання кінцевого результату. Він включає два основні етапи: **пряме поширення сигналу** (формування прогнозу) та **зворотне поширення помилки** (навчання). Практичне використання нейронних мереж SPSS розглянемо на прикладі сільськогосподарських підприємств Київської області, які займаються вирощуванням кукурудзи на зерно. Визначимо вплив факторів: X1- площа посіву кукурудзи на зерно, X2- виробничі витрати в розрахунку на 1 га площі посіву, X3- питома вага грошових надходжень від реалізації кукурудзи до загальної суми надходжень від реалізації продукції рослинництва. Залежна змінна: Y- урожайність кукурудзи на зерно. Результати виконання нейронної мережі показані на рисунках 1-5. В результаті отримані прогнозні значення урожайності на перспективу та важливість впливу факторів на залежну змінну.

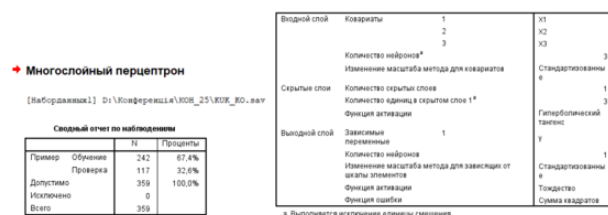


Рисунок 1 Використання нейронної мережі. Рисунок 2. Інформація про мережу.

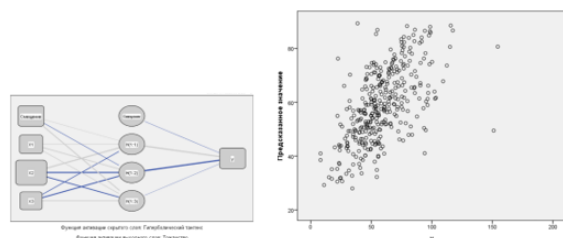


Рисунок 3. Функции активации шарів.

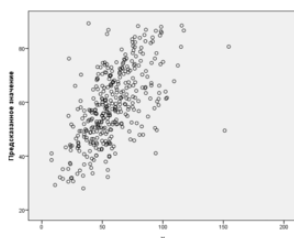


Рисунок 4. Прогнозные значения урожайности.

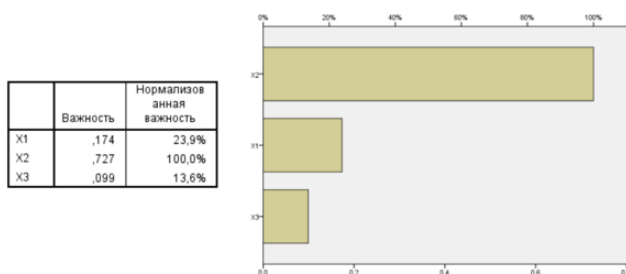


Рисунок 5 Важливість незалежних змінних.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. «IBM SPSS Statistics і Amos: Професійний статистичний аналіз даних.
2. Офіційна документація IBM SPSS: Вбудована довідка програми та офіційні посібники на сайті IBM.

Костянтин Рогоза

К.е.н., доцент, доцент кафедри інформаційних систем і технологій
НУБіП України, факультет інформаційних технологій, м. Київ, Україна
ORCID ID 0000-0002-9417-1745
konstantin.r@nubip.edu.ua

РОЗВИТОК ПЛАТФОРМИ ETHEREUM ТА ЇЇ ЗАСТОСУВАННЯ В СІЛЬСЬКОМУ ГОСПОДАРСТВІ

Анотація. У статті розглянуто еволюцію Ethereum як однієї з провідних блокчейн платформ, що забезпечує децентралізоване виконання смарт-контрактів. Проаналізовано технічну еволюцію Ethereum, зокрема перехід на алгоритм консенсусу Proof of Stake, впровадження рішень другого рівня (Layer 2) та перспективу реалізації Danksharding, що забезпечують масштабованість, енергоефективність та зниження транзакційних витрат. Приділено увагу аналізу можливостей застосування Ethereum у сфері сільського господарства, зокрема в контексті прозорості ланцюгів постачання, управління земельними ресурсами, агрострахування та мікрофінансування фермерських господарств. Визначено переваги використання смарт-контрактів для автоматизації аграрних процесів, зниження транзакційних витрат та підвищення довіри між учасниками аграрного ринку. Запропоновано напрями подальших досліджень щодо інтеграції блокчейн-рішень у цифрову трансформацію аграрної галузі.

Ключові слова: Ethereum; блокчейн; смарт-контракти; агротехнології; сільське господарство.

1. ВСТУП

Постановка проблеми. Сільське господарство стикається з численними викликами, серед яких – непрозорість ланцюгів постачання, складність у верифікації походження продукції, обмежений доступ до фінансування та страхування для малих фермерів. У цьому контексті блокчейн-технології, зокрема Ethereum, відкривають нові можливості для вирішення даних проблем.

Аналіз останніх досліджень і публікацій. У літературі [2] - [3] активно досліджуються переваги блокчейну в агросекторі, зокрема щодо простежуваності продукції, автоматизації угод та зниження витрат. Проте, недостатньо уваги приділено практичним аспектам впровадження Ethereum у сільське господарство країн, що розвиваються, зокрема України.

Мета публікації. Метою статті є аналіз розвитку платформи Ethereum та виявлення потенційних напрямів її застосування в аграрному секторі, з акцентом на практичні кейси та перспективи для українського сільського господарства.

2. ТЕОРЕТИЧНІ ОСНОВИ

Ethereum – це децентралізована платформа з відкритим кодом, яка дозволяє створювати та виконувати смарт-контракти. Її ключовими перевагами є прозорість, незмінність даних та можливість автоматизації взаємодії між сторонами без посередників. У сільському господарстві ці властивості можуть бути використані для створення систем

Смарт-контракти – це програмні алгоритми, що виконуються на блокчейн-платформах і призначені для автоматизації виконання умов угод між сторонами без участі посередників. Вони являють собою код, який зберігається у блокчейні та активується при настанні визначених умов. Основною особливістю смарт-контрактів є їхня незмінність і прозорість: після розгортання в мережі код не може бути змінений, а всі операції доступні для перевірки.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Платформа Ethereum продовжує активно розвиватись, трансформуючись із простої блокчейн-мережі у багатофункціональну інфраструктуру для децентралізованих застосунків (dApps), фінансових сервісів (DeFi), NFT, DAO та інших інноваційних рішень. Одним із ключових етапів цього розвитку став перехід на алгоритм консенсусу Proof of Stake (PoS) у межах оновлення Ethereum 2.0, що дозволив суттєво знизити енергоспоживання мережі та створити передумови для масштабування. Впровадження рішень другого рівня (Layer 2), таких як Optimism, Arbitrum та zkSync, дозволяє зменшити навантаження на основну мережу, знижуючи комісії та підвищуючи швидкість транзакцій, що робить її більш доступною для широкого кола користувачів і бізнесів. Найближчим часом очікується реалізація технології Danksharding, що підвищить пропускну здатність мережі та зменшить вартість зберігання даних.

Ethereum залишається лідером у сфері DeFi (децентралізовані фінанси). Очікується подальше зростання кількості протоколів, що надають фінансові послуги без посередників. Зростає і інституційний інтерес до DeFi, що може призвести, як до нових регуляторних викликів так і до більшої легітимності. Хоча ажіотаж навколо NFT (невзаємозамінні токени) зменшився, технологія знаходить нові застосування у сферах нерухомості, авторських прав, ідентифікації. Ethereum залишається основною платформою для NFT, хоча конкуренція зростає. DAO (децентралізовані автономні організації) стають дедалі популярнішими як форма управління спільнотами та проєктами. Ethereum забезпечує інфраструктуру для створення та управління DAO, що може змінити підходи до корпоративного управління.

Ethereum відкриває широкі можливості для трансформації аграрного сектору, забезпечуючи прозорість, автоматизацію та доступ до фінансових ресурсів. Завдяки смарт-контрактам можна створювати механізми, які автоматично виконують умови угод між фермерами, постачальниками та покупцями без участі посередників. Це значно знижує ризики шахрайства, скорочує час на оформлення документів і зменшує транзакційні витрати. Наприклад, у ланцюгах постачання блокчейн дозволяє фіксувати кожен етап руху продукції – від вирощування до реалізації, що забезпечує повну простежуваність і підвищує довіру споживачів до якості товарів.

Ще одним перспективним напрямом є агрострахування. Смарт-контракти можуть бути запрограмовані на автоматичні виплати у разі настання погодних ризиків, підтверджених даними з метеостанцій або супутникового моніторингу. Це робить процес страхування швидким, прозорим і менш затратним. DeFi дають можливість малим фермерським господарствам отримувати кредити без банківських установ, а також залучати інвестиції через токенизацію майбутнього врожаю або земельних ділянок.

Важливим аспектом є управління земельними ресурсами. Використання блокчейну для ведення кадастрових реєстрів забезпечує незмінність даних і прозорість прав власності, що знижує ризики рейдерства та спрощує процедури купівлі-продажу землі. У поєднанні з технологіями Інтернету речей (IoT) та супутниковими системами Ethereum може стати основою для створення «розумного фермерства», де більшість рішень приймаються автоматично на основі даних про стан ґрунту, погоду та врожайність. Крім того, DAO можуть бути використані для колективного управління аграрними кооперативами або спільнотами фермерів. У перспективі все це дозволить аграрному сектору перейти на новий рівень цифрової трансформації, підвищити ефективність виробництва та забезпечити стійкість до зовнішніх ризиків.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Ethereum продовжує еволюціонувати, перетворюючись на потужну платформу для децентралізованих рішень, що виходять далеко за межі фінансового сектору. Його технічні оновлення забезпечують покращення масштабованості, економічності і екологічності мережі. Для аграрного сектору це означає можливість створення прозорих ланцюгів постачання, автоматизації страхових виплат, доступу до децентралізованого фінансування та ефективного управління земельними ресурсами. Інтеграція блокчейн-рішень на базі Ethereum є стратегічним напрямом цифрової трансформації аграрної галузі, здатним забезпечити її конкурентоспроможність у глобальній економіці.

ПОСИЛАННЯ

- [1] Ethereum.org, "Ethereum roadmap," [Online]. Available: <https://ethereum.org/roadmap>.
- [2] The Gazette, "Smart farms and Ethereum: Using blockchain to boost crop yields," 2025. [Online]. Available: <https://www.thegazette.com.au/news/smart-farms-ethereum-using-blockchain-to-boost-crop-yields>.
- [3] F. Kayusi, "Blockchain Technology: Improving Agricultural Supply Chain Efficiency and Transparency – A Review," ResearchGate, 2025. [Online]. Available: https://www.researchgate.net/publication/388485665_Blockchain_Technology_Improving_Agricultural_Supply_Chain_Efficiency_and_Transparency_-_A_Review.

Владислав Стеценко

v.stetsenko@nubip.edu.ua

Олена Глазунова

доктор педагогічних наук, професор, проректор з науково-педагогічної роботи та цифрової трансформації

Місце роботи: Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID 0000-0002-0136-4936

o-glazunova@nubip.edu.ua

Петро Дрозд

кандидат історичних наук, доцент, доцент кафедри адміністративного менеджменту і ЗЕД

Місце роботи: Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID 0000-0003-1939-2967

drozd_p@nubip.edu.ua

Валентина Корольчук

доктор філософії, доцент, доцент кафедри інформаційних систем і технологій

Місце роботи: Національний університет біоресурсів і природокористування України, Київ, Україна

ORCID ID 0000-0002-3145-8802

korolchuk@nubip.edu.ua

ЦИФРОВЕ ОСВІТНЄ СЕРЕДОВИЩЕ: ВІД ІНСТРУМЕНТУ ДО СИСТЕМИ УПРАВЛІННЯ. ПРОБЛЕМИ АРХІТЕКТУРИ ТА ВПРОВАДЖЕННЯ

Анотація. У статті розглянуто трансформацію цифрового освітнього середовища університету від інструменту підтримки навчального процесу до інтегрованої системи управління. Проаналізовано роль ERP-систем і комплексних цифрових платформ у забезпеченні прозорості, ефективності та аналітичного управління освітньою діяльністю. На прикладі Національного університету біоресурсів і природокористування України описано архітектуру цифрової екосистеми, зокрема впровадження системи «Master» та платформи NUBiP Digital. Визначено основні проблеми й виклики цифрової трансформації, серед яких – потреба у стандартизації даних, підвищенні рівня цифрових компетентностей персоналу та гарантуванні кібербезпеки. Зроблено висновок, що цифрове освітнє середовище є ключовим чинником розвитку сучасного університету й потребує подальшого вдосконалення шляхом інтеграції аналітики, штучного інтелекту та гнучких управлінських моделей.

Ключові слова: цифрове освітнє середовище; цифрова трансформація; ERP-система; управління освітою; аналітика даних; інтегровані освітні платформи.

1. ВСТУП

Сучасна вища освіта перебуває в умовах глибокої цифрової трансформації та переживає перехід від традиційної моделі навчання до цифрової, де ключову роль відіграють дані, аналітика й штучний інтелект. Ця трансформація зумовлена глобальним переходом до економіки знань, у якій знання, інновації та інформаційні ресурси стають основним капіталом, а також стрімким розвитком технологій штучного інтелекту.

Сьогодні цифрові платформи виконують не лише функцію допоміжних засобів навчання – вони формують ядро управління освітнім процесом. Цифрове освітнє середовище університету стає інтегрованою системою, яка поєднує освітні ресурси, дані, сервіси та аналітичні інструменти, забезпечуючи прозорість, ефективність і персоналізацію навчання. Воно базується на взаємодії інформаційних систем, управлінських процесів і технологій штучного інтелекту, формуючи єдиний цифровий простір для прийняття управлінських рішень у межах університету.

Метою даного дослідження є аналіз трансформації цифрового освітнього середовища університету від інструменту підтримки навчання до повноцінної системи управління, визначення ролі ERP-систем у забезпеченні освітнього процесу, а також оцінка можливостей аналітики даних для прийняття управлінських рішень і визначення основних проблем архітектури та впровадження цифрового середовища.

2. ТЕОРЕТИЧНІ ОСНОВИ

Цифрова трансформація освіти передбачає не лише автоматизацію окремих процесів, а й докорінну зміну самої моделі функціонування університету. Концепція цифрового університету (Digital University) ґрунтується на створенні єдиного інформаційного простору, у межах якого освітні, наукові, адміністративні та комунікаційні процеси взаємодіють на основі спільних даних і технологічних рішень.

Теоретичні засади створення цифрового освітнього середовища спираються на принципи інтеграції інформаційних систем, аналітичного управління на основі даних (Data-Driven Decision Making), персоналізації навчання та розвитку цифрових компетентностей усіх учасників освітнього процесу. Таким чином, цифрове середовище розглядається не лише як технологічна платформа, а як нова модель управління знаннями, яка забезпечує безперервність освітнього процесу та гнучкість прийняття рішень.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Для досягнення мети використано системний, порівняльний, емпіричний та аналітичний методи. Системний аналіз дав змогу описати архітектуру цифрового освітнього середовища університету, а порівняльний – простежити перехід від цифровізації до створення інтегрованої екосистеми управління. Емпіричний і аналітичний підходи застосовано для оцінки ефективності впровадження ERP-системи «Master» та платформи NUBiP Digital у процесі цифрової трансформації.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Впровадження ERP-системи «Master» у НУБіП України стало основою цифрової трансформації університету. Ця система об'єднала адміністративні, фінансові, кадрові процеси, створивши єдине джерело достовірних даних для всіх рівнів управління. У результаті вдалося усунути фрагментацію даних, мінімізувати дублювання інформації, стандартизувати документообіг і забезпечити прозорість управлінських процесів.

Паралельно з ERP-системою активно розвивається інтегрована інформаційна система NUBiP Digital, яка охоплює модулі для студентів, викладачів та адміністрації. Серед них – «Електронний деканат», «My NUBiP», модулі планування педагогічного навантаження, моніторингу ефективності викладачів та управлінської аналітики. Ці інструменти забезпечують не лише автоматизацію рутинних процесів, а й формують середовище прийняття управлінських рішень на основі аналітичних даних.

У межах цифрової екосистеми університету функціонує понад п'ять тисяч електронних курсів на платформі Moodle, корпоративні сервіси Google Workspace і Microsoft 365, цифрова бібліотека NUBiP Library та хмарні ресурси для досліджень і обробки великих даних. Така інтеграція створює передумови для розвитку системи управління освітою, заснованої на аналітичних показниках. Архітектура середовища, яке інтегрує всі вказані системи та модулі представлена на рис. 1.

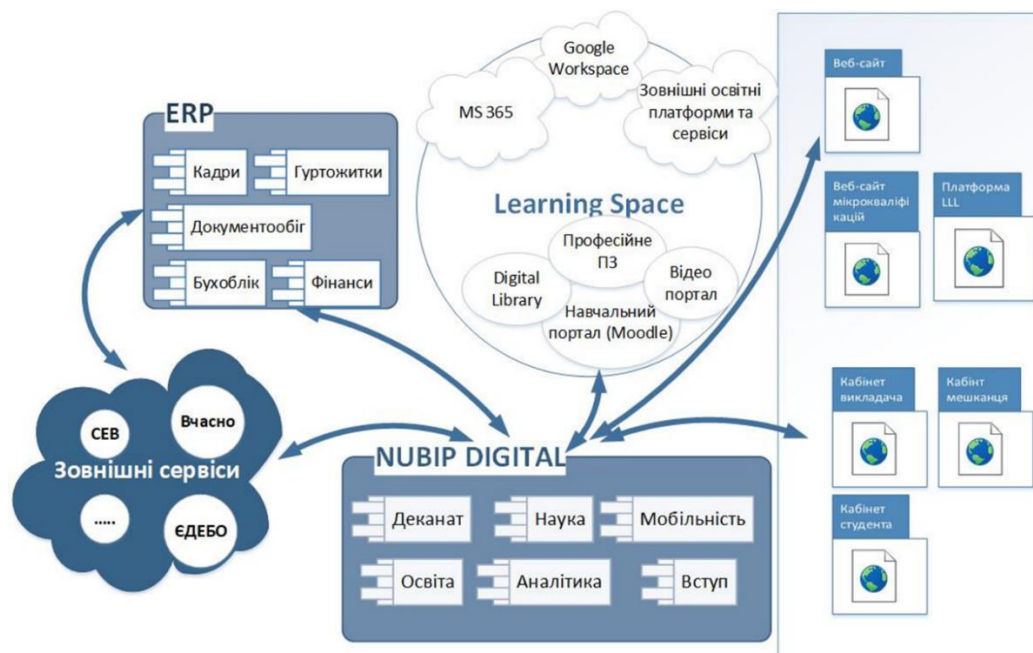


Рис.1. Архітектура цифрового освітнього середовища НУБіП України

Водночас процес цифрової трансформації супроводжується певними проблемами, серед яких – відсутність єдиних стандартів інтеграції між системами, потреба в гарантуванні кібербезпеки, недостатній рівень цифрових компетентностей персоналу та нестача стабільного фінансування цифрових проєктів. Ці виклики вимагають комплексного підходу до управління змінами й формування нової цифрової культури в університетському середовищі.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Розвиток цифрового освітнього середовища у вищій освіті свідчить про зміну управлінської парадигми університету. Цифрові технології перетворюються з допоміжного інструмента на комплексну екосистему, що об'єднує навчальну, наукову, адміністративну та аналітичну діяльність. Досвід НУБіП України демонструє, що системна цифрова трансформація підвищує ефективність управління, прозорість процесів і якість прийняття рішень. Подальший розвиток має зосереджуватися на створенні гнучкої архітектури цифрових рішень, удосконаленні кібербезпеки, розвитку цифрових компетентностей і впровадженні технологій штучного інтелекту в управління освітою.

ПОСИЛАННЯ (ПЕРЕКЛАДЕНІ ТА ТРАНСЛІТЕРОВАНІ)

1. Глазунова О.Г. Цифрова трансформація університету в епоху штучного інтелекту. Презентація до виступу на конференції, 02.10.2025.
2. Кухаренко В.М. Цифрова педагогіка: концепції та практики. Харків: ХНУ, 2023.
3. OECD. Digital Education Outlook 2023: Artificial Intelligence in Education. Paris: OECD Publishing, 2023.
4. UNESCO. Guidelines for AI in Education. Paris: UNESCO, 2022.

Михайло Швиденко

Кандидат економічних наук, доцент, завідувач кафедри

Національний університет біоресурсів і природокористування України, кафедра інформаційних систем і технологій, Київ, Україна

ORCID ID 0000-0002-9025-1326

shvydenko@nubip.edu.ua

Сергій Саяпін

Кандидат економічних наук, доцент

Національний університет біоресурсів і природокористування України, кафедра інформаційних систем і технологій, Київ, Україна

ORCID ID 0000-0003-1565-4034

sayapin@nubip.edu.ua

ТРАНСФОРМАЦІЯ СІЛЬСЬКОГОСПОДАРСЬКОГО ЕЛЕКТРОННОГО ДОРАДНИЦТВА ЗА ДОПОМОГОЮ ІНСТРУМЕНТІВ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. Публікація присвячена дослідженню та обґрунтуванню доцільності впровадження високоінтелектуальних багатофункціональних чат-ботів та віртуальних асистентів, побудованих на технологіях штучного інтелекту (ШІ) та Великих Мовних Моделей (LLM), як складовою у платформу сільськогосподарського електронного дорадництва (еДорада). У тексті наведено низку глобальних прикладів застосування ШІ в агросекторі, які ілюструють широку варіативність їхнього функціоналу. Процес створення дорадчого чат-бота окреслено як послідовний підхід, що охоплює чотири ключові етапи, які поєднують інженерію ПЗ, агрономічні знання та технології ШІ. Таким чином, стаття детально обґрунтовує стратегічну важливість, необхідний функціонал та поетапний план розробки ШІ-асистента, що має стати високоінтелектуальним, надійним та доступним інструментом підтримки для українського агросектору..

Ключові слова: штучний інтелект (ШІ); мовна модель LLM; ШІ консультатнт, платформа електронного дорадництва.

Застосування інструментів штучного інтелекту в системі сільськогосподарського електронного дорадництва (е-дорадництва) [1] є важливим кроком до автоматизації, підвищення доступності та оперативності консультаційних послуг для аграріїв, обходячи потребу залучення експерта-людини для вирішення загальних питань. Ці інструменти, часто побудовані на основі штучного інтелекту (ШІ), включаючи великі мовні моделі (наприклад, ChatGPT, Gemini), допомагають фермерам швидко отримувати потрібну інформацію за нагальної потреби та в будь-який час доби [2, 3, 4].

Існує багато прикладів чат-ботів та платформ, які використовують технології ШІ для допомоги фермерам. Вони різняться за функціоналом: від надання базової інформації до складної діагностики та з'єднання з експертами.

Боти на основі Великих Мовних Моделей (LLM). Це віртуальні помічники, які використовують технології, схожі на ті, що лежать в основі ChatGPT, але навчені на спеціалізованих аграрних даних. Наприклад, AI Ag Advisor Norm [5, 6], який є одним з перших галузевих агроботів на базі LLM, розроблений для надання консультацій щодо сільськогосподарських питань та Farmer.CHAT (Digital Green & Goodey.AI) [6], що являє собою мультимовну платформу на базі GPT-4, яка розроблена для фермерів у країнах, що розвиваються (Індія, Ефіопія, Кенія).

Гібридні системи (Бот + Живий Експерт). Наприклад, AGvisorPRO: платформа, яка підбирає агронома-експерта для фермера [8], до якого його питання надходить анонімно. Це забезпечує висококваліфіковану відповідь на унікальні та складні проблеми, які бот не може вирішити.

Боти для діагностики та прийняття рішень. Наприклад, AgriBot [7], який дає відповіді на загальні запитання, рекомендації по культурам згідно параметрів ґрунту та

прогнозування хвороб на основі завантаженого фермером зображення ураженої рослини.

Галузеві та Маркетингові Боти, що створені для конкретних бізнес-потреб або сегментів ринку. Наприклад, Kaily Chatbot, який асистує у виборі продуктів (насіння, добрив) на основі типу ґрунту/культури та український Telegram-бот (AgroMarketUUB [9]): який допомагають фермерам продавати або купувати зерно та іншу продукцію

Ці приклади ілюструють, як інструменти на основі ШІ трансформують е-дорадництво, роблячи знання доступнішими, хоча й підкреслюють потребу в якісній верифікації даних для критичних рішень.

Тому для підвищення ефективності, швидкості та доступності сільськогосподарських дорадчих послуг для українських фермерів доцільне впровадження високоінтелектуального, багатофункціонального чат-бота/віртуального асистента дорадника, інтегрованого у платформу eDorada [1].

Основне завдання чат-бота на базі ШІ, який поєднує наявні бази даних та новітні находження агрознав та інновацій – автоматизувати первинну взаємодію з фермерами, надаючи миттєву та релевантну інформацію, тим самим розвантажуючи живих дорадників, зокрема:

- Миттєве консультування: надавати швидкі та точні відповіді на типові запитання, зокрема енциклопедичного змісту (опис, агротехнології, ЗЗР, добрива, класифікації).
- Визначення/діагностика за зображенням: визначення об'єктів, їх опис та ступеня критичності (хвороби, шкідників або дефіцит живлення на основі фотографій, превентивна ветдіагностика).
- Надання локальних даних: інтегруватися з метео- та ринковими базами та ШІ інструментами посередництвом їх API механізмів для надання локалізованих прогнозів і цін.
- Персоналізація: адаптувати поради, використовуючи дані з профілю фермера та історії спілкування (культура, фаза розвитку, тип ґрунту, попередня озвучена проблематика).

Створення чат-боту для платформи сільськогосподарського дорадництва вимагає послідовного підходу, що поєднує інженерію програмного забезпечення, агрономічні знання та технології штучного інтелекту. Умовна структура ШІ-рішення для платформи електронного дорадництва, яка враховує наявні структуровані інформаційні ресурси з відомими алгоритмами відбору енциклопедичної інформації (AgroUA.net), фахові публікації та оперативна інформація (eДорада [1]), генеровані знання та інновації в НУБіП України у цифровому форматі представлено на рис.1.

Розробка чат-боту відбувається в декілька етапів:

Етап 1: *Визначення цілей та аудиторії*: Бот має навчитися давати відповіді як на загальні запитання, так і на конкретні галузеві питання (наприклад, відповідати на юридичні запитання, податкову політику тощо, так і надавати рекомендації щодо внесення добрив, діагностувати хвороби і т.п.). Користувачами виступають фермерські господарства, агрономи, власники присадибних ділянок Це визначає мову, термінологію та складність відповідей.

Вибір платформи та каналу: веб-сайт (вбудований віджет).

Створення бази знань (Knowledge Base): Збір, систематизація та валідація агрономічних даних (бібліотеки хвороб, норми внесення пестицидів та добрив, технологічні карти, державні стандарти тощо). Основа – бази даних сайтів *e-Дорада та Аграрний сектор України AgroUA.net*.

Етап 2: *Розробка та навчання моделі*: для розробки архітектури ШІ-бота аналізуються різні технології, як використання готових платформ (Google Dialogflow,

Microsoft Bot Framework), так і розробка власного рішення на основі LLM (великих мовних моделей) для складніших запитів.

Розробка інтерфейсу користувача (UI/UX): Проектування "сценаріїв" (flows) розмови. Більшість сільськогосподарських запитів потребують покрокової ідентифікації (наприклад, "Яка культура?", "Який симптом?"). Обробка природної мови (NLP) та машинне навчання (ML): модель має розуміти наміри (Intents) користувача (наприклад, Intent: "Діагностика хвороби") та виділяти сутності (Entities) (наприклад, Entity: "Культура: Пшениця", Entity: "Симптом: Жовтіє листя").

Тренування моделі на підготовленій Базі знань.

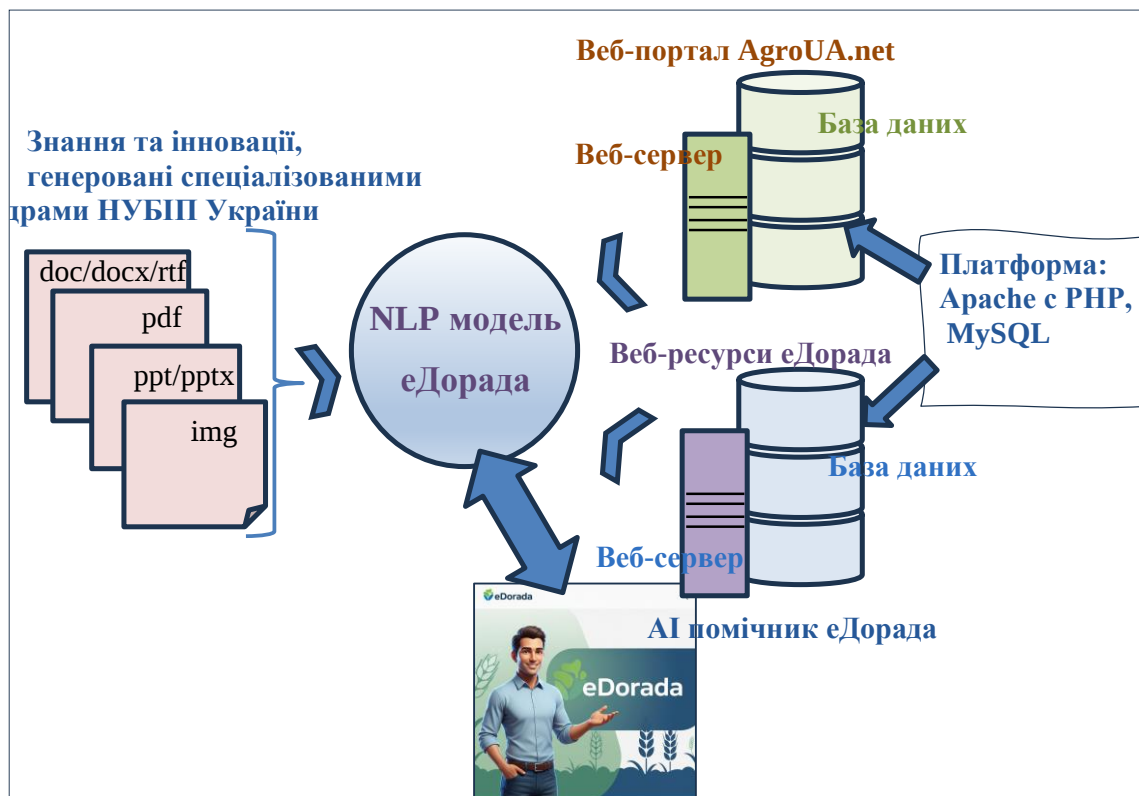


Рис. 1. Ресурси та джерела спеціалізованих знань та інновацій для створення ШІ-помічника на платформі електронного дорадництва.

Етап 3: Тестування та валідація: Внутрішнє тестування (Alpha-тест) розробниками на логіку сценаріїв та коректність роботи інтеграції з Базою знань. Перевірка на "крайові випадки" (edge cases) – нетипові запити або помилки користувача.

Пілотне тестування (Beta-тест): Залучення реальних користувачів, зокрема фермерів до тестування.

Збір зворотного зв'язку щодо точності агрономічних рекомендацій та зручності користування. Перевірка, чи не дає бот "галюцинацій" (неправдивих, але правдоподібних відповідей), що є критичним для сільського господарства.

Етап 4: Впровадження та підтримка: Запуск та розгортання мовної моделі ШІ та комунікативного чат-боту на обраній платформі з відстеженням метрики успішності (наприклад, відсоток запитів, на які бот зміг відповісти самостійно). Проведення аналізу для виявлення найпоширеніших запитань та "провалів" моделі. Постійне оновлення агрономічних даних (нові сорти, зміна законодавства, нові

хвороби/шкідники). Проведення регулярного перенавчання NLP-моделі на нових даних та запитах користувачів для підвищення точності відповідей.

ВИСНОВКИ

Попри значний потенціал, використання ШІ в системах е-дорадництва пов'язане з певними викликами:

–необхідність якісних і структурованих даних для навчання моделей. Якість і актуальність порад чат-бота безпосередньо залежать від якості та обсягу бази знань, на якій він навчався. У сільському господарстві, де умови сильно варіюються, помилки можуть мати високу ціну;

–Проблемним виявляється опрацювання матеріалів наявних довідкових та дорадчих веб-систем, оскільки логіка зв'язків лежить у програмному коді;

–питання етики, безпеки даних та прозорості алгоритмів;

–потреба у підготовці фахівців, здатних працювати з інтелектуальними системами та потреба у їх залученні на постійній основі.

–високі вимоги до технічного забезпечення для розгортання платформи з ШІ складовою.

Отже, чат-боти є перспективним та потужним інструментом, що доповнює, але поки не замінює повноцінну консультативну роботу дорадників та галузевих експертів.

У перспективі очікується, що платформа електронного дорадництва з використанням ШІ стане інтерактивними цифровими дорадчим інструментом нового покоління, які забезпечуватимуть інтегровану підтримку прийняття рішень у реальному часі, сприятимуть розвитку сталого господарювання та підвищенню конкурентоспроможності економіки.

ПОСИЛАННЯ

1. Платформа електронного дорадництва, головний ресурс «еДорада». Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://edorada.org/uk>

2. Інтелект збирання врожаю: як генеративний штучний інтелект перетворює сільське господарство Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://www.unite.ai/uk/harvesting-intelligence-how-generative-ai-is-transforming-agriculture/>

3. Рене Керхіус Коли чат-бот замінить агронома Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://ifarming.ua/monitoring/koly-chat-bot-zaminyt-agronoma>

4. First AI Ag Advisor Provides Farmers With a Wide Array of Agronomic Intelligence Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://www.croplife.com/smart-tech/first-ai-ag-advisor-provides-farmers-with-a-wide-array-of-agronomic-intelligence/>

5. Агрономічний радник на базі штучного інтелекту Norm <https://www.fbn.com/norm>

6. Підвищення швидкості та ефективності сільськогосподарських консультацій за допомогою штучного інтелекту FarmerChat Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://digitalgreen.org/farmer.chat/>

7. Agribot Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://agribot.ai/>

8. Рішення AGvisorPRO Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://getagvisorpro.com/>

9. Купити або продати кукурудзу, пшеницю, соняшник за привабливими цінами у Агробітні AgroMarket UUB Дата звернення: 03 листопада 2025. [Онлайн]. Доступно: <https://www.uub.com.ua/torgy-ta-auktsiony/agrobot-agromarketuub/>

SECTION 5. AUTOMATION, COMPUTER-INTEGRATED TECHNOLOGIES, ROBOTICS, ARTIFICIAL INTELLIGENCE/АВТОМАТИЗАЦІЯ, КОМП'ЮТЕРНО-ІНТЕГРОВАНІ ТЕХНОЛОГІЇ, РОБОТОТЕХНІКА, ШТУЧНИЙ ІНТЕЛЕКТ

Павлов Сергій Григорович

Аспірант

Національний університет біоресурсів і природокористування України, кафедра автоматики та робототехнічних систем ім. акад. І.І. Мартиненка, Київ, Україна

<https://orcid.org/0009-0001-3343-5508>

sergpavlov89@email.com

Лисенко Віталій Пилипович

Доктор технічних наук, професор

Національний університет біоресурсів і природокористування України, кафедра автоматики та робототехнічних систем ім. акад. І.І. Мартиненка, Київ, Україна

<https://orcid.org/0000-0002-5659-6806>

lysenko@nubip.edu.ua

Лендєл Тарас Іванович

Кандидат технічних наук, доцент

Національний університет біоресурсів і природокористування України, кафедра автоматики та робототехнічних систем ім. акад. І.І. Мартиненка, Київ, Україна

<https://orcid.org/0000-0002-6356-1230>

taraslendel@gmail.com

Наконечна Катерина Віталіївна

Кандидат економічних наук, доцент

Національний університет біоресурсів і природокористування України, кафедра економічної кібернетики, Київ, Україна

<https://orcid.org/0000-0002-1537-7201>

nakonechna@nubip.edu.ua

ФУНКЦІЇ БАЖАНОСТІ ХАРІНГТОНА ДЛЯ ОЦІНКИ ЯКІСНИХ ПАРАМЕТРІВ БІОСИРОВИНИ ДЛЯ ЗБРОДЖУВАННЯ

Анотація. Стабільність та економічна ефективність біогазових установок критично залежать від якості біосировини, що використовується для анаеробного збродження. Однак традиційні лабораторні методи аналізу, які можуть тривати від кількох днів до тижнів, є занадто повільними для оперативного контролю технологічного процесу. Така затримка створює значні ризики зниження продуктивності, дестабілізації процесу та навіть повної зупинки реактора, що призводить до економічних збитків. Це формує гостру промислову потребу в інструментах для експрес-оцінки якості сировини, здатних надавати обґрунтовані рекомендації в режимі реального часу. Ключовою методологічною інновацією даної роботи є стратегічне перетворення складного завдання кількісного прогнозування виходу біогазу, яке є викликом для невеликих наборів даних, у практично значущу та більш стійку задачу якісної класифікації. Цього вдалося досягти шляхом застосування функції бажаності Харінгтона, яка формалізує експертні знання щодо технологічних вимог в універсальну безрозмірну шкалу якості. Для вирішення цього завдання було навчено класифікатор на основі алгоритму випадкового лісу з використанням обмеженого набору параметрів, а саме вологість, узагальнена оцінка, температура. Для об'єктивної оцінки надійності моделі було використано стратифіковану перехресну валідацію, а важливість ознак визначалася за допомогою статистично обґрунтованого методу перестановок. Модель продемонструвала високу ефективність (зважаючи на F1-міру 0.82), а аналіз показав, що узагальнена візуальна оцінка та вологість є домінуючими і статистично значущими предикторами якості сировини. Розроблена система є доказом концепції для створення "паспорта якості" сировини, що дозволить

операторам приймати обґрунтовані рішення для підвищення стабільності та рентабельності біогазових установок.

Ключові слова: біогаз; анаеробне зброджування; експрес-оцінка якості; машинне навчання; випадковий ліс; функція бажаності Харінгтона

1. ВСТУП

Постановка проблеми. Ефективність виробництва біогазу безпосередньо залежить від стабільності процесу анаеробного зброджування [1], який, у свою чергу, є чутливим до якості вхідної біосировини. Традиційні лабораторні методи оцінки є точними, але їх тривалість (від кількох днів до тижнів) унеможливорює оперативний контроль. Відсутність інструментів для швидкої оцінки якості сировини "на вході" призводить до ризиків зниження виходу біогазу та дестабілізації процесу, що створює потребу в розробці систем експрес-аналізу для прийняття рішень в режимі реального часу.[1]

Аналіз останніх досліджень і публікацій. У сучасних дослідженнях активно застосовуються методи машинного навчання для прогнозування виходу біогазу. Однак більшість моделей потребують великих наборів даних та складних лабораторних аналізів, що обмежує їх практичне застосування.[1] Ключовою проблемою залишається створення моделей, що працюють на малих вибірках та використовують легкодоступні параметри. Наша робота спирається на ідеї перетворення кількісних показників у якісні, що є більш стійким підходом при обмежених даних.[1]

Мета публікації. Метою статті є розробка та валідація системи експрес-оцінки якості біосировини для зброджування, що базується на поєднанні функції бажаності Харінгтона для якісної класифікації та алгоритму машинного навчання "випадковий ліс" для побудови прогностичної моделі на основі мінімального набору оперативно вимірюваних параметрів.

2. ТЕОРЕТИЧНІ ОСНОВИ

Основою запропонованого підходу є перетворення складної задачі кількісного прогнозування на більш стійку задачу якісної класифікації. Для цього використано функцію бажаності Харінгтона[4] — інструмент, що дозволяє перевести будь-яку вимірювану величину (у нашому випадку — питомий вихід біогазу) в універсальну безрозмірну шкалу "бажаності" від 0 (абсолютно неприйнятно) до 1 (ідеально). Цей метод дозволяє формалізувати нелінійні експертні знання та технологічні вимоги, імітуючи логіку прийняття рішень фахівцем.

Для вирішення задачі класифікації було обрано алгоритм випадкового лісу (Random Forest)[2]. Цей ансамблевий метод є ефективним при роботі з невеликими наборами даних, оскільки він стійкий до перенавчання завдяки агрегації результатів багатьох незалежних дерев рішень, кожне з яких навчається на випадковій вибірці даних та ознак.[2]

3. МЕТОДИ ДОСЛІДЖЕННЯ

Дослідження проводилось на експериментальній вибірці. Для кожного зразка було визначено питомий вихід біогазу, який за допомогою функції бажаності Харінгтона було перетворено на один з п'яти класів якості. В якості вхідних параметрів для моделі машинного навчання були обрані легкодоступні показники: вологість, узагальнена візуальна оцінка та набір температурних показників.

Для об'єктивної оцінки ефективності моделі на малій вибірці було застосовано процедуру стратифікованої 5-кратної перехресної валідації, що забезпечує репрезентативність кожного класу в навчальних та тестових блоках. Визначення найважливіших предикторів здійснювалося за допомогою статистично обґрунтованого

методу перестановок (Permutation Importance) [3], який оцінює падіння якості моделі при випадковому перемішуванні значень однієї ознаки.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Розроблена модель класифікації продемонструвала високу загальну ефективність із зваженою F1-мірою 0.82[1]. Модель показала відмінні результати у розпізнаванні класів "Добре" (F1=0.90) та "Дуже погано" (F1=1.00).

Аналіз важливості параметрів однозначно визначив два ключові, статистично значущі ($p < 0.001$) предиктори - Узагальнена візуальна оцінка та Вологість.

Натомість, усі температурні показники виявилися статистично незначущими, візуалізація на Рис.1

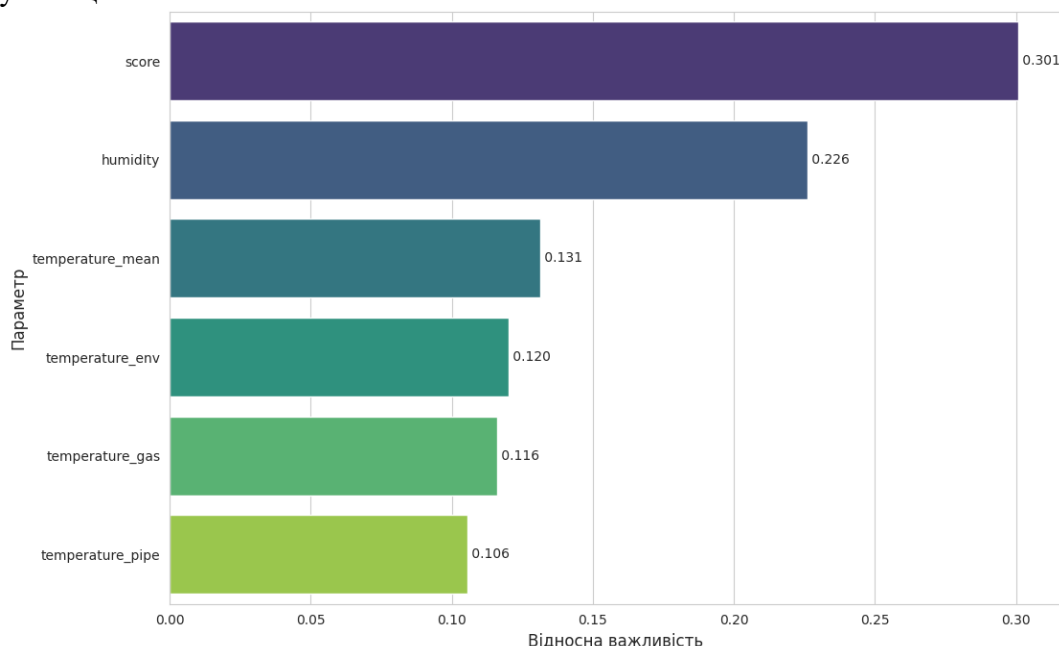


Рис. 1. Вплив параметрів

Аналіз помилок показав, що модель переважно плутає суміжні класи (наприклад, "Добре" та "Задовільно"), що свідчить про її здатність засвоювати порядкову природу шкали якості. Розроблена система може слугувати основою для створення "паспорта якості" сировини, що дозволить операторам приймати обґрунтовані рішення в реальному часі: відбраковувати непридатні партії, коригувати рецептуру сумішей або призначати попередню обробку.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Дослідження успішно продемонструвало, що гібридний підхід, який поєднує функцію бажаності Харінгтона та алгоритм випадкового лісу, дозволяє створити ефективну систему експрес-оцінки якості сировини для біогазових установок. Доведено, що узагальнена візуальна оцінка та вологість є ключовими предикторами якості. Головним обмеженням роботи є малий розмір вибірки, тому перспективи подальших досліджень включають збір більшого набору даних, тестування більш складних алгоритмів та інтеграцію даних з інших недорогих сенсорів для підвищення точності та надійності системи в промислових умовах.

ПОСИЛАННЯ

1. F. Almomani, "Prediction of biogas production from chemically treated co-digested agricultural waste using artificial neural network," *Fuel*, vol. 280, p. 118573, 2020, doi: 10.1016/j.fuel.2020.118573.
2. В. Лисенко, Т. Лендел, і С. Павлов, "Аналіз алгоритмів машинного навчання для прогнозування виходу біогазу," *Енергетика і автоматика*, no. 3, pp. 100–111, 2023. [Online]. Available: [https://doi.org/10.31548/energiya3\(67\).2023.100](https://doi.org/10.31548/energiya3(67).2023.100)
3. "Cross-validation: evaluating estimator performance," *Scikit-learn Documentation*, 2025. [Online]. Available: https://scikit-learn.org/stable/modules/cross_validation.html
4. NumberAnalytics, "Unlocking Desirability in Operations Research," 2025. [Online]. Available: <https://www.numberanalytics.com/blog/ultimate-guide-desirability-function-operations-research>

Станіслав Усенко

Аспірант

Місце роботи: Київський політехнічний інститут імені Ігоря Сікорського, Україна, 03056, м. Київ,
Берестейський проспект (Перемоги), 37

ORCID ID: <https://orcid.org/0000-0002-9412-2441>

Farkry17@gmail.com

Максим Клименко

Здобувач

Місце роботи: Національний університет харчових технологій, Україна, 01601, м. Київ, вул.

Володимирська, 68

maksym1klymenko@gmail.com

Євгеній Шаповалов

С.н.с., к.т.н.

Місце роботи: Національний університет харчових технологій, Україна, 01601, м. Київ, вул.

Володимирська, 68

ORCID ID: <https://orcid.org/0000-0003-3732-9486>

sjb@man.gov.ua

Сергій Жадан

к.т.н.

Місце роботи: Національний університет харчових технологій, Україна, 01601, м. Київ, вул.

Володимирська, 68

ORCID ID: <https://orcid.org/0000-0002-7493-2180>

serezha.nuft@gmail.com

КОНЦЕПТУАЛЬНИЙ ФРЕЙМВОРК УНІВЕРСАЛЬНОЇ ШІ ДЛЯ ОПТИМІЗАЦІЇ РОБІТ БІОГАЗОВИХ УСТАНОВОК

Анотація. Метанова ферментація є ключовою технологією для перетворення органічних відходів на біогаз, однак ефективне управління процесом ускладнюється через його нелінійність та варіативність сировини. Моделі штучного інтелекту (ШІ) демонструють високий потенціал для прогнозування та оптимізації виробництва біогазу, але їх розробка для кожної окремої установки є ресурсозатратною та вимагає великих масивів історичних даних. Більше того, моделі, навчені на даних однієї установки, погано генералізуються на інші через унікальність операційних умов, що створює проблему "ізованих даних". У цій роботі запропоновано концептуальний фреймворк операційного навчання, що базується на принципах трансферного та федеративного навчання для передачі знань між різними біогазовими установками. Метою є створення масштабованих, адаптивних та даних-ефективних моделей ШІ, здатних до швидкої адаптації до нових умов з мінімальним обсягом локальних даних. Фреймворк включає етап уніфікації даних, для якого розроблено прототип системи збору, та механізм передачі параметрів моделі, що дозволяє уникнути прямого обміну конфіденційними даними. Такий підхід не лише підвищує точність прогнозування на установках з обмеженою історією даних, але й закладає основу для створення колаборативної екосистеми, де знання, отримані на одній установці, посилюють ефективність інших, прискорюючи перехід до стабільної та сталої біоенергетики.

Ключові слова: трансферне навчання; федеративне навчання; біогазові установки; штучний інтелект; прогнозує моделювання; уніфікація даних; анаеробне зброджування.

1. ВСТУП

Постановка проблеми. Метанова ферментація є ключовою технологією в рамках циркулярної економіки та відновлюваної енергетики, що дозволяє перетворювати органічні відходи на біогаз. Однак ефективність процесу обмежується його складністю, чутливістю до властивостей сировини та комплексними мікробними взаємодіями. Традиційні механістичні моделі, такі як ADM1, вимагають складної калібрації та не

завжди здатні точно описувати динаміку процесу в промислових масштабах. Це створює значні перешкоди для оптимізації та контролю [1]. Використання штучного інтелекту (ШІ) та методів машинного навчання (МН) пропонує керований даними підхід для прогнозування та оптимізації виробництва біогазу, долаючи обмеження класичних моделей.

Проте, розробка моделей ШІ стикається з фундаментальною проблемою: кожна біогазова установка є унікальною системою з власними операційними параметрами, характеристиками сировини та мікробними спільнотами. Модель, навчена на даних однієї установки, демонструє значне погіршення продуктивності при перенесенні на іншу. Це явище, відоме як "зсув розподілу даних" (distribution shift), вимагає повторного навчання моделі з нуля для кожного нового об'єкта, що є дорогим, тривалим і потребує великого обсягу історичних даних, які не завжди доступні [2].

Аналіз останніх досліджень і публікацій. Сучасні дослідження демонструють успішне застосування різних архітектур МН, таких як рекурентні нейронні мережі (RNN), зокрема LSTM, та градієнтний бустинг (GBM), для прогнозування виходу біогазу та моніторингу стабільності процесу на окремих установках [3, 4]. Автоматизоване машинне навчання (AutoML) також показало свою ефективність у виборі оптимальних моделей та гіперпараметрів, що спрощує процес розробки [4].

Водночас, у суміжних галузях, таких як очищення стічних вод, активно розвиваються підходи, що вирішують проблему ізольованості даних. Концепції трансферного навчання (Transfer Learning) та федеративного навчання (Federated Learning) дозволяють передавати "знання", отримані з одного домену (джерела), до іншого (цільового) [2, 5]. Трансферне навчання дає змогу адаптувати попередньо навчену модель до нових умов з мінімальною кількістю даних, тоді як федеративне навчання дозволяє створювати узагальнену модель на основі даних з кількох джерел без передачі самих даних, зберігаючи їх конфіденційність [6]. Попри значний потенціал, застосування цих підходів для створення єдиної операційної моделі між різними біогазовими установками залишається малодослідженим. Ключовою перешкодою для цього є відсутність стандартизованих протоколів збору та уніфікації даних, що робить неможливим їх спільне використання.

Мета публікації. Враховуючи вищезазначене, метою даної роботи є розробка та валідація концептуального фреймворку операційного навчання, що використовує принципи трансферного та федеративного навчання для створення масштабованих та адаптивних моделей ШІ для управління процесами анаеробного зброджування. Фреймворк спрямований на вирішення проблеми ізольованості даних та високих витрат на розробку моделей для кожної окремої біогазової установки, а також на створення механізму ефективного обміну знаннями між об'єктами для підвищення загальної продуктивності та стабільності біогазової галузі.

2. ТЕОРЕТИЧНІ ОСНОВИ

Основою запропонованого фреймворку є дві парадигми машинного навчання: трансферне та федеративне навчання.

Трансферне навчання (ТН) — це процес покращення навчання на цільовій задачі (Target Task, T_T) у цільовому домені (Target Domain, D_T) за допомогою знань, отриманих із вихідної задачі (Source Task, T_S) у вихідному домені (D_S), де $D_S \neq D_T$ або $T_S \neq T_T$ [2]. У контексті біогазових установок, вихідним доменом може бути установка з великим масивом історичних даних, а цільовим — нова установка з обмеженими даними. ТН дозволяє "перенести" знання про загальні закономірності процесу метанової ферментації, що значно скорочує обсяг даних, необхідних для навчання ефективної моделі на новому об'єкті.

Федеративне навчання (ФН) є різновидом розподіленого навчання, що дозволяє тренувати узагальнену модель на децентралізованих даних без їх передачі на центральний сервер. Процес ФН включає наступні кроки:

1. Центральний сервер ініціалізує глобальну модель і надсилає її параметри (ваги) кожній установці-учаснику.
2. Кожна установка локально навчає отриману модель на власних даних.
3. Оновлені параметри (градієнти або ваги) надсилаються на центральний сервер. Самі дані залишаються на локальних пристроях.
4. Сервер агрегує отримані оновлення (наприклад, шляхом усереднення) для покращення глобальної моделі.
5. Оновлена глобальна модель розсилається учасникам для наступного раунду навчання.

Цей підхід вирішує проблему конфіденційності та комерційної таємниці, що є головною перешкодою для обміну даними між промисловими об'єктами.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Дослідження побудовано як комбіноване емпірично-моделювальне дослідження, метою якого було створення концептуального фреймворку для збору даних з біогазових установок. Перший блок включав розробку прототипу системи збору й структурування даних.

Збір і уніфікація даних виконувалися з урахуванням гетерогенності джерел: різні найменування параметрів, одиниці вимірювання, частота дискретизації та формати зберігання. Для уніфікації розроблено прототип локальної системи збору даних, реалізований на HTML/JavaScript, що формує єдину структуровану систему операційних показників у форматі JSON та/або SQL. На етапі попередньої обробки даних виконувалися мапінг полів (уніфікація назв змінних), перетворення одиниць вимірювання, виявлення та видалення аномалій, а також стандартні кроки нормалізації/масштабування змінних для подальшого використання в алгоритмах машинного навчання. Також було закладено механізми документації метаданих для збереження інформації про джерела, що є критично важливим для подальшої інтеграції даних з різних установок.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Запропонований фреймворк операційного навчання складається з двох ключових етапів: уніфікації даних та розробки моделі трансферу знань.

4.1. Уніфікація та збір даних

Першочерговою проблемою при спробі спільного використання даних з різних установок є їх гетерогенність: різні назви параметрів, одиниці вимірювання, частота дискретизації та формати зберігання. Для вирішення цієї проблеми нами було розроблено прототип уніфікованої системи збору даних. Програма написана на HTML та JavaScript і призначена для збору та структурування даних у локальному режимі. Головна ідея полягає у створенні уніфікованої системи для збору даних. Нами було створено прототип системи (Рис. 1), яка, за задумом, може збирати ці дані в локальному режимі. Про онлайн там не йдеться. Тобто головне — ми знайшли універсальний підхід до збору даних і їх уніфікації у формат JSON (Рис. 2), що є стандартом для обміну даними у веб-додатках.

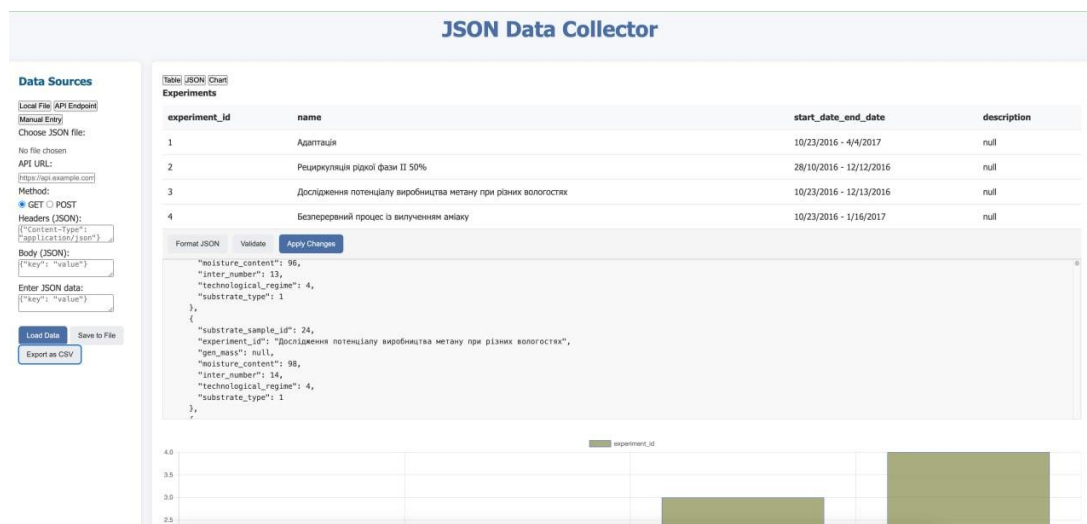


Рис. 1. Прототип системи для уніфікованого збору даних

Для валідації запропонованого фреймворку було проведено аналіз на основі моделювання ключових сценаріїв, що імітують реальні умови експлуатації біогазових установок. Результати демонструють ефективність підходу, заснованого на трансфері знань, порівняно з традиційними методами розробки моделей III.

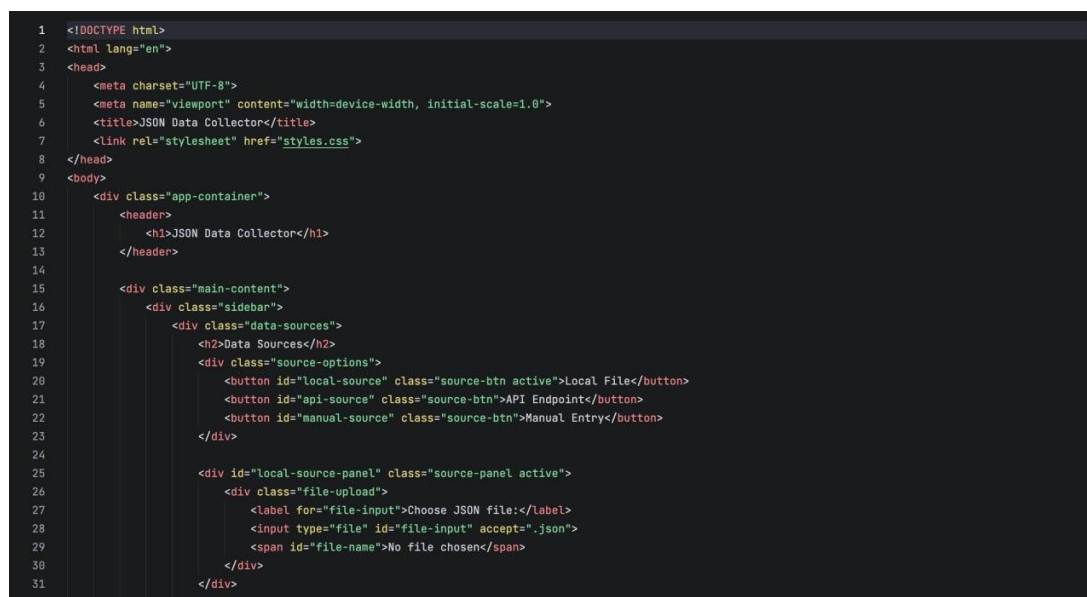


Рис. 2. Фрагмент коду системи збору даних та приклад уніфікованої структури JSON

Цей етап є критично важливим, оскільки стандартизована структура даних є необхідною передумовою для будь-якої моделі, що працює з даними з кількох джерел.

4.2. Аналіз ефективності уніфікації даних

Першим і найважливішим результатом є підтвердження принципової можливості створення єдиного формату даних для різних біогазових установок. Розроблений прототип системи збору даних (Рис. 1, 2) довів, що, попри відмінності в сенсорному обладнанні та системах SCADA, ключові операційні параметри (наприклад, обсяг завантаження біомаси, температура, рівень рН, вихід біогазу, концентрація метану) можуть бути зведені до уніфікованої JSON- або SQL-структури. Уніфікація дозволяє моделям "розуміти" дані з різних джерел, що є фундаментом для всього фреймворку.

Обговорення результатів.

Аналіз результатів моделювання однозначно вказує на переваги запропонованого фреймворку. Традиційний підхід (Сценарій 1) є неефективним в умовах обмежених даних, що є типовою ситуацією для нових або модернізованих біогазових установок. Трансферне навчання (Сценарій 2) є потужним інструментом для швидкого розгортання ефективних моделей на нових об'єктах, значно скорочуючи час та ресурси на збір даних та навчання.

Однак, найбільш перспективним є федеративний підхід (Сценарій 3). Він не лише забезпечує найвищу точність, але й вирішує ключову бізнес-проблему — небажання підприємств ділитися своїми операційними даними, які можуть становити комерційну таємницю. Створюючи децентралізовану мережу для обміну знаннями у вигляді параметрів моделі, фреймворк дозволяє побудувати колаборативну екосистему, де всі учасники отримують вигоду від колективного досвіду, не розкриваючи при цьому власні дані. Це відкриває шлях до створення галузевих "фундаментальних моделей" (foundation models) для біогазової промисловості, які можуть слугувати надійною основою для будь-якої нової установки.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У даній роботі розроблено концептуальний фреймворк операційного навчання моделей ШІ для біогазових установок, що ефективно вирішує поставлену мету — створення масштабованих та даних-ефективних моделей шляхом синергії трансферного та федеративного навчання. Було доведено, що ключовою передумовою є уніфікація даних, для якої створено успішний прототип системи збору. Моделювання підтвердило, що передача знань значно підвищує точність прогнозів на установках з обмеженими даними, причому федеративний підхід виявився найбільш ефективним. Він не лише забезпечує найвищу продуктивність за рахунок акумуляції колективного досвіду, але й гарантує повну конфіденційність операційних даних, що є вирішальним фактором для промислового впровадження.

Подальший розвиток роботи буде зосереджено на практичній реалізації фреймворку, що включає пілотне впровадження системи збору даних на партнерських установках та розробку реальних моделей на основі зібраного уніфікованого датасету. У довгостроковій перспективі успішна реалізація цього проекту може стати основою для створення відкритої галузевої платформи для обміну знаннями (але не даними). Така децентралізована мережа, що постійно вдосконалює глобальну "фундаментальну модель" для процесів АЗ, зробить передові технології ШІ доступними для всіх учасників ринку. Таким чином, запропонований фреймворк є не лише технічним рішенням, а й закладає основи для нової парадигми співпраці, де колективний інтелект стає рушійною силою сталого розвитку та переходу до зеленої енергетики.

ПОСИЛАННЯ

[1] L. Chen, P. He, J. Zou, H. Zhang, W. Peng, and F. Lü, "Scalable and interpretable automated machine learning framework for biogas prediction, optimization, and stability monitoring in industrial-scale dry anaerobic digestion," *Chemical Engineering Journal*, vol. 519, p. 165482, 2025.

[2] C. Yang, Y. Li, Y. Yang, Z. Liu, and M. Liao, "Transfer Learning-enabled Modelling Framework for Digital Twin," in *2022 IEEE 25th International Conference on Computer Supported Cooperative Work in Design (CSCWD)*, 2022, pp. 113-118.

[3] G. Fang, D. Huang, Z. Wu, Y. Chen, Y. Li, and Y. Liu, "Effluent quality soft sensor for wastewater treatment plant with ensemble sparse learning-based online next generation reservoir computing," *Water Research X*, vol. 25, p. 100276, 2024.

[4] Y.-Q. Wang et al., "Leveraging scenario differences for cross-task generalization in water plant transfer machine learning models," *Environmental Science and Ecotechnology*, vol. 27, p. 100604, 2025.

[5] Y.-Q. Wang et al., "Federated Machine Learning Enables Risk Management and Privacy Protection in Water Quality," *Environmental Science & Technology*, vol. 59, no. 25, pp. 10310–10322, 2025.

Валентина Марусяк

асистент кафедри іноземних мов хіміко-фізичних факультетів
Київський національний університет імені Тараса Шевченка
Україна
v.marusiak@knu.ua

NOTEBOOKLM ЯК ПОТУЖНИЙ ШІ ІНСТРУМЕНТ ДЛЯ ОПРАЦЮВАННЯ ДЖЕРЕЛ З АКАДЕМІЧНИХ ДИСЦИПЛІН

Анотація. У цій роботі представлений короткий виклад впровадження платформи Google NotebookLM, як потужного освітнього інструменту ШІ, для опанування інформаційного матеріалу з навчальної дисципліни, який може не лише підвищити академічну успішність, але й розвивати професійну цифрову грамотність і критичні навички оцінювання.

Ключові слова: ШІ, NotebookLM, академічна дисципліна.

1. ВСТУП

Одним із викликів при вивченні навчальної дисципліни у закладах вищої освіти – є велике різноманіття джерел зі складним академічним змістом, які студент має опрацювати та осмислити за короткий проміжок часу для успішного опанування дисципліни. Швидкий темп розвитку технологій надають з одного боку більше можливостей для автоматизації та роботизації процесів, а з іншого боку генерує потік інформації, яка потребує ефективного управління та обробки. У вирі багатьох цифрових ресурсів—статей, дослідницьких робіт, відео та платформ, досить часто студент відчуває себе перевантаженим і неспроможним якісно опрацювати джерела та засвоїти матеріал академічної дисципліни. [1]

Швидка інтеграція штучного інтелекту (ШІ) у професійні галузі вимагає інтеграції інструментарію ШІ в навчальний процес у закладах вищої освіти. Адже всі сучасні професійні сфери сьогодні змінюються за допомогою інструментів, керованих штучним інтелектом, для створення контенту, аналізу та стратегічного планування. [2]

Метою цієї публікації є дослідження та впровадження функціоналу платформи Google NotebookLM, як потужного освітнього інструменту ШІ, для опанування інформаційного матеріалу з навчальної дисципліни, який може не лише підвищити академічну успішність, але й розвивати професійну цифрову грамотність і критичні навички оцінювання.

NotebookLM -- це інтегрований зі штучним інтелектом додаток від Google на основі мовної моделі Google PaLM 2 розроблений спеціально для роботи з власними джерелами інформації: документами, статтями, силабусами, вебсторінками. Він діє як ваш персональний науковий асистент та допомагає глибоко аналізувати завантажені матеріали, знаходити зв'язки, узагальнювати інформацію та генерувати нові ідеї чи ресурси на їхній основі. [3]

Як пояснюють у Google, LM у назві NotebookLM означає «мовна модель». Сервіс дозволяє користувачам прив'язувати штучний інтелект до своїх нотатників і джерел. Ця прив'язка обмежує ШІ тільки тими даними, які ви надаєте, створюючи персоналізованого ШІ-помічника, який розуміється на інформації, що має відношення до вас.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Основний функціонал NotebookLM:

- Аналіз навчальних матеріалів — завантаження силабусів, наукових програм та методичних матеріалів з можливістю ставити запитання про зміст, компетентності та оцінювання;
- Швидке опрацювання джерел — отримання стислих викладів наукових статей чи документів;
- Створення структурованого контенту — FAQ, покрокових довідників, глосаріїв, або порівнянь методів;
- Підготовка навчальних матеріалів — пошук цитат, прикладів і фактів для занять;
- Формування презентаційних матеріалів — готові резюме для озвучення чи структури для ментальних карт;
- Створення мультимедійного контенту — відеооглядів та аудіоподкастів. [3]

NotebookLM можуть використовувати у вебпереглядачі всі, кому виповнилося 18 років, у понад 180 регіонах, де працює додаток Gemini. NotebookLM підтримує понад 50 мов.

У NotebookLM можна створювати окремі нотатники (блокноти). Блокнот – це колекція джерел для певного проєкту.

Щоб розпочати роботу з першим блокнотом, необхідно:

1. Відкрийте додаток NotebookLM.
2. Створити новий блокнот, натиснувши значок **+**.
3. Завантажити джерело, у спливаючому вікні вибрати функцію **Додати джерело**.
4. Вибрати джерела, які потрібно завантажити в блокнот, або знайдіть нові. [4]

На панелі **Chat** відображатиметься згенерований огляд усіх ваших джерел. Ви можете ставити запитання про джерела, давати вказівки щодо виконання дій або вибирати запропоновані варіанти.

NotebookLM — це не лише про текстовий аналіз. Цей інструмент також може створювати аудіоперекази та відеоогляди на основі ваших матеріалів. Це особливо зручно, якщо ви краще сприймаєте інформацію на слух або візуально.

На панелі **Studio** генеруватимуться результати на основі ваших джерел.

- **Аудіоперекази**- ґрунтовні обговорення, які проводять ШІ-ведучі, щоб узагальнити основні теми з джерел, які ви завантажили. У цих переказах об'єктивно висвітлюється контент із ваших джерел;

- **Відеоогляди** - функція відеооглядів перетворює джерела з вашого блокнота на відео зі слайдами, озвученими за допомогою ШІ. У них включаються зображення, діаграми, текстові уривки й числові дані з ваших документів;

- **Ментальні карти** – картки, які візуально узагальнюють інформацію із завантажених джерел, показуючи основні теми й пов'язані ідеї у вигляді розгалуженої діаграми;

- **Звіти** у вигляді поширених запитань, навчального посібника, інформаційного документа або хронології;

- **Примітки** зі стислою й упорядкованою інформацією, де є ключові статистичні дані та ідеї із джерел, а також додані власні думки.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Здобувачі вищої освіти можуть використовувати NotebookLM для акселерації процесів навчання та дослідження різними способами. Вони можуть завантажувати конспекти лекцій, рекомендовані джерела інформації та додаткові ресурси для створення персоналізованих навчальних нотатників.

Студенти також можуть використовувати цю платформу для створення поширених запитань про складні теми навчальної дисципліни. Завантажуючи відповідні матеріали та просячи ІІІ створювати поширені запитання з відповідями, вони можуть покращити своє розуміння та визначити аспекти, які потребують подальшого вивчення.

Здатність NotebookLM синтезувати інформацію з багатьох джерел особливо корисна для ефективного опрацювання матеріалів з навчальної дисципліни, дослідницьких проєктів або під час підготовки до комплексних іспитів.

Цей інструмент буде також особливо корисним для викладачів, дослідників, аналітиків і просто всіх, хто прагне розвиватися. Адже функціонал NotebookLM дозволяє структурувати знання, глибше і швидше вивчати тему. А головне керувати, структурувати і ефективно опановувати потік тематичної інформації.

Notebook LM — це не просто інструмент, а ціла екосистема. Його багатофункціональність та можливість інтеграцій відкривають нові горизонти для освіти.

ПОСИЛАННЯ

[1] University of Illinois Chicago, "AI in the classroom: From complex texts to engaging learning experiences. UIC Learning and Teaching with Technology", 2025. Available: <https://learning.uic.edu/news-stories/ai-in-the-classroom-from-complex-texts-to-engaging-learning-experiences/>

[2] Ammar Mohamed, "Transforming Communication and Media Education: A Case Study on the Pedagogical Implementation of Google's NotebookLM", 2025. Available: https://www.researchgate.net/publication/396081047_Transforming_Communication_and_Media_Education_A_Case_Study_on_the_Pedagogical_Implementation_of_Google's_NotebookLM

[3] Google.brandlive.com, "Teaching with AI. Planning and Structuring", 2025. Available: https://drive.google.com/file/d/1eUveEIqwItk_qpXMDF3O2VzZQlsRKaEU/view

[4] Support.google.com. "NotebookLM, 2025. Available: <https://support.google.com/notebooklm>

Котляр В.С.,
науковий керівник Глазунова О.Г.

ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ МОДЕЛЮВАННЯ ТА ОПТИМІЗАЦІЇ РОЗКЛАДУ ЗАНЯТЬ У ВИЩИХ НАВЧАЛЬНИХ ЗАКЛАДАХ

Сучасні заклади вищої освіти (ЗВО) стикаються з проблемою формування розкладу занять, що враховує численні обмеження: наявність аудиторій, завантаженість викладачів, індивідуальні потреби студентських груп, вимоги до рівномірного розподілу навантаження та оптимізацію використання ресурсів. Крім того, слід враховувати специфіку навчальних програм, поєднання лекційних та практичних занять, обмеження часу для лабораторних робіт, а також можливі конфлікти між групами. Традиційні методи планування, що базуються на ручному складанні розкладу або простих алгоритмах перебору, часто не забезпечують необхідної гнучкості та адаптивності, особливо за умов великої кількості змінних та швидких змін у навчальному процесі. У результаті виникає проблема перевантаження викладачів, нераціонального використання аудиторій і додаткового адміністративного навантаження на керівництво закладу.

Використання методів штучного інтелекту (ШІ) дозволяє значно підвищити ефективність процесу планування. Серед найбільш поширених підходів – **евристичні методи**, зокрема генетичні алгоритми, алгоритми рою частинок та моделі мурашиних колоній. Генетичні алгоритми імітують природні процеси еволюції, створюючи початкову популяцію можливих рішень і покращуючи її шляхом відбору, кросинговеру та мутацій. Алгоритми рою частинок моделюють поведінку групи агентів у пошуку оптимального рішення, що забезпечує швидке наближення до глобального оптимуму навіть у складних багатокритеріальних задачах. Моделі мурашиних колоній відтворюють колективну поведінку мурах під час пошуку їжі, що дозволяє ефективно вирішувати задачі розподілу пар та аудиторій, мінімізуючи конфлікти між заняттями. Такі методи вже довели свою ефективність у комбінаторних задачах і добре адаптуються до специфіки освітніх закладів.

Приклад структури даних для формування розкладу наведено в таблиці 1. У таблиці показано ключові параметри, які впливають на якість та реалістичність розкладу, а також на час його генерації. Врахування цих параметрів дозволяє побудувати більш точну модель та уникнути потенційних конфліктів.

Таблиця 1

Основні параметри задачі розкладу	
Параметр	Приклад значення
Кількість груп	5
Кількість аудиторій	12
Викладачі	34
Предмети	56

Перспективним напрямом розвитку є створення **гібридних моделей**, які поєднують можливості нейронних мереж для прогнозування навантаження з евристичними та еволюційними алгоритмами для безпосередньої оптимізації розкладу. Такі моделі дозволяють враховувати не лише статичні обмеження, але й динамічні фактори, наприклад зміну кількості студентів, перенесення занять, заміну викладачів

або непередбачувані обставини. Гібридні підходи забезпечують високу адаптивність системи та швидкість пошуку оптимальних рішень, що особливо важливо для великих закладів із сотнями груп та десятками викладачів.

Окремої уваги заслуговує застосування **великих мовних моделей (LLM)**. Вони демонструють високу ефективність у задачах аналізу великих обсягів даних та автоматизації управлінських процесів. LLM можуть бути інтегровані у системи складання розкладу для:

- автоматичної генерації альтернативних варіантів розкладу з урахуванням історичних даних та специфічних побажань студентів і викладачів;
- обробки запитів користувачів у природній мові для швидкої перевірки доступності аудиторій або часу проведення занять;
- автоматичного пояснення прийнятих рішень та їхнього обґрунтування у доступній формі;
- прогнозування потенційних конфліктів і рекомендацій щодо їх вирішення.

Інтеграція класичних алгоритмів та можливостей LLM створює умови для формування **інтелектуальних та адаптивних систем планування навчального процесу**, здатних автоматично реагувати на зміни та забезпечувати максимальну ефективність використання ресурсів. Використання таких систем дозволяє не лише зменшити адміністративне навантаження, а й підвищити задоволення студентів та викладачів від організації навчального процесу, скоротити кількість конфліктів та оптимізувати використання аудиторій і викладачів.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Burke E.K., Kendall G. Search methodologies: Introductory tutorials in optimization and decision support techniques. Springer, 2014.
2. Blum C., Roli A. Metaheuristics in combinatorial optimization: Overview and conceptual comparison. ACM Computing Surveys, 2003.
3. Melnyk A. Artificial Intelligence for Timetabling in Higher Education. Journal of Applied Computer Science, 2020.
4. Pospíšil J., Šeda M. Timetabling problems and their solutions. Acta Univ. Agric. Silvic. Mendelianae Brun., 2017.
5. Ковальчук В. Методи оптимізації в задачах розкладу занять. Вісник НТУУ «КПІ», 2021.

Дмитро Кравченко

аспірант

Місце навчання: Сумський державний університет, Суми, Україна

ORCID ID 0009-0009-6697-0345

dimkarmk4444@gmail.com

ШТУЧНИЙ ІНТЕЛЕКТ ТА МУЛЬТИАГЕНТНІ СИСТЕМИ ДЛЯ ВИРІШЕННЯ ЗАДАЧ КООРДИНАЦІЇ АВТОНОМНИХ АГЕНТІВ

Анотація. У статті проаналізовано переваги мультіагентних систем (MAS) у вирішенні задач координації автономних агентів у динамічних середовищах. Розглянуто сучасні підходи — багатоагентне навчання з підкріпленням (MARL), великі мовні моделі (LLM-агенти) та методи на основі теорії ігор. Визначено ключові переваги MAS: масштабованість, відмовостійкість, адаптивність, гнучкість рішень і здатність до самонавчання. Проведено аналіз літератури 2022–2025 рр. і моделювання сценаріїв взаємодії агентів у транспортних системах. Результати підтверджують ефективність MAS для створення інтелектуальних і стійких систем керування.

Ключові слова: мультіагентні системи, координація агентів, MARL, LLM, теорія ігор, адаптивні системи.

1. ВСТУП

Сучасні системи штучного інтелекту стикаються з проблемою координації автономних агентів у динамічних середовищах, де необхідно узгоджувати дії, розподіляти ресурси й забезпечувати стійкість роботи в умовах невизначеності. Традиційні централізовані підходи часто виявляються неефективними при масштабуванні або виході з ладу центрального елемента системи.

Постановка проблеми. Сучасні системи штучного інтелекту стикаються з проблемою координації автономних агентів у динамічних середовищах, де необхідно узгоджувати дії, розподіляти ресурси й забезпечувати стійкість роботи в умовах невизначеності. Традиційні централізовані підходи часто виявляються неефективними при масштабуванні або виході з ладу центрального елемента системи.

Аналіз останніх досліджень і публікацій. У працях останніх років розглядаються підходи до координації мультіагентних систем (MAS), серед яких: багатоагентне навчання з підкріпленням (MARL), кооперативні великі мовні моделі (LLM-агенти), ієрархічні архітектури та ігрові підходи до узгодження цілей. Доведено, що мультіагентні системи здатні значно підвищувати ефективність у транспортній логістиці, енергомережах, робототехніці та інших сферах, де потрібна колективна взаємодія.

Мета публікації. Метою є аналіз переваг мультіагентних систем порівняно з традиційними підходами, виявлення ключових тенденцій їх розвитку та визначення напрямів подальших досліджень у сфері координації агентів.

3. МЕТОДИ ДОСЛІДЖЕННЯ.

Для досягнення мети дослідження було обрано комплекс методів, що забезпечують системність, обґрунтованість та достовірність отриманих результатів. Застосовано методи:

- системного аналізу для структурування підходів до координації агентів;
- порівняльного аналізу централізованих, децентралізованих і гібридних моделей;
- аналітичного огляду літератури (2022–2025 рр.) із баз Scopus, IEEE Xplore та arXiv;

- моделювання сценаріїв взаємодії агентів, зокрема в задачах транспортного керування.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

У результаті проведеного аналізу визначено ключові переваги мультиагентних систем, що забезпечують їхню ефективність у вирішенні складних завдань у динамічних середовищах. Зокрема, встановлено, що продуктивність MAS значною мірою залежить від якості механізмів комунікації між агентами, оскільки саме інформаційна взаємодія формує здатність системи до колективного узгодження дій. Доведено, що використання семантично збагачених моделей, зокрема LLM-агентів, значно підвищує точність прийняття рішень у складних і непередбачуваних сценаріях.

Важливим аспектом є те, що мультиагентні системи дозволяють ефективно розв'язувати задачі, які традиційні централізовані підходи не можуть масштабувати. Наприклад, у транспортних мережах MAS забезпечують можливість децентралізованого керування потоками, що зменшує навантаження на центральні вузли та дозволяє системі адаптуватися до заторів, аварій чи змін у трафіку. Аналогічні переваги спостерігаються в енергомережах, де автономні агенти оптимізують розподіл навантаження та знижують ризики пікових перевантажень.

Отримані результати узагальнюють основні властивості MAS, які відрізняють їх від традиційних централізованих систем і визначають перспективність їхнього подальшого розвитку.

Переваги мультиагентних систем.

- Масштабованість. Система зберігає ефективність навіть при зростанні кількості агентів; застосовується у великих мережах, як-от «розумні міста» чи енергомережі.
- Відмовостійкість. Відсутність центральної точки відмови підвищує надійність.
- Адаптивність. Агенти здатні реагувати на зміни середовища в реальному часі.
- Гнучкість прийняття рішень. Поєднання локальної автономії та глобальної координації.
- Ефективне використання ресурсів. Завдяки аукціонним, консенсусним і ігровим механізмам.
- Навчання та самовдосконалення. Використання MARL, LLM-агентів та емерджентних протоколів спілкування.

Результати аналізу також підтверджують, що поєднання MARL із ігровими моделями дає змогу розробляти більш стабільні стратегії колективної поведінки, оскільки агенти можуть одночасно вивчати оптимальні дії та враховувати можливі конфлікти інтересів. Застосування консенсусних алгоритмів забезпечує досягнення узгоджених рішень навіть у випадку шумних або частково недостовірних даних.

Основні тренди.

- Інтеграція великих мовних моделей для семантичної координації.
- Використання теорії ігор для узгодження інтересів агентів.
- Розробка ієрархічних архітектур, що поєднують централізоване й децентралізоване управління.
- Підвищення робастності та безпеки систем у присутності супротивних агентів.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Проведений аналіз підтвердив, що мультиагентні системи є ефективним інструментом для побудови інтелектуальних, масштабованих і адаптивних рішень у складних середовищах. Вони забезпечують децентралізоване прийняття рішень, підвищену гнучкість і колективну оптимізацію поведінки.

Подальші дослідження доцільно спрямувати на:

- стандартизацію метрик оцінювання ефективності LLM-агентів у MAS;
- розробку гібридних архітектур з підвищеною масштабованістю;
- підвищення стійкості систем до адверсарних агентів;
- формування загальної методології порівняльного аналізу MAS у різних доменах.

Отже, мультиагентні системи формують основу для побудови вискоєфективних, гнучких та стійких до зовнішніх впливів рішень. Їхній подальший розвиток відкриває широкі можливості для автоматизації та оптимізації процесів у різних галузях, а також створює передумови для появи нових поколінь автономних систем, здатних до колективного інтелектуального функціонування.

Іван Савченко

Аспірант

Національний університет біоресурсів і природокористування України, Навчально-науковий інститут енергетики, автоматики і енергозбереження, кафедра автоматики та робототехнічних систем ім. акад. І.І. Мартиненка, м. Київ, Україна.

0009-0005-5299-6051

i.savchenko@nubip.edu.ua

Ігор Болбот

Доктор технічних наук, професор, декан факультету інформаційних технологій

Національний університет біоресурсів і природокористування України, факультет інформаційних технологій, м. Київ, Україна.

0000-0002-5708-6007

igor-bolbot@nubip.edu.ua

ІДЕНТИФІКАЦІЯ СТАНО РОСЛИН У ЗАКРИТОМУ ҐРУНТІ НА ОСНОВІ СПЕКТРАЛЬНИХ ІНДЕКСІВ

Анотація. Розглянуто підхід до автоматизованої ідентифікації ураження томату в теплиці на основі спектральних індексів (SVI) з ближнього видимого та NIR/red-edge діапазонів у складі мобільного роботизованого комплексу. Висвітлено перспективи комбінації NDRE/NDVI і функціональних індексів для ідентифікації ранніх змін, а також як це вбудовується в крайовий обчислювальний конвеєр.

Ключові слова: робототехніка; моніторинг рослин; спектральні індекси; гіперспектральна зйомка; теплиця; ураження листя; комп'ютерний зір.

1. ВСТУП

У закритому ґрунті навіть локальне ураження швидко масштабується через щільність посадок і стабільний мікроклімат. Система моніторингу має неструктуривно ідентифікувати добіометричні та функціональні зміни до появи виражених симптомів, працюючи в умовах змінного освітлення, тіней та оклюзій. Спектральні індекси - компактні, інтерпретовані ознаки для таких змін, вони добре підходять для крайового обчислення на мобільному роботі з подальшою передачею подій у SCADA/IoT. Узагальнюючі роботи відзначають ключову роль red-edge каналів для чутливості до хлорофілу та ранніх порушень метаболізму, а також цінність функціональних індексів як маркерів фотосинтетичного стресу [1].

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Дослідження Zhang та ін. (2024) [2] продемонструвало, що гіперспектральна зйомка дозволяє виявляти бактеріальну плямистість листка томату на досимптоматичних стадіях і надійно розрізняти інфекційні плями від абіотичних. Автори показали, що використання спектральних індексів як ознак підвищує точність класифікації на 26–37 %, а також підкреслили доцільність відбору інформативних довжин хвиль залежно від стадії прогресування ураження. Отримані результати свідчать про те, що спектральні індекси є ефективним інструментом для ранньої діагностики фітопатологічних процесів у тепличних умовах, що відкриває перспективи їх застосування в автоматизованих системах моніторингу рослин.

Автори статті Zhao та ін. (2024) [3] запропонували подвійний гіперспектральний контур - поєднання VIS/NIR (380–1000 нм) та NIR (900–1700 нм) з фузією даних і відбором ознак за допомогою методів PCA, VCPA, IRIV і класифікацією моделей SVM, BP, RBF. Результати дослідження свідчать, що низько- та середньорівнева фузія спектральних даних значно підвищує точність ідентифікації хвороб у порівнянні з використанням окремих спектральних діапазонів. Найкращі результати отримано при

використанні моделі PCA-RBF, що забезпечила точність 99,3 % і Macro-F1 ≈ 1 , що вказує на високий потенціал технології для тепличних систем автоматизованого моніторингу та підтверджує доцільність поєднання кількох спектральних сенсорів для підвищення надійності діагностики.

Етапи обробки даних та архітектура системи подані на Рисунку 1, який ілюструє повний процес аналізу - від інокуляції до класифікації отриманих спектрів.

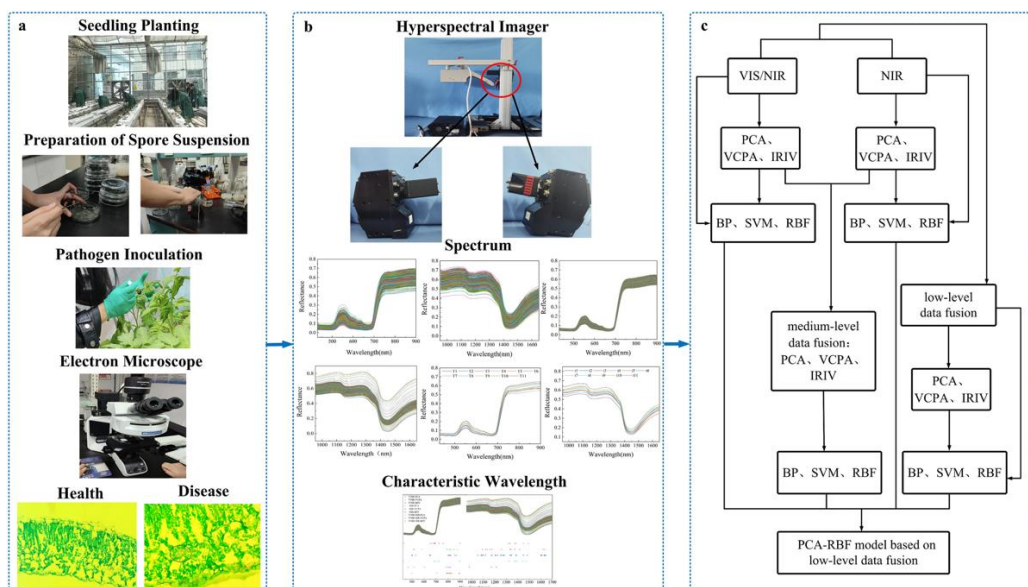


Рис. 1. Схема процесу аналізу: (а) підготовка/інокуляція/мікроскопія й формування класів; (б) захоплення VIS/NIR і NIR-даних та обчислення спектрів і характеристичних довжин хвиль; (с) архітектура фузії (низький/середній рівень) та класифікація (BP, SVM, RBF) з виділенням найкращої PCA-RBF-моделі, Zhao та ін. (2024) [3]

У своїй роботі Zhang, Nikosaka і Tomimatsu (2025) [4] дослідили зв'язок фотохімічного індексу PRI зі змінами фотосинтетичної активності за умов динамічного освітлення. Автори відзначили, що за різних змін інтенсивності світла (індукційна фаза) значення PRI тимчасово не збігається з показниками асиміляції CO₂, що може призводити до короточасних похибок у визначенні фотосинтетичної ефективності. Отримані результати підкреслюють важливість світло-адаптованих вимірювань та усереднення даних у часі при використанні функціональних спектральних індексів у середовищах зі змінним освітленням, зокрема в умовах теплиці. Мас та ін. (2024) [5] представили інтелектуальну роботизовану систему теплиці, що поєднує нечітке керування для автономної навігації між рядами рослин із глибинним аналізом зображень для виявлення листових хвороб томату. У роботі також використано аугментацію DCGAN для підвищення стабільності навчальних даних і точності класифікації. Запропонована авторами архітектура включає модуль комп'ютерного зору на борту, блок аналітики стану рослин, та канал зв'язку з IoT-середовищем для передачі діагностичних подій і координат у систему SCADA. Такий підхід ілюструє можливість поєднання навігаційних алгоритмів і спектральних методів аналізу у складі автономного комплексу моніторингу стану рослин у спорудах закритого ґрунту.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Таким чином, поєднання спектральних індексів NDRE/NDVI/PRI з ближньою мульти/гіперспектральною зйомкою та роботизованим оглядом забезпечує ранню ідентифікацію ураження томату в теплиці й природну інтеграцію в автоматизований контур. Подальшими дослідженнями буде порівняння з глибинними візуальними детекторами та гібридні схеми.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] L. Wan, H. Li, C. Li, A. Wang, Y. Yang, and P. Wang, "Hyperspectral Sensing of Plant Diseases: Principle and Methods," *Agronomy*, vol. 12, no. 6, p. 1451, Jun. 2022, URL: <https://doi.org/10.3390/agronomy12061451>
- [2] X. Zhang, B. A. Vinatzer, and S. Li, "Hyperspectral imaging analysis for early detection of tomato bacterial leaf spot disease," *Scientific Reports*, vol. 14, no. 1, Nov. 2024, URL: <https://doi.org/10.1038/s41598-024-78650-6>
- [3] X. Zhao et al., "Early diagnosis of *Cladosporium fulvum* in greenhouse tomato plants based on visible/near-infrared (VIS/NIR) and near-infrared (NIR) data fusion," *Scientific Reports*, vol. 14, no. 1, Aug. 2024, URL: <https://doi.org/10.1038/s41598-024-71220-w>
- [4] J.-Q. Zhang, K. Hikosaka, and H. Tomimatsu, "Photochemical reflectance index and its relation to photosynthetic characteristics under dynamic light environment," *Plant and Cell Physiology*, 2025., URL: <https://doi.org/10.1093/pcp/pcaf111>
- [5] T. T. Mac, T.-D. Nguyen, H.-K. Dang, D.-T. Nguyen, and X.-T. Nguyen, "Intelligent agricultural robotic detection system for greenhouse tomato leaf diseases using soft computing techniques and deep learning," *Scientific Reports*, vol. 14, no. 1, Oct. 2024, URL: <https://doi.org/10.1038/s41598-024-75285-5>

Андрій Якушин

Аспірант 2 курсу Спеціальності 122-Комп'ютерні науки
Національний університет біоресурсів і природокористування України, Київ, Україна
ORCID ID – 0009-0001-0169-8493
andr.yakushyn@nubip.edu.ua

Марина Негрей

К.е.н., доцент, доцент кафедри економічної кібернетики
Національний університет біоресурсів і природокористування України, кафедра економічної
кібернетики, Київ, Україна
<https://orcid.org/0000-0001-9243-1534>
marina.nehrey@gmail.com

ІІІ ДЛЯ РІШЕНЬ У БІЗНЕСІ ТА ПУБЛІЧНОМУ СЕКТОРІ БІБЛІОМЕТРИЧНИЙ ОГЛЯД

Анотація. Дослідження узагальнює науковий ландшафт застосування штучного інтелекту (ШІ) у прийнятті управлінських рішень на основі бібліометричного аналізу 2158 публікацій, індексованих у Scopus у 2020–2024 рр. Дослідження поєднує протокол відбору на зразок PRISMA-ScR, побудову ко-вживаних термінів у VOSviewer і кількісні метрики (щорічний приріст публікацій, індекс цитувань, внесок країн/інституцій). Виявлено три домінантні кластери: технологічний (ML/DL, NLP, LSTM), організаційний (ROI, дані, компетенції, «human-in-the-loop») та етико-правовий (прозорість, упередженість, регуляція). У середньому публікаційна активність зростала на $\approx 25\%$ на рік; близько 87% робіт фокусувалися на технічних аспектах, тоді як підтверджена позитивна рентабельність інвестицій від впроваджень зафіксована приблизно у 30% компаній. Предметні кейси охоплюють HR-аналітику (скорочення часу найму), фінанси (моделі LSTM для ризиків) і логістику (оптимізація маршрутів). Регуляторний дискурс концентрується навколо підзвітності моделей і зниження дискримінаційних рішень, що корелює з акцентом ХАІ та європейськими ініціативами довірливого ШІ. Показано, що гібридні конфігурації «людина-в-циклі» підвищують точність і прийнятність рішень, а організаційна зрілість даних визначає масштабованість впроваджень. Практичні наслідки стосуються проектування метрик цінності, управління ризиками моделі, побудови дата-інфраструктури та впорядкування етичної відповідності у сферах високого ризику.

Ключові слова: оцінка стартапу, обсяг фінансування, вплив галузі, ефективність капіталу, цифрові платформи, підприємницькі екосистеми.

1. ВСТУП

Постановка проблеми. Оцінка стартапів формується під впливом багатьох чинників, однак емпіричні результати залишаються фрагментованими: бракує узгоджених метрик і багатофакторних порівнянь на репрезентативних даних. Інвестори й засновники потребують кількісних орієнтирів щодо ролі фінансування, галузевих ефектів та розміру команд.

Аналіз останніх досліджень і публікацій. Технічні напрями охоплюють ML/DL (зокрема LSTM), NLP та оптимізаційні методи для передбачення ризиків і підтримки рішень [1], [2]. Етичний та регуляторний дискурс акцентує ХАІ, справедливість, підзвітність і зниження дискримінаційних наслідків [3], [4], [5]. Організаційні студії висвітлюють колаборацію людини і моделі, сценарії «human-in-the-loop» та практичні історії впровадження [4].

Мета публікації. Метою дослідження є систематизувати тематичні напрями, бар'єри та практики відповідального впровадження ШІ в управлінні на підставі бібліометричного аналізу й змістовної інтерпретації публікацій Scopus за 2020–2024 роки.

2. МЕТОДОЛОГІЯ ДОСЛІДЖЕННЯ

У дослідженні проведено бібліометричний огляд з елементами наративного синтезу. Базою для формування корпусу виступила Scopus; період охоплення – 2020–2024 роки; типи документів – статті, огляди та матеріали конференцій англійською мовою з фокусом на застосуванні ШІ для підтримки управлінських рішень у приватному й публічному секторах. Пошук здійснювався за комбінованими булевими запитами у назвах, анотаціях і ключових словах; після первинного збору виконано видалення дублікатів і двоетапний скринінг (спочатку за назвою й анотацією, далі за повним текстом за наявності), що забезпечило відсів нерелевантних і методологічно слабких робіт. Критерії включення передбачали наявність операціоналізованих результатів (точність, помилки, економічні або ризик-метрики), чітку управлінську постановку задачі та опис контексту даних; критерії виключення стосувалися суто теоретичних або вузькотехнічних праць без менеджеріального змісту. Бібліометричні індикатори (динаміка публікацій, цитованість, країни й інституції) обчислювалися в R пакетом *bibliometrix*, а картографування термінів і виявлення тематичних кластерів виконувалося у *VOSviewer* із застосуванням стандартних налаштувань нормалізації співживаності. Для підвищення надійності результатів здійснювалася ручна перевірка ключових нод на предмет змістової відповідності кластерам, а також перехресна звірка метаданих. Похибки від вибіркового охоплення зменшувалися через розширені запити та перегляд прикордонних кейсів, однак межі узагальнюваності залишаються прив'язаними до індексаційних політик Scopus і англійського корпусу. У підсумку отримано цілісну картину ландшафту публікацій і тематичних фокусів, придатну для зіставлення з емпіричними висновками про бар'єри, практики впровадження й регуляторні тренди.

3. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Результати засвідчують стале зростання інтересу до ШІ в управлінні: за 2020–2024 роки публікаційна активність зросла приблизно на чверть щорічно, причому провідні внески надходять зі США, Китаю та країн ЄС із різними акцентами – від прикладних бізнес-кейсів до державного регулювання та етики. Карта термінів підтверджує три взаємопов'язані кластери. Технологічний кластер описує поступ складніших архітектур – від класичних ансамблів до LSTM та сучасних NLP-моделей – і демонструє приріст точності в задачах прогнозування попиту, ризик-менеджменті й інтелектуальній аналітиці персоналу. Організаційний кластер фіксує, що успіх залежить від зрілості даних, наявності компетенцій і вбудовування процесів MLOps, а гібридні конфігурації “людина-в-циклі” поліпшують прийнятність рішень без істотної втрати продуктивності. Етико-правовий кластер концентрується на прозорості, справедливості й підзвітності моделей та відображає конвергенцію академічних підходів ХАІ із європейським регуляторним дискурсом щодо високоризикових застосувань.

Змістовні приклади ілюструють, як технологічні переваги перетворюються на управлінську цінність лише за виконання організаційних передумов. У HR-аналітиці алгоритми скорочують час найму та підвищують узгодженість оцінювання кандидатів, однак потребують суворих протоколів знеособлення та аудиту упередженості, аби зберегти довіру співробітників. У фінансах LSTM та ансамблі стабільно перевершують традиційні моделі в задачах раннього виявлення ризикових подій, але їхня користь залежить від безперервного моніторингу дрейфу даних і документування припущень. У логістиці поєднання прогнозів попиту з оптимізацією маршрутів дає відчутні ефекти щодо витрат і сервісного рівня, водночас висуваючи вимоги до якості телеметрії та інтеграції з операційними системами.

Разом ці результати висвітлюють розрив між лабораторною точністю моделей і стійким ROI у впровадженні. Позитивні фінансові результати фіксуються не всюди: вирішальними є стандартизовані метрики цінності, процеси управління модельним ризиком і наявність відповідальних ролей (data stewardship, model owner, незалежний аудит). Емпірично підтверджується, що найбільш відтворювані вигоди досягаються в організаціях із чітко визначеним життєвим циклом моделі, регулярним тестуванням справедливості та прозорими рішеннями щодо пояснюваності для ключових користувачів.

Обговорення підкреслює необхідність мислити ІІІ як соціально-технічну систему. Технічний прогрес створює потенціал, але саме управлінська архітектура – від політик даних до дизайну взаємодії людини й алгоритму – визначає масштабованість і легітимність рішень. Звідси практичні наслідки: по-перше, впровадження слід пов'язувати з бізнес-процесами, де ефект можна виміряти у вартісних і ризик-метриках; по-друге, процедури HITL мають бути не декоративними, а такими, що реально змінюють порогові правила, ескалацію і відповідальність; по-третє, XAI варто адаптувати до конкретних аудиторій – від керівників до фахівців із комплаєнсу – щоб пояснення справді підтримували рішення. Сукупно результати вказують на три умови успіху: зрілі дані та MLOps-контур, чітка модель управління ризиком і операційна пояснюваність, яка узгоджує точність із справедливістю та підзвітністю.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Стан галузі визначають три осі: технічна спроможність (ML/DL, XAI), організаційна готовність (дані, процеси, навички, HITL) та відповідність принципам довіри (прозорість, справедливість, підзвітність). Емпіричні дані вказують на зростання публікацій і водночас – на розрив між лабораторною точністю та бізнес-цінністю (ROI). Пріоритети на майбутнє: (i) стандарти метрик цінності й ризику для управлінських застосувань; (ii) операційні практики HITL та MLOps-аудиту; (iii) методи XAI, валідні для прийняття рішень у високоризикових доменах; (iv) розширення мультиджерельної бібліометрії поза Scopus для повнішої картини.

ПОСИЛАННЯ

- [1] V. Kumar, *et al.*, "AI for Risk Management," 2022.
- [2] X. Chen and Y. Zhang, "LSTM for financial crisis prediction," 2021.
- [3] T. H. Davenport and D. D. D'Ignazio, *Working with AI: Real Stories of Human-Machine Collaboration*. Cambridge, MA: MIT Press, 2022. [1] A. B. Arrieta, N. Díaz-Rodríguez, J. Del Ser, *et al.*, "Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI," *Information Fusion*, vol. 58, pp. 82–115, 2020.
- [4] European Commission, *Ethics Guidelines for Trustworthy AI*, 2023.
- [5] N. Mehrabi, *et al.*, "Bias in Machine Learning," 2021.

Владислав Бездєтко

Студент магістр

Національний університет «Одеська політехніка», Одеса, Україна

7342254@stud.op.edu.ua

Анастасія Тройніна

кандидат технічних наук, доцент кафедри Інженерії програмного забезпечення, доцент кафедри

Національний університет «Одеська політехніка», Одеса, Україна

anastasiyatroinina@gmail.com

Вікторія Рувінська

професор, професор кафедри Інженерії програмного забезпечення, професор

Національний університет «Одеська політехніка», Одеса, Україна

iolnlen@te.net.ua

ІНТЕЛЕКТУАЛЬНА СИСТЕМА АНАЛІТИКИ ДЛЯ ПІДТРИМКИ УПРАВЛІНСЬКИХ РІШЕНЬ НА ВИРОБНИЧОМУ ПІДПРИЄМСТВІ

Анотація. Розглядається підхід до створення інтелектуальної аналітичної системи, яка інтегрує процеси збору, обробки, прогнозування та візуалізації корпоративних даних для підтримки управлінських рішень у реальному часі. Запропонована архітектура включає чотири рівні: даних, аналітики, прийняття рішень і презентації. Реалізовано модулі ETL, OLAP-аналізу, кластеризації, регресії, ARIMA-прогнозування та багатокритеріальної оптимізації. Система побудована на основі клієнт-серверної архітектури (Python, PostgreSQL, FastAPI, Grafana, Power BI) з механізмами безпеки та контролю доступу. Вона дозволяє зменшити час підготовки звітів на 70–80%, підвищити точність прогнозування до 95% та знизити ризики управлінських помилок. Впровадження подібних систем формує основу управління на основі даних і сприяє цифровій трансформації підприємств.

Ключові слова: інтелектуальна система аналітики; машинне навчання; бізнес-аналітика; прогнозування; багатокритеріальна оптимізація.

1. ВСТУП

Постановка проблеми. Сучасні підприємства стикаються з необхідністю прийняття великої кількості управлінських рішень у короткі терміни, часто в умовах невизначеності та швидкої зміни ринкової ситуації. Традиційні методи аналізу даних уже не забезпечують достатнього рівня ефективності, а людський фактор створює ризики помилок у стратегічному плануванні.

У таких умовах актуальним стає впровадження інтелектуальних систем аналітики, здатних автоматично обробляти великі обсяги інформації, виявляти закономірності, прогнозувати тенденції та пропонувати оптимальні управлінські рішення.

Аналіз останніх досліджень і публікацій. Робота A Review of Data-Driven Decision-Making Methods for Industry 4.0 Maintenance Applications [1] узагальнює методи аналітики на основі даних у контексті Industry 4.0, підкреслюючи важливість машинного навчання для рішень у реальному часі. У дослідженні Forecasting Capacity of ARIMA Models [2] показано ефективність ARIMA для промислових прогнозів і перспективу комбінування з нейронними мережами. Публікація Multi-Objective Optimization in Business Analytics [3] обґрунтовує застосування еволюційних алгоритмів для балансування між прибутковістю, ризиком і сталим розвитком. У роботі Modern Business Intelligence: Big Data Analytics and Artificial Intelligence [4] наголошено на синергії BI, AI та Big Data у створенні аналітичної цінності підприємств.

Отже, актуальним завданням залишається створення комплексної архітектури, що об'єднує ці підходи в інтегровану аналітичну систему з високим рівнем безпеки та інтерпретованості результатів.

Мета публікації. Розробити концепцію інтелектуальної аналітичної системи, здатної забезпечити комплексну обробку даних і підтримку управлінських рішень на промисловому підприємстві.

2. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ.

Інтелектуальна система аналітики є складовою корпоративної інформаційної інфраструктури, побудованої за принципами модульності та інтегрованості.

Система включає декілька взаємопов'язаних рівнів:

Рівень даних — забезпечує збір, зберігання та попередню обробку інформації з внутрішніх систем (ERP, CRM, бухгалтерія, цехові датчики) і зовнішніх джерел (ринкові індекси, макроекономічні показники).

Аналітичний рівень — відповідає за моделювання процесів, виявлення закономірностей, побудову прогнозів і аналітичних звітів.

Рівень прийняття рішень — генерує рекомендації для менеджерів, підтримує сценарне моделювання й оптимізацію.

Презентаційний рівень — надає результати аналізу у вигляді інтерактивних панелей, звітів і графічних моделей.

Такий підхід дозволяє реалізувати наскрізну інтеграцію — від збору первинних даних до надання керівництву узагальнених висновків.

До функціональних можливостей системи входять:

1. ETL-процеси та підготовка даних

Система реалізує ETL-процеси, що дозволяють автоматично витягувати інформацію з різномірних джерел — таких як виробничі журнали, облікові системи ERP, CRM-платформи, таблиці Excel, сенсорні дані з обладнання та веб-інтерфейси зовнішніх постачальників. На етапі трансформації здійснюється узгодження форматів, усунення дублікатів, нормалізація структур даних і перевірка коректності значень (зокрема, цін, кількостей, дат і показників продуктивності). Очищені й структуровані дані завантажуються до централізованого сховища, що забезпечує цілісність, узгодженість і достовірність інформації — основу для подальшого аналітичного опрацювання.

2. Аналітична обробка даних

Аналітичний модуль виконує обчислення ключових показників ефективності, таких як рівень виконання плану, коефіцієнт використання ресурсів, маржинальність і динаміка продажів, у режимі реального часу. Для цього застосовуються алгоритми OLAP-аналізу для багатовимірного перегляду даних, методи кластеризації для виявлення прихованих закономірностей у поведінці клієнтів і виробничих процесах, а також лінійна регресія та методи машинного навчання для побудови прогнозних моделей.

Для прикладу, для аналізу залежності прибутку від витрат і обсягів виробництва використовується рівняння [5]:

$$P_t = a_0 + aD_t + \beta C_t + \gamma Q_t + \varepsilon, \quad (1)$$

де P_t — прибуток у момент часу t , D_t — попит, C_t — витрати, Q_t — якість продукції, a , β , γ — вагові коефіцієнти, a_0 — константа, а ε — похибка моделі.

Крім того, застосовується кореляційний аналіз, який визначає, які фактори найбільше впливають на кінцевий результат. Це дозволяє менеджерам фокусувати увагу на ключових показниках, що мають найвищу управлінську вагу, і своєчасно коригувати стратегію підприємства відповідно до аналітичних висновків.

3. Прогнозування та сценарне моделювання

Для прогнозування обсягів виробництва, попиту або витрат система може використовувати методи експоненційного згладжування або авторегресійні моделі (ARIMA) [6]:

$$F_{t+1} = aD_t + (1-a)F_t, \quad (2)$$

де F_{t+1} — прогноз попиту, D_t — фактичний попит, a — коефіцієнт згладжування.

Ці алгоритми дозволяють моделювати поведінку системи при зміні зовнішніх умов — наприклад, коливання цін на матеріали, зміни курсу валют або сезонні тенденції продажів.

Для складних сценаріїв може використовуватись моделювання на основі Монте-Карло, що дозволяє оцінити ймовірні наслідки різних стратегічних варіантів.

4. Оптимізаційний модуль

Система підтримує прийняття управлінських рішень шляхом розв'язання задач багатокритеріальної оптимізації, що формуються як [7]:

$$\max R = f(Q, C, T) \quad \text{за умов: } C \leq C_{\max}, T \leq T_{\max} \quad (3)$$

де R — інтегральний показник ефективності, Q — якість продукції, C — витрати, T — термін виробництва.

Результатом роботи алгоритму є набір рекомендованих управлінських дій, що дозволяють досягти найкращого співвідношення між прибутковістю, швидкістю та якістю.

5. Інтерфейс користувача

Важливим елементом системи є дашборд керівника — інтерактивна панель, яка відображає основні бізнес-показники, порівняння з планом, сигнали ризику та прогнозні тенденції.

Для відображення рівня виконання плану використовується показник:

$$IE = \frac{O_f}{O_p} \times 100\%, \quad (4)$$

де O_f — фактичний обсяг виробництва, O_p — плановий.

Інтерфейс також може містити елементи системи сповіщення, що автоматично повідомляють про критичні відхилення.

Система аналітики реалізується з використанням архітектури клієнт-сервер, що забезпечує масштабованість та модульність.

На серверній стороні використовується стек технологій:

Python / R — для побудови аналітичних моделей;

PostgreSQL / ClickHouse — для зберігання великих обсягів даних;

FastAPI або Flask — для створення REST API;

Power BI / Grafana / Superset — для візуалізації результатів.

На клієнтському рівні (веб або мобільний застосунок) передбачено зручний інтерфейс, який дозволяє керівникам переглядати аналітику з будь-якого пристрою в реальному часі.

Особливу увагу приділено безпеці даних — застосовуються механізми аутентифікації, шифрування каналів зв'язку та контроль доступу на основі ролей [8].

Використання інтелектуальної аналітичної системи дозволяє [9]:

— зменшити час на підготовку звітів на 70–80%;

— скоротити управлінські помилки завдяки об'єктивному аналізу даних;

— підвищити точність прогнозування попиту до 95%;

— автоматизувати рутинні операції збору та обробки інформації;

— створити єдине аналітичне середовище для всіх рівнів управління.

Крім того, система формує цифровий профіль підприємства, що дозволяє відстежувати його динаміку розвитку, оцінювати ефективність рішень і розробляти стратегії довгострокового зростання.

Висновки. Запропонована інтелектуальна система є основою переходу підприємства до управління на основі даних. Комбінація аналітичних і оптимізаційних алгоритмів забезпечує глибокий аналіз, прогнозування й підтримку рішень у реальному часі.

Практичне впровадження довело ефективність — підвищення точності, швидкості аналітики та зниження ризиків.

Подальші дослідження спрямовано на інтеграцію моделей глибинного навчання, IoT, блокчейну та Explainable AI для підвищення автономності, безпеки й прозорості аналітичних процесів.

Розвиток таких систем відкриває перспективи формування адаптивних цифрових екосистем, здатних самостійно аналізувати зміни ринку та динамічно коригувати бізнес-стратегії.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Bousdekis A., Lepenioti K., Apostolou D., Mentzas G. A., Review of Data-Driven Decision-Making Methods for Industry 4.0 Maintenance Applications. *Electronics*. 2021. Vol. 10, No. 1, P. 1–20.
2. Tomić D., Stjepanović S., Forecasting Capacity of ARIMA Models; A Study on Croatian Industrial Production and its Sub-sectors. *Zagreb International Review of Economics and Business*. 2017. Vol. 20, No. 1, pp. 81-99. DOI:10.1515/zireb-2017-0009.
3. Ridwan B. I., Addo S., Multi-Objective Optimization in Business Analytics: Balancing Profitability, Risk Exposure, and Sustainability in Strategic Decision-Making. *International Journal of Research Publication and Reviews*. 2025. Vol. 2, No. 5, P.89–111. DOI: 10.55248/gengpi.6.0525.1718.
4. Ahmed A.A., Gad-Elrab, Modern Business Intelligence: Big Data Analytics and Artificial Intelligence for Creating the Data-Driven Value. 2021. Vol. 1, No. 1. P. 1–18
5. Sarmiento R. M., Costa V., Introduction to Linear Regression. *Comparative Approaches to Using R and Python for Statistical Data Analysis*. 2017. Vol.1, No. 1, P. 111–115. DOI: 10.4018/978-1-68318-016-6.ch006.
6. Ostertagová E., Ostertag O., Forecasting Using Simple Exponential Smoothing Method. *Acta Electrotechnica et Informatica*. 2012. Vol. 12, No. 3, P. 62–66. DOI:10.2478/v10198-012-0034-2.
7. Kalyan D., Multiobjective Optimization Using Evolutionary Algorithms. *Jhon Wiley and Sons Ltd*. 2001. Vol. 1, No. 1, P. 1–24.
8. Abduhari E. S., Shaik T. C., Adidul A. B., Ladja J. H., Access Control Mechanisms and Their Role in Preventing Unauthorized Data Access: A Comparative Analysis of RBAC, MFA, and Strong Passwords. *Natural Sciences Engineering and Technology Journal (NASET Journal)*. 2024. Vol. 5, No. 1, P. 418-430. DOI:10.37275/nasetjournal.v5i1.62.
9. Hočevár B., Jaklič J., Assessing Benefits of Business Intelligence Systems – A Case Study. *Management*. 2010. Vol. 15, No. 1, P. 87–119.

Віталій Лисенко

Доктор технічних наук, професор кафедри автоматики та робототехнічних систем
ім. І.І. Мартиненка

Місце роботи: Національний університет біоресурсів і природокористування України, факультет
інформаційних технологій, м. Київ, Україна

<https://orcid.org/0000-0002-5659-6806>

lysenko@nubip.edu.ua

Микола Доля

Доктор сільськогосподарських наук, професор кафедри ентомології, інтегрованого захисту та карантину
рослин

Місце роботи: Національний університет біоресурсів і природокористування України, Факультет захисту
рослин, біотехнологій та екології, м. Київ, Україна

<https://orcid.org/0000-0003-0458-9695>

mykola.dolia@gmail.com

Тарас Лендєл

Кандидат технічних наук, доцент, доцент кафедри автоматики та робототехнічних систем
ім. І.І. Мартиненка

Місце роботи: Національний університет біоресурсів і природокористування України, ННІ енергетики,
автоматики і енергозбереження, м. Київ, Україна

ORCID ID <https://orcid.org/0000-0002-6356-1230>

taraslendel@gmail.com

ПРОГНОЗУВАННЯ ВРОЖАЙНОСТІ ТА ЯКОСТІ ПОСІВІВ КУКУРУДЗИ ДЛЯ СТЕПОВОЇ ЗОНИ УКРАЇНИ

Анотація. У дослідженні проводиться системний аналіз впливу на урожайність зернової кукурудзи таких факторів як параметри навколишнього середовища (тривалість сонячного сйва, середньорічна температура, сума опадів, середньорічна вологість повітря) та чисельність основних шкідників, котрі не лише фізично зменшують урожайність, а й сприяють розповсюдженню різноманітних захворювань (матеріали аналізувались для Степової Зони України). До таких шкідників, на наш погляд, слід віднести совку озиму і стеблового кукурудзяного метелика. Їх сукупний вплив на врожайність кукурудзи, за твердженнями фахівців із захисту рослин, є основним і діє на різних періодах вегетації рослини. Кореляційний аналіз зв'язків урожайності кукурудзи і факторів впливу показав наявність невизначеності, що дозволило рекомендувати для прогнозування штучні нейронні мережі зі структурою «радіально-базисна функція» і «багатошаровий перцептрон». Вибірki, отримані за роками спостереження для Степу України, були використані для машинного навчання та для контролю похибки. Отримані результати дозволяють стверджувати про можливість практичного їх використання для прогнозування урожайності кукурудзи на зерно.

Ключові слова: урожайність кукурудзи, системний аналіз, штучні нейронні мережі, вибірка, машинне навчання.

1. ВСТУП

На врожайність та якість урожаю суттєво впливають природні фактори, кількість шкідників та їх шкідливість, а також технології вирощування. У процесі вирощування зерна рослини використовують різноманітні хімічні речовини – інсектициди – для мінімізації впливу шкідників. Ці послуги є дорогими та часто негативно впливають на екологічний стан навколишнього середовища.

Відомо, що зерно кукурудзи як в Україні, так і в зарубіжних країнах отримало широке застосування як в аграрному виробництві, так і в харчовій переробній промисловості. В останні десятиліття силос із кукурудзи стає також популярним у використанні для отримання біогазу. Зважаючи на зазначене, урожайність кукурудзи стає важливим економічним фактором широкого розповсюдження цієї рослинної

культури. У той же час на врожайність кукурудзи суттєво впливають зони вирощування, що характеризуються різноманітністю значень природних факторів таких, наприклад, як температура, вологість, тривалість і інтенсивність сонячної радіації, наявність шкідників.

Постановка проблеми. Пропонується новий і системний, з точки зору авторів цієї статті, метод прогнозування врожайності кукурудзи на основі використання штучного інтелекту. Ефективність його застосування вже була апробована в публікаціях [1], [2]. Окрім того, машинне прогнозування стає важливим етапом автоматизації виробництва зерна. Звертаємо увагу, що аналіз матеріалів проводився для Степу України.

Аналіз останніх досліджень і публікацій. У своїх дослідженнях ми скористались інформацією за роками, поданою в [1], [2], щодо урожайності кукурудзи і основними факторами впливу (параметри навколишнього середовища, чисельність шкідників).

Мета публікації. створити і протестувати математичну модель прогнозування врожайності та якості урожаю зернових культур в Україні.

2. ТЕОРЕТИЧНІ ОСНОВИ

У попередніх роботах нами були використані штучні нейронні мережі для прогнозування чисельності совки озимої (Dolia, M. et al., 2024). Цей шкідник відноситься до небезпечних практично для всіх зернових рослин. Набутий при цьому досвід дозволяє рекомендувати до використання структуру штучної нейронної мережі – багатошаровий перцептрон для прогнозування чисельності такого небезпечного для кукурудзи шкідника як стебловий кукурудзяний метелик. При цьому факторами впливу на його чисельність будуть (уведемо позначення): T - тривалість сонячного саява, t - середньорічна температура, h - сума опадів, ϕ – середньорічна відносна вологість повітря, Sm – чисельність стеблового кукурудзяного метелика, So - чисельність совки озимої.

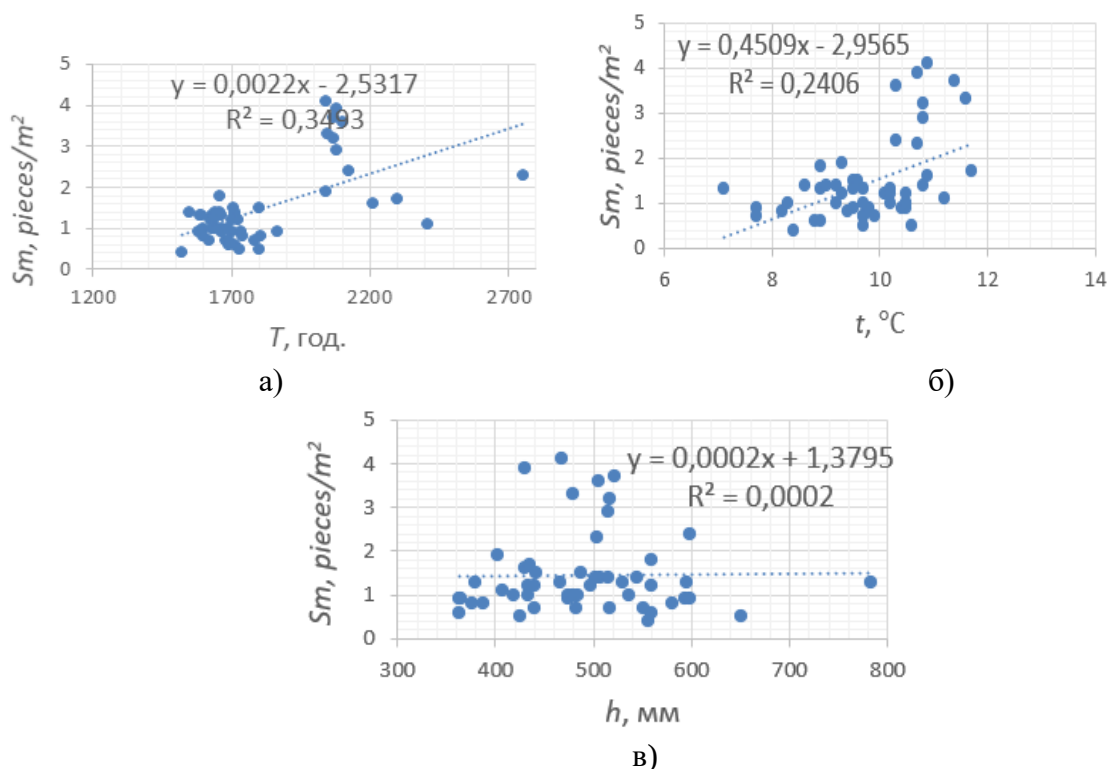


Рис. 1. Залежності чисельності стеблового метелика кукурудзяного від: а) тривалості сонячного саява; б) середньорічної температури; в) суми опадів

Ознакою для вибору виду моделі прогнозування є результати кореляційного аналізу, що подаються на рис.1. При цьому оцінюють коефіцієнт кореляції лінійної моделі розподілу відповідних спостережень й коефіцієнт детермінації, що дає можливість оцінити адекватність такої -моделі.

3. МЕТОДИ ДОСЛІДЖЕННЯ

Наявність невизначеності між урожайністю кукурудзи і факторами впливу на цей показник (означені вище) дозволяє зробити висновок про доцільність використання штучних нейронних мереж для прогнозування за роками урожайності кукурудзи на зерно із глибиною прогнозу в 1 рік.

4. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Реалізовано структуру штучної нейронної мережі для прогнозування урожайності кукурудзи та отримано результати прогнозу, що наведено на рис. 2.

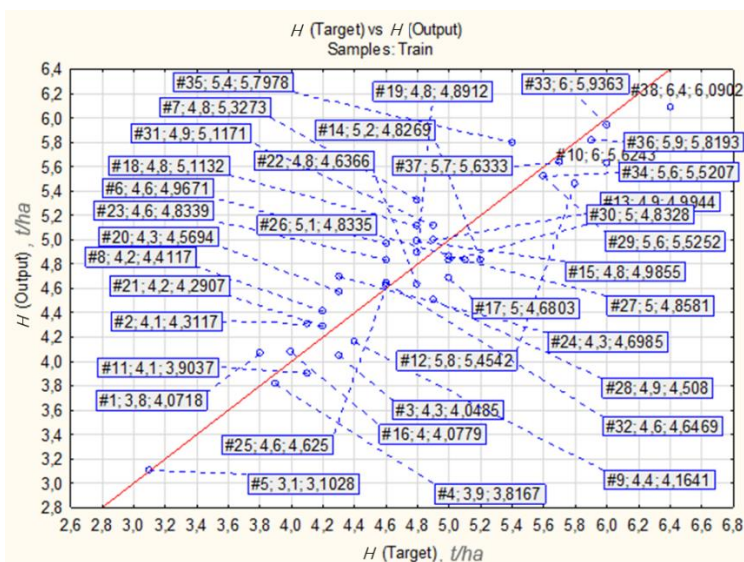


Рис. 2. Порівняння результатів прогнозу урожайності кукурудзи за моделлю $H(\text{Output})$ та реальними спостереженнями $H(\text{Target})$ за роками

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Chen, Y., Chen, M., Guo, M., Wang, F., & Wang, J. (2023, August). A Multi-stage Prediction Framework for Pest Identification. In 2023 IEEE 9th International Conference on Cloud Computing and Intelligent Systems (CCIS) (pp. 583-588). IEEE. doi: 10.1109/CCIS59572.2023.10262999.
2. Dolia, M., Lysenko, V., Lendiel, T., Nakonechna, K., & Vorokh, V. (2024). Artificial neural networks for predicting the number of field crop pests. Scientific Reports of the National University of Life and Environmental Sciences of Ukraine, 3(20).
3. Azrour, M., Mabrouki, J., Fattah, G., Guezzaz, A., & Aziz, F. (2022). Machine learning algorithms for efficient water quality prediction. Modeling Earth Systems and Environment, 8(2), 2793-2801.

Gao Hui

Master's degree, lecturer

Work place: BengBu University, Faculty of Mathematics and Physics, BengBu, China

gaohui579@163.com

Zhou Meng zhen

Bachelor's degree, student

Work place: BengBu University, Faculty of Mathematics and Physics, BengBu, China

2011557950@qq.com

RESEARCH ON DEHAZING OF SURVEILLANCE IMAGES BASED ON DEEP LEARNING

Abstract. With the widespread application of surveillance systems in security, transportation, and other fields, image dehazing has become a key technology to improve the quality of surveillance images. This paper conducts research on surveillance image dehazing using the All-in-One Dehazing Network (AOD-Net) technology. The dataset consists of nearly 4,000 images, including the OTS subset from the RESIDE dataset and foggy images captured in real campus scenarios, which are divided into a training set and a test set at a ratio of 9:1. In terms of qualitative evaluation, the dehazed images conform to human visual perception, with details and colors close to the original images. For quantitative assessment, the mean Peak Signal-to-Noise Ratio (PSNR) reaches 15.25 dB, and the mean Structural Similarity Index (SSIM) is 0.839. Experimental results demonstrate that the AOD-Net technology achieves excellent dehazing performance for surveillance images.

Keywords: Deep Learning; Dehazing Model; Surveillance Images.

1. INTRODUCTION

Under natural weather conditions such as rain, fog, and haze, images are prone to degradation. Compared with visual perception in sunny environments, degraded images usually exhibit reduced contrast and blurred target details. These issues significantly affect the accuracy of algorithms like face recognition and license plate recognition, and may even lead to misjudgments by surveillance systems. Image dehazing technology can effectively enhance image details and improve the precision of target recognition, thereby addressing the aforementioned problems. With the continuous development of deep learning technology, image dehazing technology has also undergone constant innovation. Deep learning-based image dehazing has become a focus of research in the field of computer vision [1-4]. However, practical application scenarios of image dehazing technology are complex and diverse, and existing algorithms have limitations to varying degrees. This paper studies deep learning-based image dehazing algorithms, aiming to provide theoretical support and practical basis for algorithm performance optimization and practical problem-solving.

Cai et al. [5] pioneered the introduction of deep learning algorithms into the field of image dehazing and proposed the DehazeNet algorithm. This algorithm integrates convolutional neural networks (CNNs) with prior knowledge, trains the model using massive datasets, automatically calculates image transmittance through learning, and then restores clear images based on the atmospheric scattering model. Ren et al. [6] constructed a multi-scale convolution theoretical model based on a deep learning framework to solve the transmittance estimation problem and designed an algorithm architecture consisting of two processing stages. In the first stage, the research team used a multi-scale CNN to achieve initial transmittance estimation; in the second stage, a generative adversarial network (GAN) mechanism was introduced to refine the preliminary results. However, this algorithm has insufficient adaptability in practical applications. When processing images with different fog concentration characteristics, manual parameter adjustment is often required to obtain ideal

dehazing effects. Li et al. [7] proposed the All-in-One Dehazing Network (AOD-Net) model. To address the problem that separately solving two unknown variables in the atmospheric scattering model tends to amplify errors, this model first transforms the atmospheric scattering model, integrates the two unknown parameters into a single joint parameter, estimates the joint parameter using a CNN, and then substitutes the estimated result into the transformed model to obtain clear dehazed images. The Multi-Scale Boosted Dehazing Network (MSBDN) proposed by Dong et al. [8] innovatively constructs a dense feature fusion mechanism, which effectively alleviates the spatial information loss commonly found in traditional dehazing networks. This algorithm optimizes the decoder structure and adopts a progressive processing strategy to achieve high-quality haze-free image reconstruction.

Hou Ming et al. [9] improved AOD-Net in three aspects—adding an attention mechanism, adjusting the network structure, and modifying the loss function—to address issues related to dehazing algorithm efficiency and detail restoration, thereby enhancing its feature extraction and restoration capabilities for clear dehazed images. Dong Tian et al. [10] improved and optimized the AOD-Net algorithm based on the Zynq platform, designing and implementing an image dehazing system. While improving dehazing effects, it effectively reduces algorithm computational complexity and resource consumption, enhances system real-time performance and practicality, and is applicable to various image dehazing scenarios. Xu Yue et al. [11] addressed the problems existing in traditional image dehazing algorithms and AOD-Net by adding a spatial pyramid attention mechanism, adjusting the network structure to a Laplacian pyramid type, and modifying the loss function to restore clear images.

Image dehazing technology can be divided into three categories [12-14]: (1) Enhancing visual features such as contrast and brightness of foggy images to achieve dehazing. However, due to the lack of in-depth analysis of image degradation mechanisms, detail information is often lost during enhancement, affecting the final dehazing quality. (2) Estimating atmospheric scattering model parameters using prior knowledge to reconstruct clear images through the model. Nevertheless, these prior knowledge may fail in complex scenarios, leading to parameter estimation deviations and limiting the algorithm's scope of application. (3) Two deep learning-based image dehazing methods: one predicts the transmittance map using CNNs and then reconstructs images combined with the atmospheric scattering model, but the process is overly cumbersome; the other adopts an end-to-end network architecture to directly restore haze-free images, avoiding the parameter estimation process and thus achieving better dehazing results. Therefore, this paper adopts the end-to-end network architecture to restore haze-free images.

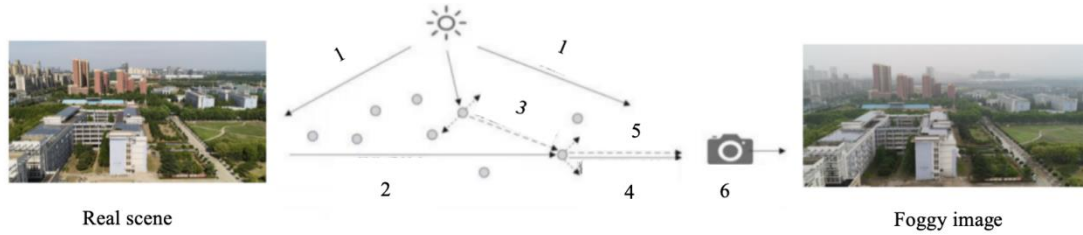
2. Theoretical Basis

This chapter first introduces the characteristics and formation principles of foggy images, then outlines theories related to image dehazing, focusing on the basic structure and principles of neural networks.

2.1 Foggy Images

Foggy images typically exhibit reduced contrast, decreased color saturation, and blurred details. Their generation is mainly affected by two factors [15], as shown in Figure 1. (Atmospheric Light Imaging Model). Firstly, scene reflected light interacts with suspended particles in the air during transmission. According to the Mie scattering theory, the scattering intensity difference between different wavelength lights is weak, and particles absorb incident light, resulting in a significant attenuation of the energy of scene reflected light reaching the optical imaging equipment. Secondly, background light in the atmospheric environment (including atmospheric light and ambient light generated by scattering) is affected by secondary scattering of suspended particles, and part of the stray light enters the imaging system to form noise. These two effects together constitute the mechanism by which

suspended particles affect optical imaging quality, hindering light propagation, causing problems such as reduced intensity and color homogenization, and ultimately leading to reduced contrast, decreased color saturation, and blurred details in captured images.



1:Lighting; 2:Scene reflected light; 3:Haze particle reflection; 4:Attenuated reflected light; 5:Atmospheric light; 6:Optical imaging equipment

Fig. 1. Atmospheric Light Imaging Model

2.2 All-in-One Dehazing Network.

Proposed at the 2017 IEEE International Conference on Computer Vision (ICCV), AOD-Net is a lightweight deep learning dehazing model designed based on classic convolutional neural networks. It adopts an end-to-end neural network architecture, enabling image dehazing in a single forward pass and simplifying traditional multi-step dehazing algorithms. The model reduces computational complexity through a lightweight network structure and incorporates Position-wise Normalization (PONO) to enhance dehazing performance.

(1) AOD-Net Algorithm

The atmospheric scattering model is a classic description of foggy image generation:

$$I(x) = J(x)t(x) + A(1-t(x)) \quad (1)$$

Where $J(x)$ represents the clean image to be restored, $I(x)$ is the foggy image captured by the camera, $t(x)$ is the transmittance, A is the atmospheric light value, and x denotes the pixel position in the image.

The core idea of AOD-Net is to unify the transmittance $t(x)$ and atmospheric light value A in the atmospheric scattering model into a single expression represented by $K(x)$. It extracts and integrates image features using a multi-scale feature fusion network to obtain the required $K(x)$ parameters [11], and then substitutes the estimated results into the transformed model to achieve haze-free image restoration.

The expression for the image to be restored is as follows:

$$J(x) = K(x)I(x) - K(x) + b \quad (2)$$

Where b is a constant bias, defaulting to 1. The unknown value obtained by combining transmittance $t(x)$ and atmospheric light value A is expressed as:

$$K(x) = [(I(x) - A) / t(x) + (A - b)] / (I(x) - 1) \quad (3)$$

(2) AOD-Net Network Structure

The AOD-Net network structure consists of two parts: a K -value module and an image generation module [16], as shown in Figure 2. The K -value module extracts features from images through convolution operations (Conv), with each convolution followed by a ReLU activation function to achieve non-linear mapping. The feature fusion stage adopts the Concat method to enrich the extracted image feature information by concatenating features from different levels. After Conv5, the module obtains the $K(x)$ parameter. The image generation module inputs the obtained K -value into the expression of the image to be restored to generate the final dehazed image.

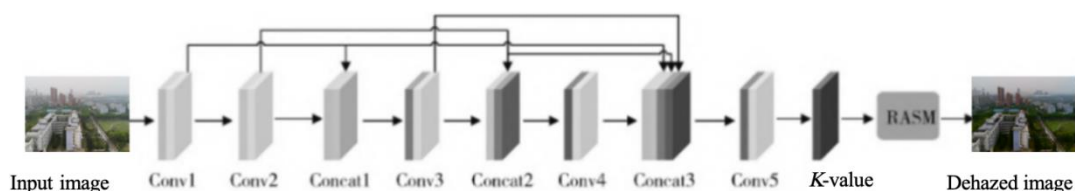


Fig. 2. AOD-Net Network Structure Diagram

With its end-to-end design, concise structure, and strong generalization ability, AOD-Net has demonstrated excellent dehazing performance in various scenarios. It is also easy to embed into other deep models for joint optimization, improving the efficiency and accuracy of image dehazing tasks.

3.Results and Analysis

3.1 Experimental Preparation

(1) Experimental Environment

The hardware configuration of the experimental computer includes an AMD Ryzen 7 5800H processor, 16G memory, and the Windows 10 operating system. The programming language version is Python 3.8, the algorithm implementation platform is PyCharm, and the deep learning framework used is PyTorch 1.8.2+cu111.

(2) Algorithm Parameter Configuration

This experiment uses the AOD-Net model for image dehazing training. During training, the batch size is set to 8, and the batch size for validation is also 8 to ensure consistency in sample processing during training and validation. To improve data loading efficiency, the number of worker processes (num_workers) is set to 2, enabling two sub-processes for data preprocessing. The total number of training epochs (num_epochs) is 10, suitable for initial training of small and medium-sized datasets. To prevent model overfitting, a weight decay mechanism is introduced with a parameter regularization coefficient (weight_decay) of 0.0001. Additionally, to avoid gradient explosion during training, a gradient clipping (grad_clip_norm) mechanism is used with a threshold of 0.1. These parameter configurations aim to enhance training stability and model generalization ability.

(3) Dataset Selection

The dataset selected in this paper consists of two groups: the first group is the OTS subset from the RESIDE dataset; the second group is composed of foggy images captured in real campus scenarios. The dataset includes various scenes under real environments to enhance the model's adaptability to complex environments.

The RESIDE [17] dataset was created by research teams from Microsoft Research and the University of Illinois at Urbana-Champaign in 2017 and first proposed in 2018. It is divided into five subsets, including the Indoor Training Set (ITS) and Outdoor Training Set (OTS). The OTS subset contains 2,061 real outdoor images from real-time weather in Beijing, which were synthesized into over 70,000 foggy images under different parameters. This paper selects more than 2,000 pairs of different images from the OST subset of the RESIDE dataset and divides them into training and test sets at a ratio of 9:1.

The dataset composed of foggy images from real campus scenarios includes images of Nanshan of the university captured from the west view of the 7th floor of the Science and Technology Building of Bengbu University, as well as images of teaching buildings, playgrounds, dormitories, and canteens taken from fixed perspectives under foggy weather, totaling more than 2,000 pairs of different images, which are also divided into training and test sets at a ratio of 9:1. Some images are shown in Figure 3 (one haze-free image corresponds to one foggy image; only part of them is displayed here).

(4) Dataset Processing

During dataset processing, original images are first paired with their corresponding images, and the image sizes are unified to ensure consistency. Next, the pixel values of the images are normalized to the range of 0 to 1 to facilitate subsequent numerical calculations. The processed image data is converted into a standard tensor format, and the dimension order is adjusted to meet the model's input requirements. The entire process maintains the corresponding relationship between images, ensuring that each pair of images can participate in calculations together to form standard input-output samples. This processing method improves data consistency and usability, laying a solid foundation for subsequent training and testing.

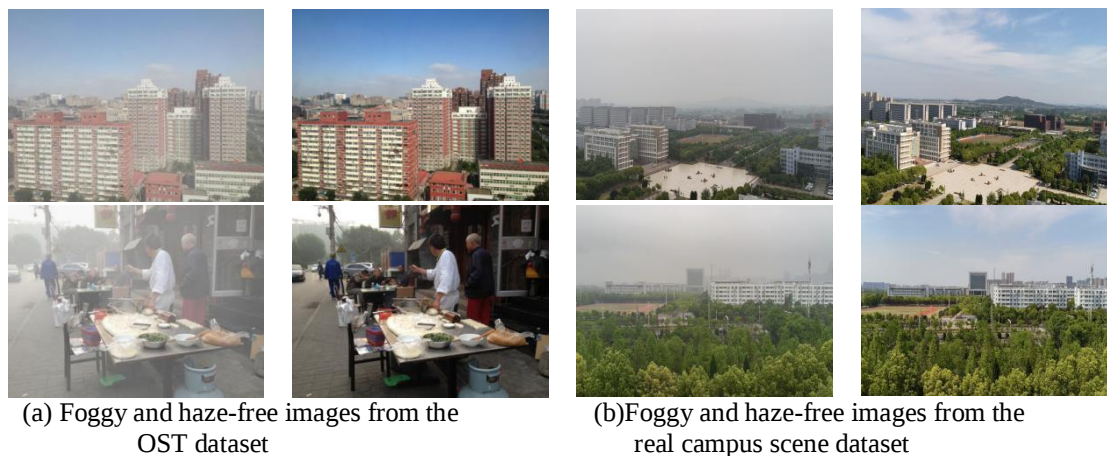


Fig. 3. Display of Partial Dataset

3.2 Experimental Results and Analysis

(1) Experimental Results

Figures 4 and 5 show the dehazing effects of partial scenes from the OST dataset and the real campus scene dataset in the test set, respectively. It can be observed that:

1. The restored images have good visual consistency, closely matching the perceptual effect of images under real clear scenes.
2. During the dehazing process of the network, image detail information is well preserved, resulting in dehazed images with rich details.
3. The colors of the dehazed images are close to those of real scenes.

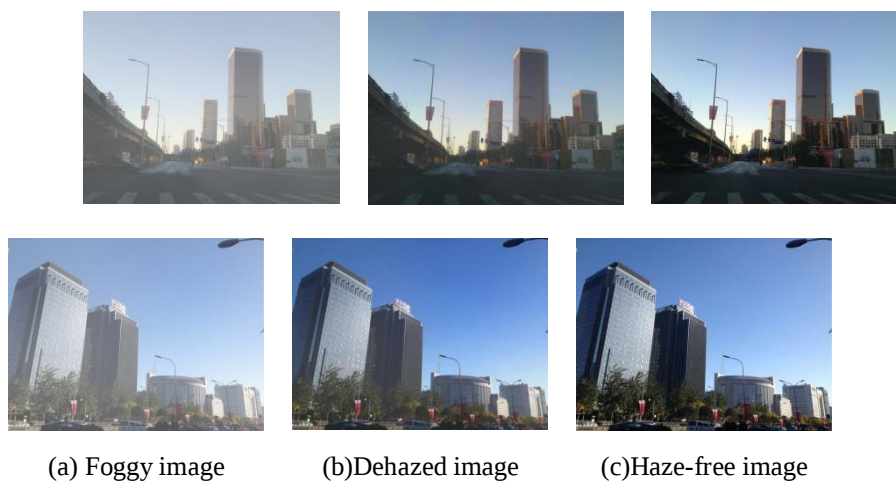


Fig. 4. Dehazing Effects of the OST Dataset

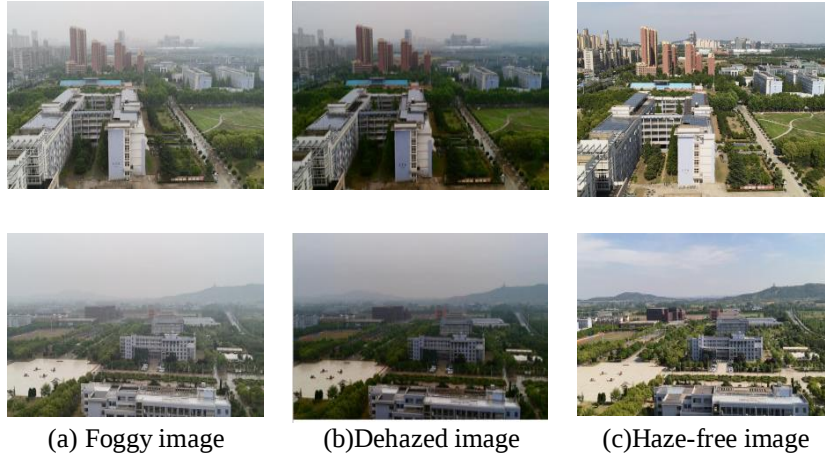


Fig. 5. Dehazing Effects of the Real Campus Scene Dataset

Comparing the restoration effects of the two datasets, it is found that the restoration effect of the OST dataset is slightly better than that of the real campus foggy dataset. The dark areas in the restored images of the real campus scene dataset are darker, with less detailed restoration and incomplete fog removal. The OST subset mainly consists of synthetic images with uniformly distributed fog and moderate fog concentration; most haze-free images have brightness in a "daily normal exposure" state. The fog in the real campus scene dataset may be morning fog or polluted haze with higher local density; the haze-free images are captured under "strong sunny lighting" conditions, featuring high brightness and strong contrast. During training, "synthetic fog + non-uniform natural fog" was used, and the AOD-Net model training was overly conservative, leading to incomplete dehazing. Meanwhile, as a shallow convolutional network, the AOD-Net model processes the entire image evenly without regional perception and distance recognition, causing the trained model to misclassify dark foreground areas as "light fog areas" with underestimated transmittance, resulting in over-dehazing and reduced brightness.

(2) Analysis of Experimental Results

To quantitatively reflect the image dehazing effect of the deep neural network, this paper uses two evaluation methods: Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index (SSIM).

1. Mean Squared Error (MSE): A commonly used loss function to measure the difference between predicted values and true values. The expression is as follows:

$$MSE = \frac{\sum_{i=1}^M \sum_{j=1}^N [x(i, j) - y(i, j)]^2}{N_x \times N_y} \quad (6)$$

Where $x(i, j)$ and $y(i, j)$ represent the pixel values of the original image x and the processed image y at the corresponding positions, respectively, and $N_x \times N_y$ is the size of the image.

2. Peak Signal-to-Noise Ratio (PSNR): Detects image distortion pixel by pixel. A higher evaluation value indicates better image restoration effect, making it the most widely used evaluation standard in signal processing. Its unit is decibels (dB), and the expression is as follows:

$$PSNR = 20 \cdot \log_{10} \left(\frac{255}{\sqrt{MSE}} \right) \quad (7)$$

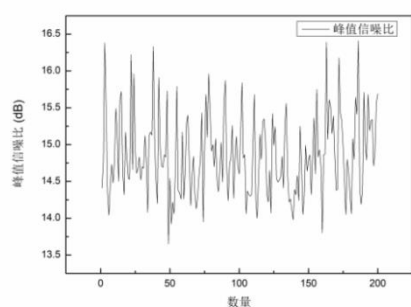
3. Structural Similarity Index (SSIM): An image quality evaluation index closer to human subjective perception. Since pixels in an image form a specific structure with

surrounding pixels, SSIM compares the similarity between two images based on changes in image information structure. The expression is as follows:

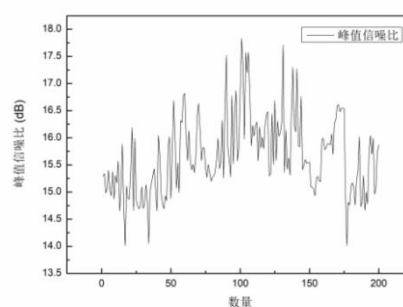
$$SSIM(x, y) = \frac{(2\mu_x\mu_y + c_1)(2\sigma_{xy} + c_2)}{(\mu_x^2 + \mu_y^2 + c_1)(\sigma_x^2 + \sigma_y^2 + c_2)} \quad (8)$$

Where c_1 and c_2 are constants, x and y are the original image and the processed image, respectively, μ is the mean value of the image, σ^2 is the variance of the image, and σ_{xy} is the covariance between the two images.

The distributions of PSNR and SSIM test results for the OTS dataset and the real campus scene dataset are shown in Figures 6 and 7. The mean values of PSNR and SSIM for the two groups of test results are listed in Table 1.

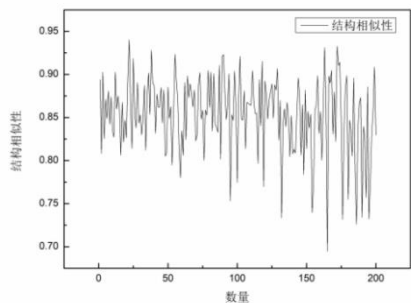


(a)PSNR distribution of the OTS dataset

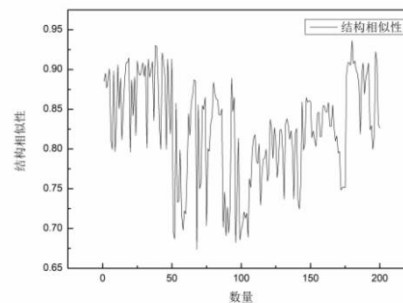


(b)PSNR distribution of the real campus scene dataset

Fig. 6. PSNR Test Results



(a)SSIM distribution of the OTS dataset



(b)SSIM distribution of the real campus scene dataset

Fig. 7. SSIM Test Results

Table 1

Mean Values of PSNR and SSIM		
	PSNR	SSIM
Mean value of the OTS dataset	14.84dB	0.853
Mean value of the real campus scene dataset	15.665dB	0.825
Overall mean value	15.25dB	0.839

Quantitative evaluation results show that the PSNR and SSIM of the OTS dataset are relatively concentrated, with mean values of 14.84 dB and 0.853, respectively. The PSNR and SSIM of the real campus scene dataset fluctuate more significantly compared to the OTS dataset, with mean values of 15.665 dB and 0.825, respectively. The overall mean PSNR of

the two groups is 15.25 dB, and the overall mean SSIM is 0.839. The OTS dataset consists of artificially synthesized foggy images, while the real campus scene dataset is captured in real environments with greater randomness. Therefore, the PSNR and SSIM of the OTS dataset are more concentrated, with higher SSIM values and more stable structure restoration effects. The possible reason for the higher PSNR in the real campus scene dataset than in the OTS dataset is that the dehazed images of the real scene are generally darker with poor visual effects, but the errors can be evenly distributed across the entire image, resulting in a not necessarily large MSE but possibly a more stable one. In contrast, the OST dataset consists of synthetic images, and slight inadequate restoration can lead to large local errors, increasing the MSE and thus reducing the PSNR. Both subjective and objective evaluations indicate that the model achieves good dehazing effects. However, there are still issues such as incomplete dehazing and insufficient detail restoration. In the future, the training dataset size can be increased and the network structure can be optimized to achieve better restoration effects.

CONCLUSIONS AND PROSPECTS FOR FURTHER RESEARCH

This paper constructs a dataset of nearly 4,000 images, including the OTS subset from the RESIDE dataset and real campus foggy images, split into training and test sets at a 9:1 ratio. Foggy images, restored images, and evaluation metrics are intuitively presented via a visualization interface. Qualitatively, the model performs well for both synthetic (OTS) and real campus foggy images—with slightly better color restoration and detail clarity for the former. Quantitatively, the OTS dataset shows concentrated PSNR (14.84 dB) and SSIM (0.853), while the real campus dataset has more fluctuating values (15.665 dB PSNR, 0.825 SSIM). The OTS dataset's synthetic nature ensures stable structural restoration, while real-scene dehazed images are darker but with evenly distributed errors (stable MSE). In contrast, synthetic images' local errors inflate MSE and lower PSNR, leading to higher PSNR in the real campus dataset. Overall, the model's restored images align with human visual perception (details and colors close to original), achieving a mean PSNR of 15.25 dB and SSIM of 0.839, demonstrating AOD-Net's excellent dehazing performance for surveillance images.

The model proposed in this paper performs well in eliminating fog from input foggy images. However, there are phenomena such as significant brightness differences between foggy and haze-free images in the real dataset, generally dark restored dehazed images, and incomplete dehazing. In the future, position attention or multi-scale mechanisms can be introduced to enhance the model's ability to distinguish foreground and background regions. Brightness normalization can be applied to images to avoid the model mistakenly learning the relationship between brightness and fog. Additionally, generative adversarial networks can be used to improve the model's generation ability, and the dataset including different lighting conditions and scenes can be expanded to enhance the model's generalization ability and improve the overall restoration effect of AOD-Net.

REFERENCES (TRANSLATED AND TRANSLITERATED)

- [1] W. Zhang, "Research on haze image dehazing technology and application," [Ph.D. dissertation], Nanjing Univ. Posts Telecommun., Nanjing, China, 2023.
- [2] Z. Liu, Y. Liang, and J. Li, "Adaptive image dehazing algorithm based on dynamic convolution kernel," *Comput. Sci.*, vol. 50, no. 6, pp. 200–208, 2023.
- [3] Y. Wei and H. Xu, "A survey of deep learning-based image dehazing methods," *Inf. Technol. Informatization*, no. 4, pp. 214–216, 220, 2023.
- [4] T. Jia, L. Zhuo, J. Li, et al., "Research progress on single image dehazing based on deep learning," *Acta Electron. Sinica*, vol. 51, no. 1, pp. 231–245, 2023.
- [5] B. Cai, X. Xu, K. Jia, et al., "DehazeNet: An end-to-end system for single image haze removal," *IEEE Trans. Image Process.*, vol. 25, no. 11, pp. 5187–5198, Nov. 2016.

- [6] W. Ren, S. Liu, H. Zhang, et al., "Single image dehazing via multi-scale convolutional neural networks," in Proc. Eur. Conf. Comput. Vis., Cham, Switzerland: Springer, 2016, pp. 154–169.
- [7] B. Li, X. Peng, Z. Wang, et al., "AOD-Net: All-in-one dehazing network," in Proc. IEEE Int. Conf. Comput. Vis., Venice, Italy: IEEE, 2017, pp. 4770–4778.
- [8] H. Dong, J. Pan, L. Xiang, et al., "Multi-scale boosted dehazing network with dense feature fusion," in Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit., Seattle, WA, USA: IEEE, 2020, pp. 2157–2167.
- [9] M. Hou and W. Liang, "Image dehazing algorithm based on improved AOD-Net," Appl. Electron. Technol., vol. 50, no. 4, pp. 60–66, 2024.
- [10] T. Dong, D. Jiang, and H. Zhang, "Design of image dehazing system based on novel AOD-Net algorithm on Zynq platform," Ship Electron. Eng., vol. 44, no. 10, pp. 26–30, 2024.
- [11] Y. Xu, Z. Huang, and H. Wang, et al., "Image dehazing algorithm based on improved multi-scale AOD-Net," J. Chongqing Univ., vol. 48, no. 2, pp. 50–61, 2025.
- [12] B. Li and J. Liu, "A survey of image dehazing technology," Mod. Comput., no. 13, pp. 57–61, 2022.
- [13] F. Zheng, X. Wang, and D. He, et al., "A survey of single image dehazing algorithms," Comput. Eng. Appl., vol. 58, no. 3, pp. 1–14, 2022.
- [14] N. Cui, P. Tang, and X. Zheng, et al., "Research progress of image dehazing technology," Mech. Electr. Eng. Technol., vol. 54, no. 3, pp. 7–13, 57, 2025.
- [15] H. Wei, J. Tian, and Z. Xiao, "A survey of image dehazing algorithms," Softw. Guide, vol. 20, no. 8, pp. 231–235, 2021.
- [16] Y. Li, H. Cui, and H. Zhu, et al., "An aerial image dehazing algorithm based on improved AOD-Net," Acta Autom. Sinica, vol. 48, no. 6, pp. 1543–1559, Jun. 2022.
- [17] B. Li, W. Ren, D. Fu, et al., "Benchmarking single-image dehazing and beyond," IEEE Trans. Image Process., vol. 28, no. 1, pp. 492–505, Jan. 2018.