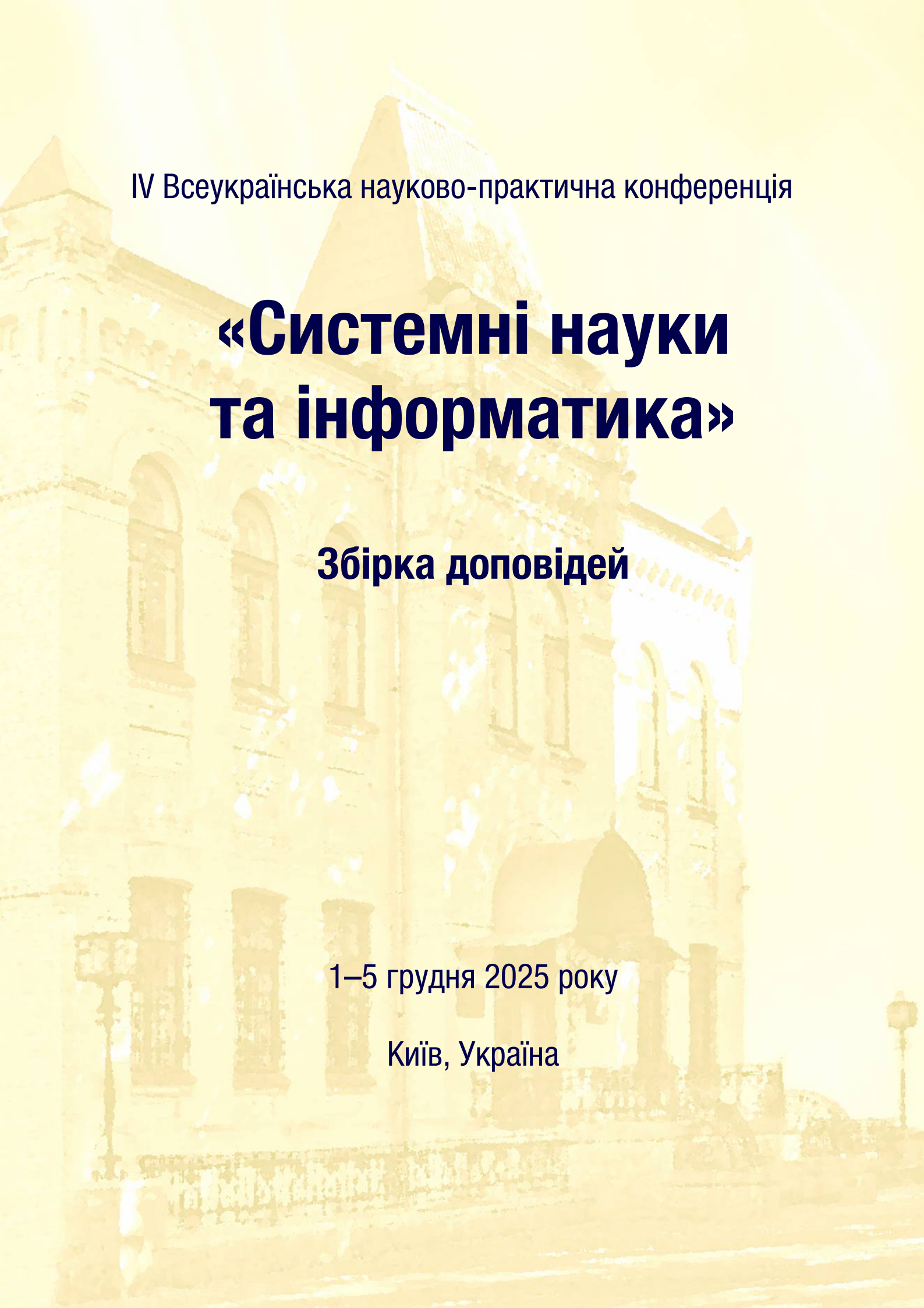




IV Всеукраїнська науково-практична конференція

«Системні науки та інформатика»

Збірка доповідей

1–5 грудня 2025 року

Київ, Україна

- С40 Системні науки та інформатика:** збірка доповідей IV науково-практичної конференції «Системні науки та інформатика», 1–5 грудня 2025 року, Київ [Електронне видання]. – К., НН ІПСА КПІ ім. Ігоря Сікорського, 2025. – 347 с.

До збірки увійшли матеріали доповідей, представлених на IV науково-практичній конференції «Системні науки та інформатика» (1–5 грудня 2025 року, м. Київ, Україна).

Наведені доповіді, присвячені питанням системного аналізу складних систем різної природи, інтелектуальних сервіс-орієнтованих розподілених обчислень, систем і методів штучного інтелекту. Розглядаються новітні досягнення в галузі розробки та досліджень математичних методів, моделей, прогресивних інформаційних технологій для потреб освіти та науки, оборони, промисловості, економіки та навколишнього середовища.

Для фахівців в області системного аналізу, штучного інтелекту, сучасних інформаційних технологій, а також для аспірантів і студентів старших курсів вищої школи відповідних спеціальностей.

© Навчально-науковий Інститут прикладного системного аналізу КПІ ім. Ігоря Сікорського, 2025.

© Колектив авторів, 2025.

Співголови програмного комітету конференції:

Єфремов К.В.

Касьянов П.О.

Пишнограсв І.О.

Романенко В.Д.

Члени програмного комітету конференції:

Бідюк П.І.

Кисельов Г.Д.

Положаєнко С.А.

Гавриленко В.В.

Корабльов М.М.

Савченко І.О.

Гожий О.П.

Купенко О.П.

Тимошук О.Л.

Данилов В.Я.

Литвиненко В.І.

Цюцюра С.В.

Джигирей І.М.

Мухін В.Є.

Калініна І.О.

Петренко А.І.

Верстка збірки: Савченко І.О.

ЗМІСТ

Секція 01

Системний аналіз і управління

Розробка апаратно-програмного комплексу для інтелектуальної підтримки фокус-пулінгу (трекінгу)	8
<i>Бабич В.С., Савастьянов В.В.</i>	
Система оцінки кредитного ризику на основі машинного навчання та багатокритеріальної підтримки прийняття рішень	14
<i>Белік А.А., Недашківська Н.І.</i>	
Розробка рекомендаційної системи на основі моделі GraphSAGE графових нейронних мереж	20
<i>Бичков Д.В., Недашківська Н.І.</i>	
Інтелектуальна система автоматизованого складання розкладу занять з використанням адаптивних гібридних метаевристик	28
<i>Гнідобор С.М., Тимощук О.Л.</i>	
Прогнозування фінансових показників із використанням зовнішніх інформаційних факторів	33
<i>Заваріхін В.О., Тимощук О.Л.</i>	
Інтелектуальні системи для прогнозування у банківській сфері	39
<i>Міщенко А.С., Гуськова В.Г.</i>	
Інтелектуальна система підтримки прийняття рішень для автоматизованої торгівлі на фінансових ринках на прикладі криптовалют	44
<i>Мороз В.В., Канцедаль Г.О.</i>	
Система прогнозування і аналізу сезонних часових рядів у маркетингових дослідженнях	47
<i>Пащенко А.І., Тимощук О.Л., Стусь О.В.</i>	
Модель для аналізу впливу погодних факторів на формування цін зернових культур в Україні	56
<i>Тихолоз А.А., Тимощук О.Л.</i>	
Інформаційна система для оцінювання кредитних ризиків на основі ансамблевих моделей	61
<i>Халімончук Р.А., Гуськова В.Г.</i>	

Секція 02

Системний аналіз фінансового ринку

Intelligent decision support systems for financial risk modeling and analysis	68
<i>Tymoshchuk O.L., Bidyuk P.I., Levenchuk L.B., Guskova V.G.</i>	

Рекомендаційні системи з використанням текстів оглядів користувачів та ансамблювання на основі стекінгу	76
<i>Артеменко Є.В., Недашківська Н.І.</i>	
Гібридна рекомендаційна система для підбірки вакансій	81
<i>Бабанов Д.Є., Тимощук О.Л.</i>	
Моделювання зміни ціни криптоактивів в залежності від їхньої емісії	86
<i>Біль К.І., Терентьєв О.М.</i>	
Моделювання торгових сигналів на основі аналізу фінансових часових рядів за допомогою мереж Колмогорова-Арнольда	92
<i>Деркач А.С., Шаповал Н.В.</i>	
Застосування методів математичного моделювання в системах підтримки прийняття рішень з повоєнного відновлення	99
<i>Дуда В.О., Терентьєв О.М.</i>	
Спотова та ф'ючерсна торгівля як засіб хеджування для алгоритмічної торгівлі на криптобіржі у східноєвропейському регіоні	104
<i>Загарук А.Я., Касьянов П.О.</i>	
Кластеризація регіонів України за динамікою повітряних тривог методом k-середніх для часових рядів (DTW/Soft-DTW) у контексті релокацій бізнесу	109
<i>Кабанова М.І., Кузнєцова Н.В.</i>	
Інформаційна система підтримки прийняття рішень для підвищення рівня інформаційної безпеки організації	117
<i>Коломієць Д.С.</i>	
Інтелектуальна система прогнозування фінансових часових рядів та робастного аналізу прогнозованої невизначеності	124
<i>Макітрук М.Т., Селін Ю.М.</i>	
Інтелектуальна система підтримки прийняття рішень для аналізу економічних ризиків	129
<i>Мельник М.С., Левенчук Л.Б.</i>	
Поведінкові моделі оцінювання успішності студентів у смарт-класах	137
<i>Ткач В.С., Кузнєцова Н.В.</i>	
Розробка та дослідження прогнозних моделей у системах алгоритмічної торгівлі	143
<i>Ярінко Б.Б.</i>	
Система оцінювання і прогнозування врожайності зернових культур методами машинного навчання на основі супутникових і кліматичних даних	150
<i>Ярко А.Ю., Кузнєцова Н.В.</i>	

Секція 03

Інтелектуальні сервіс-орієнтовані розподілені обчислювання

Децентралізована система контролю ланцюгів поставок гуманітарної допомоги на основі технології блокчейн	156
<i>Андрушко А.І., Гіоргізова-Гай В.ІІІ.</i>	

Розробка сервіс-орієнтованої системи координації рою дронів на основі алгоритму мурашиних колоній (АСО) та стигмергійної взаємодії	160
<i>Будзинський А.В., Булах Б.В.</i>	
Дослідження впливу токенизації в малоресурсних мовах на якість перекладу	164
<i>Головач А.А., Кислий Р.В.</i>	
Моніторинг станів інформаційних ресурсів для реалізації адаптивного управління захищеністю комп'ютерних систем	169
<i>Казаков В.В., Мухін В.Є.</i>	
Оцінка ефективності роботи агентів в медичних системах	176
<i>Кайзер Н.В., Безносик О.Ю.</i>	
Private cloud computing in the context of AI	184
<i>Крайнік М.В., Письменний І.О.</i>	
Вплив стратегій декодування токенів у великих мовних моделях на якість Text-to-SQL парсингу	189
<i>Курганський Л.С., Шантала Р.В.</i>	
Інтелектуальні агенти для систем автоматизації управління бізнес-процесами	194
<i>Мухін В.Є., Лимар О.В.</i>	
Методи підтримки динамічного керування завантаженості готельного комплексу	201
<i>Мухін В.Є., Поліщук А.Б.</i>	
Система генерації шаблонів документів з LLM компонентами	206
<i>Нікітін В.С., Гіоргізова-Гай В.Ш.</i>	
Прогнозне моделювання поведінки агентів та адаптивне управління енергоресурсами для підвищення безпеки та автономності рою дронів	213
<i>Ніколайчук О.Г., Харченко К.В.</i>	
Застосування методу FAIR для кількісного аналізу ризиків у системах інформаційної безпеки	218
<i>Пакуля С.О., Мухін В.Є.</i>	
Опис та порівняльний аналіз технологій для потокової передачі відео-контенту для E-Learning платформ	223
<i>Потанчук О.М., Кисельов Г.Д.</i>	
Різонінг великих мовних моделей в медичних системах	229
<i>Прокопенко Д.М., Кислий Р.В.</i>	
Розробка засобів автоматизованого створення доступних ІТ-графічних елементів	233
<i>Ревякін Д.Е., Кисельов Г.Д.</i>	
Методи динамічної композиції платформи медичної діагностики на основі рекрутингу LLM-агентів	238
<i>Ренський М.М., Письменний І.О.</i>	
Методологія створення доступного навчального контенту для дистанційного навчання	243
<i>Сенчук І.О., Кисельов Г.Д.</i>	

Підвищення ситуаційної обізнаності рою дронів шляхом контекстного аналізу стигмергії	248
<i>Стернійчук С.Ю., Булах Б.В.</i>	
Рекомендаційна система для вибору сталої серверної архітектури веб-додатку	252
<i>Товкач М.А., Гіоргізова-Гай В.Ш.</i>	
Адаптивні засоби захисту комп'ютерних систем на основі апарата нейронних мереж	259
<i>Трень А.С., Мухін В.Є.</i>	
Аналіз методів та інструментів симуляції квантових алгоритмів	266
<i>Фіцайло Г.Ю., Кисельов Г.Д.</i>	
Інтеграція України в мережу PERPOL: глобальний контекст, національна специфіка та результати пілотних проєктів	272
<i>Цешевський В.В., Зеленков А.В.</i>	
Методи глибокого навчання нейронних мереж для аналізу та прогнозування динаміки ситуацій в імітаційних системах	276
<i>Шевчук В.В., Харченко К.В.</i>	
Децентралізована координація рою дронів через стигмергію і ройовий інтелект без GPS зв'язку	282
<i>Шостак В.О., Петренко А.І.</i>	

Секція 04

Системи і методи штучного інтелекту

Proposed architecture for product review summarisation: a modular pipeline for aspect-based insight extraction and synthesis	288
<i>Yaroslav Panasiuk</i>	
Акустична класифікація БПЛА в умовах доменного зсуву та дефіциту даних	292
<i>Вільгурін В.І., Сидорський В.С., Данилов В.Я.</i>	
Прунінг згорткових нейронних мереж за допомогою інтерпретованості мереж Колмогорова-Арнольда	297
<i>Єфанов І.С., Шаповал Н.В.</i>	
Розпізнавання іменованих сутностей в українських текстах в умовах обмеженої розмітки	303
<i>Кашперова С.В., Шаповал Н.В.</i>	
DDI-TransMDL	308
<i>Мацуєв Р.О.</i>	
Оптимізація Visual SLAM у динамічних середовищах з використанням YOLOv8 та TSDF Fusion	313
<i>Мірошніченко М.А.</i>	
Деформовані згорткові мережі	317
<i>Орленко А.С., Чумаченко О.І.</i>	
Інтелектуальна система розпізнавання емоцій за виразами обличчя на основі згорткових нейронних мереж	322
<i>Панасенко В.В., Тимошук О.Л.</i>	

Інтелектуальна система діагностики остеопорозу на основі обробки DXA зображень	329
<i>Семенов О.Г.</i>	
Виявлення контрастних об'єктів з використанням комбінації напівконтрольованого та неконтрольованого навчання	334
<i>Степанчук Д.К., Пишинограєв І.О.</i>	
Гібридні підходи для моделювання та прогнозування фінансових даних	338
<i>Харитонов С.В., Гуськова В.Г.</i>	
Адаптивний метод структурного прунінгу для оптимізації великих мовних моделей	343
<i>Швець В.О., Шаповал Н.В.</i>	

РОЗРОБКА АПАРАТНО-ПРОГРАМНОГО КОМПЛЕКСУ ДЛЯ ІНТЕЛЕКТУАЛЬНОЇ ПІДТРИМКИ ФОКУС-ПУЛІНГУ (ТРЕКІНГУ)

Бабич В.С.¹, Савастьянов В.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ varvara.babich13@gmail.com [0000-0002-9496-4454],

² vvs.in.ua@gmail.com [0000-0002-2052-0420]

У роботі досліджено підхід до автоматизації фокусування у кінематографічних системах на основі технологій комп'ютерного зору. Запропоновано апаратно-програмний комплекс, який поєднує стереокамеру Luxonis OAK-D Pro Auto Focus, модуль Raspberry Pi та систему ARRI ECS. Розроблені алгоритми визначення глибини та фільтрації даних забезпечують точність до 1 % і стабільність роботи у реальному часі, що підтверджує практичну ефективність і наукову новизну підходу.

Ключові слова: автоматизоване фокусування; комп'ютерний зір; стереозір; Luxonis OAK-D Pro; Raspberry Pi; обробка зображень.

1. ВСТУП

Фокусування під час кінозйомки є одним із ключових процесів, що визначає чіткість, глибину та художню виразність зображення. Робота з сучасною кінематографічною оптикою великого формату ускладнює цей процес, оскільки мала глибина різкості вимагає від оператора фокусу високої точності та швидкої реакції. У динамічних сценах навіть незначні відхилення або затримка корекції можуть призвести до втрати різкості кадру, що негативно впливає на якість кінопродукту та потребує повторних дублів.

З поширенням технологій комп'ютерного зору [1] та сенсорів глибини з'явилися можливості для часткової або повної автоматизації фокусування. Інтелектуальні системи здатні аналізувати об'єкт зйомки, визначати його просторове положення та коригувати фокус у режимі реального часу. Перспективним напрямом є створення відкритих, адаптивних рішень, які можуть інтегруватися з професійним кінообладнанням, зберігаючи при цьому можливість модифікацій, розширення та дослідницького використання.

У межах цієї роботи запропоновано підхід до розроблення апаратно-програмного комплексу, що поєднує засоби стереозору, алгоритми комп'ютерного зору [2] та механізми фільтрації вимірювань для автоматизованого керування фокусом камери. Метою дослідження є створення системи, здатної визначати дистанцію до об'єкта зйомки, стабілізувати отримані дані та передавати їх у систему ARRI ECS для точного позиціонування фокусу.

Об'єктом дослідження є процес керування фокусом у динамічних сценах, предметом – методи сенсорного вимірювання та алгоритми комп'ютерного зору, що забезпечують обчислення глибини та відстеження об'єктів. Наукова новизна полягає у використанні відкритої платформи стереозору в поєднанні з професійною системою керування кінообладнанням, а практична значущість – у можливості застосування комплексу для підвищення точності, стабільності та автоматизації процесу фокус-пулінгу.

2. МЕТОДИ ТА ТЕХНІЧНА РЕАЛІЗАЦІЯ

Розроблений комплекс спирається на методи комп'ютерного зору, стереовимірювань та обробки даних у режимі реального часу. Основним сенсорним модулем є стереокамера Luxonis OAK-D Pro Auto Focus [3], яка формує карту глибини шляхом порівняння зображень двох камер. Просторове положення об'єкта визначається за класичною моделлю триангуляції (1) [4]:

$$Z = \frac{f \cdot B}{d}, \quad (1)$$

де Z – відстань до об'єкта, f – фокусна відстань оптики, B – базис стереопари, d – диспаратність між відповідними точками зображень.

Для визначення області аналізу застосовується автоматичне формування ROI [5]. Система спочатку виконує детекцію обличчя, потім – детекцію тіла, а якщо об'єкт не знайдено, застосовується фіксована зона в центрі кадру.

Для формального опису механізму переходів між моделями введемо ймовірності: p_f – впевненість детекції обличчя, p_b – впевненість детекції тіла, p_0 – фіксована вага центральної зони (fallback), та пороги T_f і T_b , відповідно.

Тоді вибір ROI описується через максимізацію активної моделі (2):

$$ROI = \underset{p_0}{\operatorname{argmax}} \begin{pmatrix} p_f \cdot 1(p_f \geq T_f), \\ p_b \cdot 1(p_b \geq T_b), \end{pmatrix}, \quad (2)$$

де $1(\cdot)$ – індикаторна функція, що дорівнює 1, коли умова істинна, і 0 – інакше.

Такий підхід формалізує механізм переходу між моделями: при високій впевненості система використовує ROI по обличчю; при недостатній – переходить до ROI по тілу; за відсутності детекцій – застосовує центральний ROI.

Оскільки виміряна відстань містить випадкові коливання, спричинені шумами стереопари та локальними втраченими кореспонденціями, у системі використано комбінований підхід до цифрового згладжування. Першим етапом є медіанне фільтрування (3), яке видаляє імпульсні викиди, визначаючи проміжне значення як медіану просторового вікна області інтересу:

$$\tilde{d}_t = \operatorname{median}(d_{t-k}, \dots, d_t, \dots, d_{t+k}), \quad (3)$$

де d_t – виміряна відстань для кадру t , $\tilde{d}_t$ – згладжене значення після придушення імпульсних викидів, k – радіус просторового вікна.

Другим етапом є експоненційне згладжування (EMA) (4), яке усуває високочастотний шум і забезпечує плавну еволюцію сигналу у часі:

$$D_t = \alpha \tilde{d}_t + (1 - \alpha) D_{t-1}, \quad (4)$$

де D_t – фінальний відфільтрований вимір відстані, $\alpha \in (0,1)$ – коефіцієнт згладжування, що визначає швидкість реакції фільтра.

Комбінація цих двох етапів дозволяє приглушувати локальні викиди та стабілізувати плавні зміни сигналу. Такий підхід забезпечує надійну роботу керування фокусом, оскільки механізми фокус-пулінгу чутливі до різких стрибків відстані.

Після теоретичної обробки дані нормалізуються та перетворюються у формат, сумісний із системою ARRI ECS [7], яка здійснює керування фокусним механізмом камери ARRI Alexa Mini LF. Така інтеграція забезпечує можливість автоматизованого фокус-пулінгу на основі даних комп'ютерного зору.

Узагальнену структуру роботи розробленого комплексу подано на рис. 1.

Процес автоматизованого фокусування камери

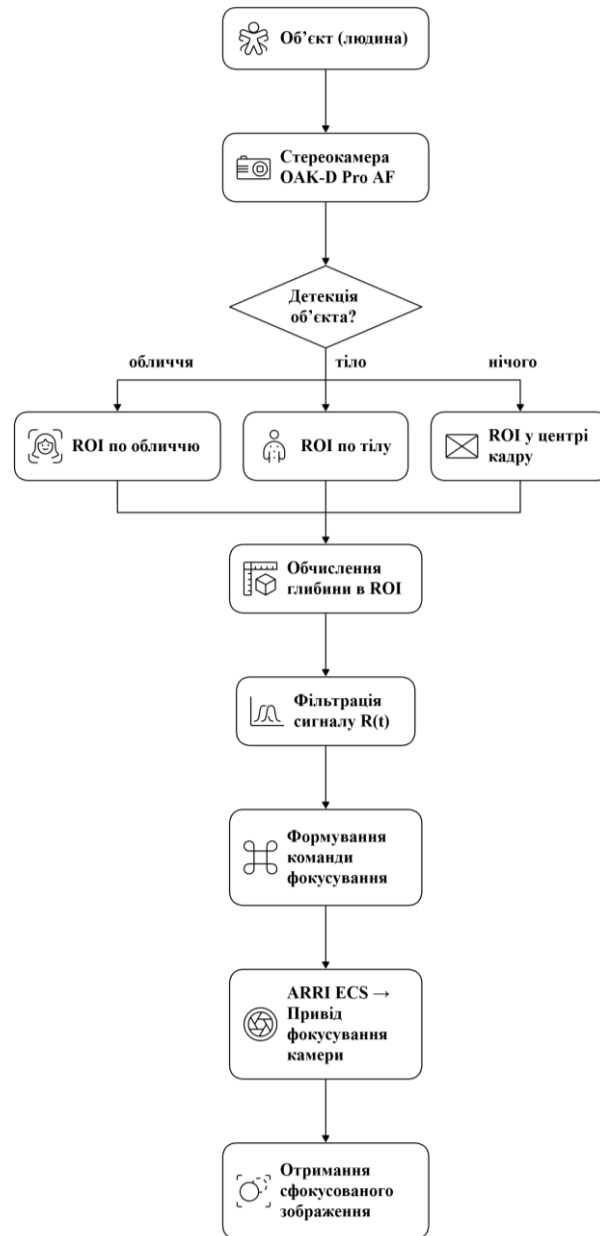


Рисунок 1. Процес автоматизованого фокусування камери

3. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

Експериментальні дослідження були проведені для оцінювання точності визначення глибини, стабільності роботи системи та здатності комплексу забезпечувати коректне фокусування у динамічних умовах. Тестування здійснювалося у контрольованих сценаріях із заздалегідь відомими еталонними відстанями та різними траєкторіями руху об'єкта.

Першим етапом було виконано порівняння виміряних значень дистанції з еталонними, що дозволило оцінити точність роботи стереовимірювань. Як видно з рис. 2, виміряна крива

достатньо добре повторює реальну траєкторію, а характерні для стереозору коливання не порушують загальної збіжності. Похибка у робочому діапазоні дистанцій залишається в межах, прийнятних для практичного використання у задачах автоматизованого фокусування.

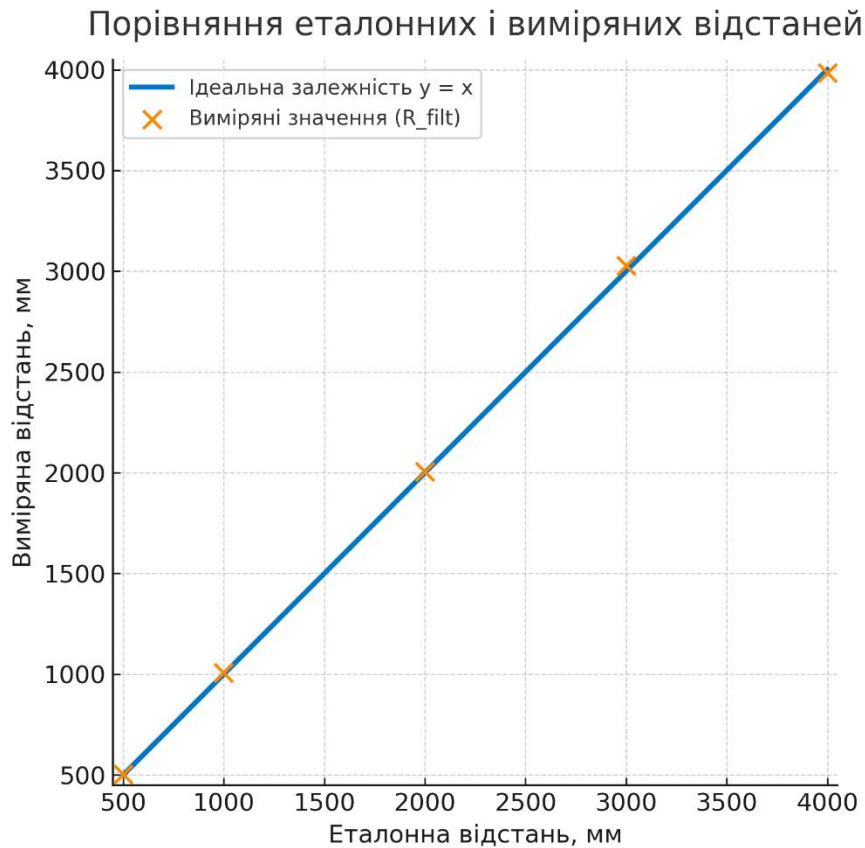


Рисунок 2. Порівняння еталонних і вимірених відстаней до об'єкта

Другий етап досліджень був присвячений оцінці впливу фільтрації на стабільність сигналу глибини. Порівняння необроблених і відфільтрованих даних (див. рис. 3) показало, що фільтрація суттєво зменшує амплітуду випадкових шумів, згладжує динаміку вимірювань і забезпечує більш передбачувану зміну дистанції при русі об'єкта. Це є критичним для приводу фокусування, який чутливий до стрибків та локальних викидів.

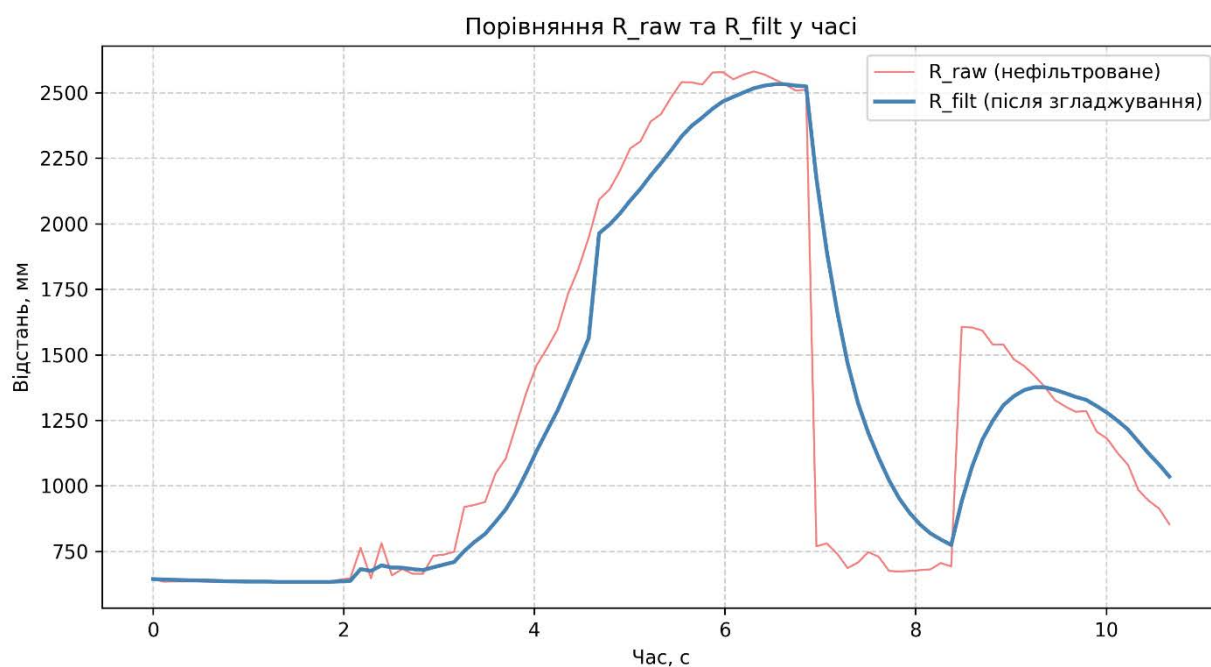


Рисунок 3. Порівняння необробленої та відфільтрованої відстані R у часі при русі об'єкта

На завершальному етапі результати точності було подано у табличній формі. Узагальнені дані (див. табл. 1) демонструють конкурентоспроможність запропонованого рішення відносно інших систем і підтверджують відповідність точності вимірювання вимогам до автоматизованого фокус-пулінгу.

Таблиця 1. Порівняння еталонних і вимірних значень відстані.

№	Еталонна відстань, м	Середнє R_filt, мм	Абсолютна похибка, мм	Відносна похибка, %
1	0.5	503	3	0.6
2	1.0	1008	8	0.8
3	2.0	2006	6	0.3
4	3.0	3028	28	0.9
5	4.0	3985	15	0.4

Отримані результати узгоджуються з тенденціями, описаними в сучасних оглядових роботах зі стереозору [8], та підтверджують ефективність поєднання методів комп'ютерного зору зі стереовимірюванням та демонструють можливість застосування розробленого комплексу як основи для автоматизованого керування фокусом у професійних кінематографічних системах.

4. ВИСНОВКИ

У роботі представлено апаратно-програмний комплекс автоматизованого фокусування, що поєднує технології комп'ютерного зору, стереовимірювання та цифрової фільтрації сигналів. Розроблена система забезпечує визначення відстані до об'єкта в реальному часі, адаптивний вибір області інтересу (ROI), стабілізацію отриманих даних та передачу параметрів у систему ARRI ECS для керування приводом фокусування професійної кінокамери.

Експериментальні дослідження засвідчили, що комплекс коректно працює в робочому діапазоні до 5 метрів, забезпечуючи середню похибку близько 5 % у межах 0,5–5 м. Розроблені методи фільтрації – медіанне згладжування та експоненційне згладжування (ЕМА) – зменшили амплітуду шумових коливань у 2–3 рази, підвищивши плавність і стабільність сигналу $R(t)$, що є критично важливим для надійного керування механізмом фокусування.

Окремо відзначимо, що система коректно працює з адаптивним вибором ROI на основі рівня впевненості детекторів: при наявності детекції використовується область обличчя, при недостатній впевненості – область тіла, а за відсутності валідних ознак – центральна зона кадру. Такий підхід забезпечує стійкість оцінки дистанції у змінних умовах освітлення, положення оператора та взаємодії зі сценою.

Отримані результати демонструють перспективність використання розробленого комплексу як основи для створення інтелектуальних систем автоматизованого фокус-пулінгу в кінематографічних застосуваннях. Розроблений підхід поєднує простоту апаратної інтеграції з високою точністю та стабільністю обчислюваної дистанції, що відкриває можливості для подальшої розробки професійних автономних інструментів керування фокусом.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Szeliski, R. Computer Vision: Algorithms and Applications (2nd ed.). – Springer, 2022.
2. Goodfellow, I., Bengio, Y., Courville, A. Deep Learning. – MIT Press, 2016.
3. Luxonis Documentation. OAK-D Pro Auto Focus and DepthAI SDK. – Luxonis, 2024.
4. Li, J. et al. Practical Stereo Matching... – CVPR, 2022.
5. Intel OpenVINO Toolkit. Face Detection and Depth Estimation Models. – Intel, 2024.
6. Zhang, K. et al. Beyond a Gaussian Denoiser. – IEEE TIP, 2017.
7. ARRI GmbH. Electronic Control System (ECS) Overview. – 2024.
8. Zhang, S. et al. Review of Stereo Matching Based on Deep Learning. – Pattern Recognition Letters, 2024.

СИСТЕМА ОЦІНКИ КРЕДИТНОГО РИЗИКУ НА ОСНОВІ МАШИННОГО НАВЧАННЯ ТА БАГАТОКРИТЕРІАЛЬНОЇ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ

Белік А.А.¹, Недашківська Н.І.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ bielik.andriy@lil.kpi.ua,

² nedashkovskaya.nadezhda@lil.kpi.ua [0000-0002-8277-3095]

Метою роботи є розробка гібридного підходу до оцінювання кредитного ризику, що поєднує моделі машинного навчання з багатокритеріальним аналізом рішень. На основі підготовлених даних кредитних заявок побудовано та порівняно кілька алгоритмів прогнозування дефолту, після чого отримані оцінки інтегровано з ключовими характеристиками позичальника методом TOPSIS. Застосування запропонованого підходу дозволяє формувати обґрунтовані кредитні рішення та підвищувати ефективність ризик-менеджменту у фінансових установах.

Ключові слова: кредитний ризик, скорингова модель, машинне навчання, ймовірність дефолту, метод багатокритеріальної підтримки прийняття рішень, оцінка кредитоспроможності.

1. ВСТУП

У сучасних умовах цифровізації та зростання складності економічних процесів питання точного та відтворюваного оцінювання кредитоспроможності позичальників набуває особливої значущості. Кредитний ризик залишається одним із ключових факторів, що визначають стабільність банківської системи, а невірна оцінка ймовірності дефолту може призвести до формування проблемних активів та погіршення фінансової стійкості установ. Традиційні методи аналізу кредитних заявок, засновані на експертних судженнях або спрощених показниках, дедалі частіше виявляються недостатніми в умовах зростання обсягів даних та динамічності зовнішнього середовища [1].

Сучасні підходи на основі машинного навчання дозволяють точніше прогнозувати ризик неповернення, враховуючи широкий спектр соціально-демографічних, поведінкових та фінансових характеристик позичальників [2]. Проте результати таких моделей потребують інтеграції з практичними вимогами кредитної політики, що включає оцінку забезпечення, галузеві ризики, дохідність кредитної угоди та інші бізнес-критерії.

У цьому контексті актуальною є розробка систем, які поєднують прогностичні можливості математичних моделей із формалізованими методами підтримки прийняття рішень. Інтеграція оцінки ймовірності дефолту з багатокритеріальним аналізом дозволяє сформувати узгоджений інтегральний показник ризику, який одночасно враховує як статистичні, так і бізнес-орієнтовані фактори.

Представлена робота спрямована на створення гібридної системи оцінювання кредитного ризику, яка поєднує моделі машинного навчання з методами багатокритеріального вибору. Система реалізує повний цикл обробки даних – від підготовки ознак і тренування моделей до формування інтегрованого скорингового рішення – та забезпечує можливість налаштування ваг критеріїв і застосування вето-правил. Практична значущість підходу

полягає у підвищенні точності скорингу та покращенні якості управління кредитними ризиками у банківських і фінтех-структурах.

2. ТЕОРЕТИЧНІ ОСНОВИ ТА МЕТОДИ ДОСЛІДЖЕННЯ

Розроблена гібридна система оцінювання кредитного ризику ґрунтується на інтеграції методів машинного навчання з багатокритеріальним аналізом рішень. На першому етапі виконується всебічна підготовка даних, що передбачає аналіз структури вибірки, перевірку повноти та узгодженості показників, виявлення аномальних значень і визначення основних закономірностей у профілі клієнтів. Подальша обробка включає очищення та уніфікацію даних, агрегування кредитних записів одного позичальника, імпутацію пропусків, перетворення категоріальних ознак у числові представлення та масштабування кількісних характеристик, а також формування цільової змінної дефолту на основі кредитної історії. Цільова змінна є бінарною, проте характеризується суттєвим дисбалансом: частка дефолтних клієнтів становить близько 2%. Для підвищення стійкості моделі до такого дисбалансу застосовано алгоритм Synthetic Minority Oversampling Technique (SMOTE), який генерує синтетичні приклади міноритарного класу. У результаті формується коректно структурований масив даних, оптимізований для побудови моделей машинного навчання та подальшого багатокритеріального аналізу.

Далі здійснюється моделювання ймовірності дефолту на основі декількох алгоритмів машинного навчання: логістична регресія, дерева рішень, ансамблеві методи (Random Forest, Balanced Random Forest, Bagging), градієнтні бустинги (XGBoost, LightGBM), багатопарова нейромережа (MLP) [3]. Найкраща модель визначається за метриками ROC-AUC та суміжними показниками на тестовій вибірці. Її прогнози значення PD використовуються як один із критеріїв у подальшому багатокритеріальному аналізі.

Механізм інтеграції побудовано на методі TOPSIS, який дає змогу поєднувати PD із важливими бізнес-показниками: рівнем річного доходу, стажом роботи, наявністю нерухомості, характером кредитної поведінки [4]. Частина ознак інтерпретується як benefit-критерії (більше – краще), інші як cost-критерії (більше – гірше). Ваги критеріїв обираються з урахуванням пріоритетності ризику, надійності та фінансової спроможності клієнта.

Після нормалізації та вагового масштабування показників метод TOPSIS формує позитивно-ідеальний та негативно-ідеальний варіанти, до яких вимірюється відстань кожного клієнта. На основі дистанцій S_i^+ та S_i^- обчислюється інтегральний індекс:

$$C_i = \frac{S_i^-}{S_i^+ + S_i^-},$$

який відображає узгодженість клієнта з «ідеальним» профілем позичальника.

На завершальному етапі формується набір порогових правил, які визначають рішення за заявкою. Зокрема, надмірно високе значення PD може виступати в ролі автоматичного вето незалежно від інших показників, тоді як високий інтегральний бал MCDM-score свідчить про загальну стабільність клієнта та перевагу позитивних характеристик у сукупному профілі.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Проаналізовано набір даних Credit Card Approval Prediction, який містить інформацію про соціально-демографічні характеристики клієнтів та їхню історію кредитних зобов'язань [5]. Набір є прикладом реальної банківської бази, очищеної від персональних даних, і призначений для побудови скорингових моделей машинного навчання. Після етапу попередньої обробки даних було сформовано вибірку, що містить повний профіль клієнтів за

соціально-демографічними, фінансовими та анкетними характеристиками. Початковий оглядовий аналіз дозволив визначити базові закономірності у структурі позичальників (рис.1). Розподіл віку, сімейного стану та наявності активів показав помітні відмінності між групами клієнтів та підтвердив релевантність цих параметрів для подальшого моделювання кредитного ризику.

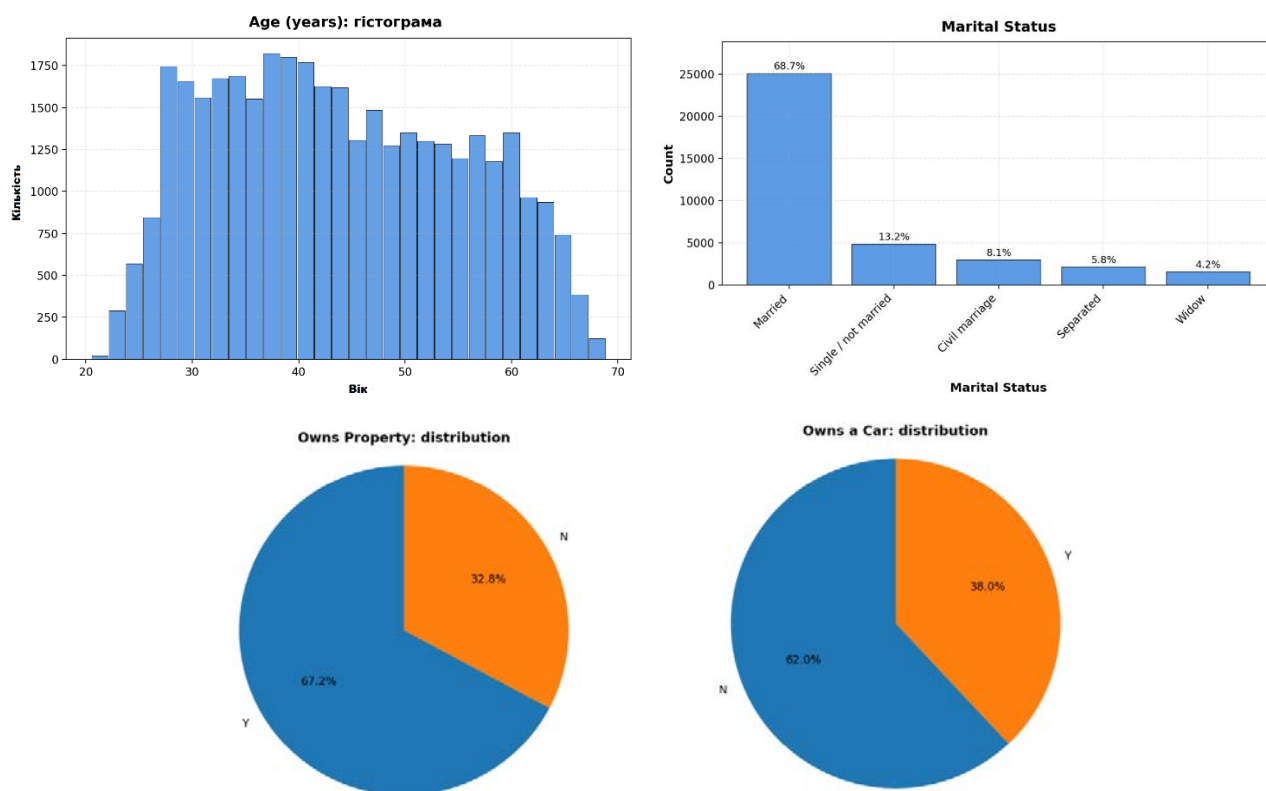


Рисунок 1. Приклад оглядових графіків (розподіл віку, сімейного стану та наявності активів)

На наступному етапі було здійснено порівняння кількох алгоритмів машинного навчання, які формують прогноз ймовірності дефолту (PD). Результати тестування моделей наведено у підсумковій таблиці метрик, де відображено їхню дискримінаційну здатність та якість ранжування позичальників за ризиком (рис. 2). Порівняння показало перевагу ансамблевих методів на основі дерев рішень, насамперед Balanced Random Forest, який продемонстрував найкращі значення прогнозовної якості.

Model	Accuracy	Precision	Recall	F1 score	AUC-ROC	Gini	KS-статистика	Brier Score
Balanced Random Forest	98,62%	98,36%	98,88%	98,62%	99,54%	99,09%	97,35%	1,44%
Bagging	98,14%	97,99%	98,30%	98,15%	99,50%	99,00%	96,46%	1,96%
XGBoost	95,76%	95,52%	96,03%	95,77%	99,17%	98,33%	91,89%	3,41%
LightGBM	94,79%	94,35%	95,29%	94,82%	98,56%	97,12%	89,79%	5,50%
Decision Tree	97,27%	97,35%	97,18%	97,26%	97,57%	95,14%	94,57%	2,70%
MLP	91,33%	90,65%	92,16%	91,40%	96,80%	93,59%	83,07%	6,60%
Logistic Regression	61,78%	62,08%	60,58%	61,32%	68,06%	36,11%	25,77%	22,21%
Naive Bayes	58,76%	56,15%	80,01%	65,99%	64,49%	28,97%	21,97%	30,96%

Рисунок 2. Порівняльна таблиця метрик моделей машинного навчання

Після вибору найкращої моделі її прогнозні значення PD було інтегровано у схему багатокритеріального аналізу TOPSIS разом із показниками доходу, стажу роботи та наявності

нерухомості. На основі MCDM-score сформовано три групи рішень: Accept, Review та Reject (рис. 3). Розподіл PD у межах цих трьох класів показав очікувану впорядкованість: клієнти класу Accept мають найнижчі значення PD, Review займає проміжне положення, а Reject характеризується найвищими ризиками.

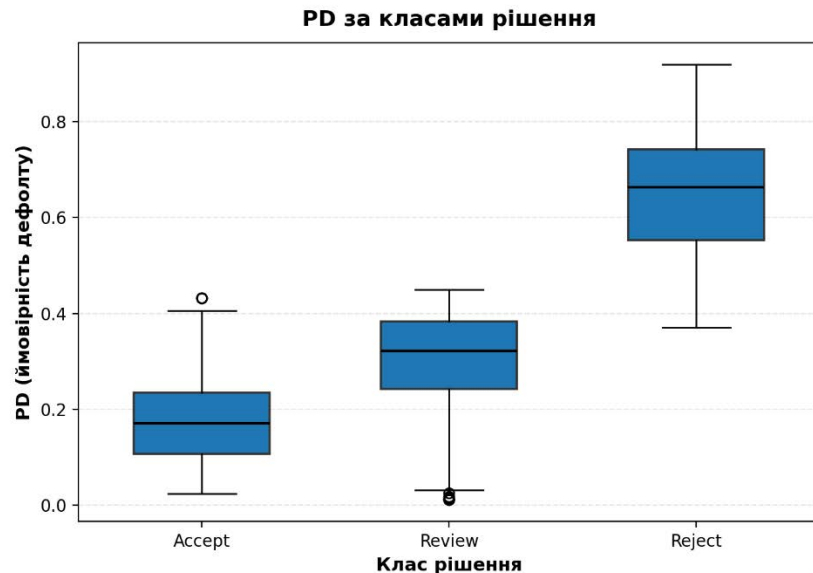


Рисунок 3. Розподіл PD за класами рішення (Accept / Review / Reject)

Важливою частиною аналізу стала оцінка взаємозв'язку між PD та MCDM-score (рис. 4). Побудована діаграма розсіювання продемонструвала стійку негативну залежність: із зростанням PD інтегральний бал зменшується, що підтверджує узгодженість моделі машинного навчання з критеріальною логікою TOPSIS.

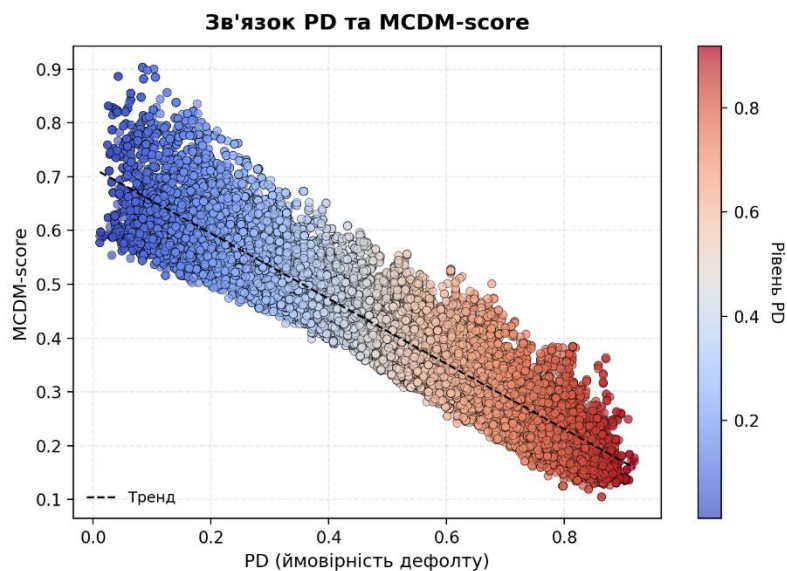


Рисунок 4. Зв'язок PD та MCDM-score

Для демонстрації практичної застосовності моделі реалізовано модуль оцінювання окремого позичальника (рис. 5). Після введення ключових характеристик клієнта система

автоматично обчислює ймовірність дефолту, показники benefit/cost-критеріїв, інтегральний MCDM-score та видає рекомендацію щодо кредитного рішення.

Анкета клієнта

Вік, років: Стаж роботи, років:

Розмір сім'ї: Стать:

Освіта:

Тип доходу:

Сімейний стан:

Тип житла:

Рівень доходу (0–100):

Рівень доходу: 50-й перцентиль

Бінарні характеристики

☒ Має авто ☒ Має нерухомість

☐ Робочий телефон ☒ Домашній телефон

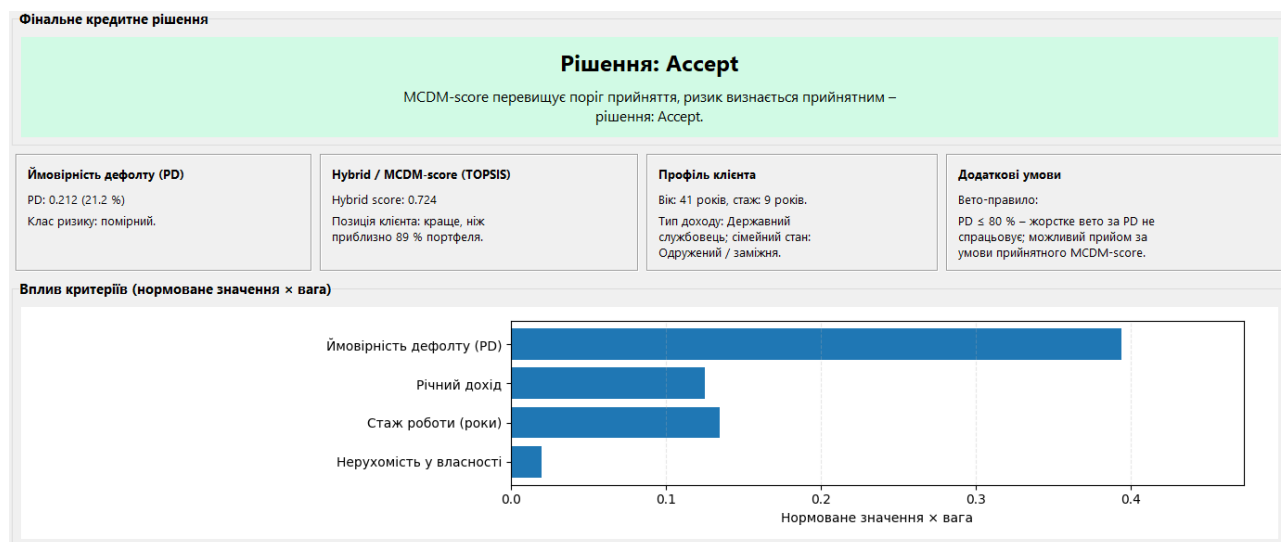
☐ Електронна пошта

Модель PD та метод MCDM

Модель ML (PD):

Метод MCDM:

(а)



(б)

Рисунок 5. Результат прогнозування (б) та оцінювання для окремого позичальника (а) в програмному модулі

4. ВИСНОВКИ

У роботі запропоновано гібридний підхід до оцінювання кредитного ризику, який поєднує прогностичні можливості сучасних алгоритмів машинного навчання з формалізованими методами підтримки прийняття рішень. На основі реального набору кредитних заявок реалізовано повний процес: від підготовки даних і аналізу їх структури до побудови моделей, що прогнозують надійність позичальників, та подальшого інтегрування цих результатів у багатокритеріальну схему оцінювання.

Порівняння декількох моделей показало, що найстабільніші й найточніші результати забезпечують ансамблеві методи на основі дерев рішень. Отримані прогностні оцінки було доповнено іншими важливими показниками – рівнем доходу, трудовим стажем та наявністю власних активів, – що дозволило сформулювати комплексний індикатор для підтримки рішення щодо видачі кредиту.

Впровадження інтегрального показника та порогових правил дало змогу формалізувати процедуру оцінювання ризику й підвищити прозорість процесу ухвалення рішення. Аналіз структури отриманих рішень підтвердив, що запропонована система послідовно розмежовує стабільних клієнтів і клієнтів із підвищеним ризиком, а розроблений інтерфейс забезпечує зручність застосування моделі у практичній діяльності аналітика.

Отримані результати свідчать про перспективність поєднання прогностних моделей із багатокритеріальними підходами в задачах кредитного скорингу. Подальший розвиток роботи може бути пов'язаний із розширенням набору критеріїв, урахуванням економічних умов та адаптацією вагових коефіцієнтів під специфіку окремих фінансових установ.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Viblyi P. I., Zharchynska A. Y. Credit risk analysis of commercial banks and methods of their assessment. *Management and Entrepreneurship in Ukraine: the stages of formation and problems of development*. 2019. Vol. 1, no. 1. P. 8-13. <https://doi.org/10.23939/smeu2019.01.008>.
2. Default avoidance on credit card portfolios using accounting, demographical and exploratory factors: decision making based on machine learning (ML) techniques / N. Sariannidis et al. *Annals of Operations Research*. 2019. Vol. 294, no. 1-2. P. 715–739. <https://doi.org/10.1007/s10479-019-03188-0>.
3. Framework for multi-criteria assessment of classification models for the purposes of credit scoring / P. Ziemba et al. *Journal of Big Data*. 2023. Vol. 10, no. 1. <https://doi.org/10.1186/s40537-023-00768-7>.
4. García V., Sánchez J. S., Marqués A. I. Synergetic application of multi-criteria decision-making models to credit granting decision problems. *Applied Sciences*. 2019. Vol. 9, no. 23. P. 1-15. <https://doi.org/10.3390/app9235052>.
5. Credit Card Approval Prediction. URL: <https://www.kaggle.com/datasets/rikdifos/credit-card-approval-prediction> (дата звернення: 10.11.2025).

РОЗРОБКА РЕКОМЕНДАЦІЙНОЇ СИСТЕМИ НА ОСНОВІ МОДЕЛІ GRAPH SAGE ГРАФОВИХ НЕЙРОННИХ МЕРЕЖ

Бичков Д.В.¹, Недашківська Н.І.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹bychkov.denys@lil.kpi.ua,

²nedashkovskaya.nadezhda@lil.kpi.ua [0000-0002-8277-3095]

Метою роботи є дослідження сучасних підходів до рекомендацій, підготовка даних великого обсягу, побудова гетерогенного графа взаємодій «користувач-артист-жанр» та розробка моделі GraphSAGE з навчанням за допомогою BPR-loss. У роботі виконано аналіз датасету Last.fm 360K, створено ембедінги вузлів графа, проведено оцінювання моделі за метриками Recall@K та NDCG@K, а також виконано візуалізацію ембедінгів і графової структури.

Ключові слова: рекомендаційні системи, графові нейронні мережі, GraphSAGE, імпліцитні дані, музичний контент, BPR-loss.

1. ВСТУП

Стрімкий розвиток музичних стримінгових сервісів зумовив потребу у високоточних і швидких рекомендаційних системах, здатних обробляти великі масиви взаємодій користувачів із контентом. Традиційні методи колаборативної та контентної фільтрації мають низку обмежень, зокрема проблему розрідженості даних і холодного старту. У цьому контексті графові нейронні мережі є перспективним інструментом, оскільки дозволяють моделювати складні зв'язки між користувачами, виконавцями та жанрами. Було виконано побудову гетерогенного графа взаємодій, розроблено та навчено модель GraphSAGE, проведено оцінювання її якості й досліджено можливості застосування отриманих результатів у вигляді стартап-проєкту. Робота спрямована на створення сучасної персоналізованої системи рекомендацій музичного контенту та аналіз її ефективності.

2. СУЧАСНИЙ СТАН ПРОБЛЕМИ ТА ПОСТАНОВКА ЗАДАЧІ

2.1. Сучасний стан проблеми

Згідно з аналізом сучасних досліджень, рекомендаційні системи для музичного контенту традиційно базувалися на методах колаборативної (CF) та контентної фільтрації (CBF). Колаборативні підходи, включно з матричною факторизацією та алгоритмами ALS чи BPR, ефективно використовують історію прослуховувань, проте мають обмеження, пов'язані з розрідженістю даних і проблемою холодного старту. Контентні системи покращують рекомендації для нових треків, оскільки використовують жанрові й аудіо-ознаки, але їхня точність часто нижча через обмежену інформативність метаданих.

Сучасний тренд у науковій літературі – перехід до гібридних рішень і графових моделей. Представлення даних у вигляді графа «користувач-трек-жанр» дає змогу врахувати багатозв'язні структури, недоступні класичним методам. Графові нейронні мережі (GNN), такі як GraphSAGE [1], GCN [2], GAT дозволяють агрегувати інформацію з сусідніх вузлів і формувати якісні латентні подання для рекомендацій. Моделі на кшталт LightGCN [3], NGCF [4] та PinSAGE [5] демонструють вищу продуктивність порівняно з традиційними CF-алгоритмами, особливо у випадках розрідженості даних або складних структур зв'язків.

Огляд наукових джерел свідчить, що GNN-методи – перспективний напрям розвитку рекомендаційних систем для музики, оскільки забезпечують масштабованість, можливість індуктивних висновків для нових треків та більш точну персоналізацію, ніж класичні підходи.

2.2. Постановка задачі

Дано: множина користувачів (users) U ; множина музичних виконавців (artists, items) I ; множина жанрових тегів (genre) P ; множина неявних взаємодій (implicit-feedback) $E=\{(u, i)\}$.

Нехай $G=(U \cup I \cup P, E_{UI} \cup E_{IP})$ – гетерогенний граф.

Потрібно: побудувати модель рекомендацій, яка для кожного користувача u формує ранжований список виконавців $i \in I$ на основі графу G .

Критерії оптимальності – мінімізація функції втрат BPR-loss та максимізація Recall@K, NDCG@K, Precision@K.

Обмеження: імпліцитні дані; холодний старт; розрідженість графа; індуктивність моделі GraphSAGE.

Тип задачі – implicit-feedback item ranking (link prediction).

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

3.1. Вибір та опис датасету

Для реалізації рекомендаційної системи було обрано масштабний відкритий набір даних Last.fm 360K, який містить історію прослуховувань понад 360 тисяч користувачів [6].

Особливістю цього датасету є те, що він відображає імпліцитні вподобання – кількість прослуховувань слугує непрямым свідченням інтересу до артиста. Відсутність прослуховування не означає негативної оцінки, а лише те, що користувач досі не натрапив на цього виконавця.

В роботі використано API Last.fm для автоматичного збору жанрових тегів для кожного артиста.

3.2. Підготовка та обробка даних для GNN-моделі

На етапі підготовки даних було реалізовано кілька ключових кроків для формування вхідного графа, який використовувався для навчання графової нейронної мережі. Основні етапи включали: очищення та фільтрацію, індексацію, побудову графа у форматі PyTorch Geometric, а також інтеграцію жанрової інформації.

На етапі очищення та фільтрації основним джерелом даних слугувала таблиця usersha1-artmbid-artname-plays.tsv [6], яка містить інформацію про кількість прослуховувань кожного виконавця кожним користувачем. Етапи очищення:

- видалено записи з відсутніми або некоректними значеннями;
- відфільтровано взаємодії з нульовою кількістю прослуховувань;
- застосовано фільтри активності (рис. 1).

Очищення дозволило зосередитись на найбільш інформативних користувачах і виконавцях, уникнувши шуму від надто рідкісних взаємодій. У результаті залишилося близько 1400 користувачів та понад 10 тисяч артистів, з приблизно 90 тисячами пар прослуховувань.

```
user_counts = listens['user'].value_counts()
artist_counts = listens['artist'].value_counts()
listens = listens[listens['user'].isin(user_counts[user_counts >= 10].index)]
listens = listens[listens['artist'].isin(artist_counts[artist_counts >= 5].index)]
```

Рисунок 1. Фільтри активності

Далі було виконано індексацію та побудову словників (рис. 2). Оскільки `user_id` та `artist_id` в датасеті є строковими (SHA1-хешами), для побудови графа їх було переіндексовано до послідовних чисел. Це значно оптимізує розмірність ембедінгів і дозволяє ефективно формувати граф.

Базова структура графа – двочастковий граф "користувач ↔ виконавець". Для цього використано формат `edge_index` бібліотеки PyTorch Geometric, де кожне ребро (u, v) представляє взаємодію між користувачем і артистом (зі зміщенням ідентифікаторів артистів на `+num_users` для унікальної нумерації) (рис. 3).

```
user2id = {u: i for i, u in enumerate(listens['user'].unique())}
artist2id = {a: i for i, a in enumerate(listens['artist'].unique())}

listens['user_id'] = listens['user'].map(user2id)
listens['artist_id'] = listens['artist'].map(artist2id)
```

Рисунок 2. Індексція та побудова словників

```
edge_index = []
for _, row in listens.iterrows():
    u = row['user_id']
    i = row['artist_id'] + num_users
    edge_index.append([u, i])
    edge_index.append([i, u])
```

Рисунок 3. Побудова графа user-item

Всі ребра приведено до `torch.Tensor` і передано у Data-структуру.

Для збагачення графа було використано API Last.fm. У коді реалізовано цикл для збору тегів артистів, після чого збережено топ-3 тегів на артиста (рис. 4).

```
API_KEY = "b25b959554ed76058ac220b7b2e0a026" # офіційний public test key
artist_tags = {}
for artist in tqdm(list(artist2id.keys())[:1000]): # обмежено до 1000 запитів
    try:
        url = f"http://ws.audioscrobbler.com/2.0/?method=artist.gettoptags&artist={artist}&api_key={API_KEY}"
        response = requests.get(url, timeout=4)
        tags = response.json().get('toptags', {}).get('tag', [])
        tag_names = [t['name'] for t in tags[:3]] # до 3 найпопулярніших тегів
        artist_tags[artist2id[artist]] = tag_names
    except:
        continue
```

100% |██████████| 1000/1000 [04:11<00:00, 3.97it/s]

Рисунок 4. Отримання тегів з API Last.fm

Після отримання тегів побудовано вузли жанрів та ребра типу `item ↔ genre` (рис. 5).

У результаті граф став гетерогенним: окрім користувачів і артистів, він включає вузли жанрів (~100–150) та відповідні ребра, що дозволяють моделі враховувати семантичну близькість між виконавцями.

Для об'єктивного оцінювання якості рекомендацій застосовано розбиття набору даних за схемою Leave-One-Out: для кожного користувача один позитивний приклад було відкладено в тестову вибірку. У коді на рис. 6 сформовано позитивні приклади.

```

#Додаємо вузли жанрів і ребра item <-> genre
genre2id = {}
genre_edges = []
for item_id, tags in artist_tags.items():
    for tag in tags:
        if tag not in genre2id:
            genre2id[tag] = len(genre2id) + num_users + num_items
        gid = genre2id[tag]
        genre_edges.append([item_id + num_users, gid])
        genre_edges.append([gid, item_id + num_users])

all_edges = edge_index + genre_edges
tensor_edge_index = torch.tensor(all_edges, dtype=torch.long).t().contiguous()
graph_data = Data(edge_index=tensor_edge_index, num_nodes=num_users + num_items + len(genre2id))

```

Рисунок 5. Створення вузлів жанрів і ребра item ↔ genre

```

user_pos_items = listens.groupby('user_id')['artist_id'].apply(set).to_dict()
all_items = torch.arange(num_items, device=device)

```

Рисунок 6. Позитивні приклади

Тестовий набір для метрик (Recall, NDCG тощо) формувався окремо, по 1 позитивному прикладу + випадково згенеровані негативні (артисти, яких користувач не слухав).

3.3. Архітектура моделі та навчання з BPR-loss

У розробленій системі рекомендацій використано підхід графових нейронних мереж (GNN) – потужного інструменту для роботи з даними, що мають структуру графа [2]. Основною ідеєю є навчання ембедінгів (векторних подань) для кожного вузла (користувача, виконавця, жанру) так, щоб у побудованому векторному просторі відображалась семантична близькість.

В основі реалізована модель GraphSAGE – один з популярних варіантів GNN, запропонований для узагальнення на великі графи [1]. На відміну від GCN [2], де кожен шар обробляє всю матрицю суміжності, GraphSAGE (рис. 7) агрегує ознаки лише з сусідів вузла, що дозволяє ефективно працювати з великими графами.

```

class RecGNN(torch.nn.Module):
    def __init__(self, num_nodes, emb_dim=64):
        super().__init__()
        self.emb = torch.nn.Embedding(num_nodes, emb_dim)
        self.conv1 = SAGEConv(emb_dim, emb_dim)
        self.conv2 = SAGEConv(emb_dim, emb_dim)

    def forward(self, edge_index):
        x = self.emb.weight
        x = self.conv1(x, edge_index)
        x = F.relu(x)
        x = self.conv2(x, edge_index)
        return x

```

Рисунок 7. Модель GraphSAGE

Модель GraphSAGE складається з наступних елементів:

- шару Embedding, який ініціалізує випадкові вектори для кожного вузла;
- двох шарів SAGEConv, які виконують агрегацію інформації з графа;
- функції активації ReLU між шарами.

Ембедінги оновлюються на основі зв'язків у графі, що дозволяє враховувати не лише безпосередні взаємодії (user–item), а й жанрові зв'язки (item–genre).

Для навчання використано функцію втрат BPR-loss, орієнтовану на рейтинг замість передбачення конкретного значення [7]. Вона підходить для неявних даних (implicit feedback), коли відомо лише, що користувач прослухав трек, але невідомо наскільки цей трек йому сподобався. Основна ідея: користувач має отримувати вищий рейтинг для обраного (позитивного) треку, ніж для випадкового (негативного). В коді реалізовано mini-batch sampling для пришвидшення навчання (рис. 8).

```
for epoch in range(1, 6):
    np.random.shuffle(users_all)
    model.train()
    total_loss = 0.0

    # обчислюємо з один на batch користувачів (перед backward).
    # Це потрібно перераховувати перед КОЖНИМ optimizer.step(), але не для кожного користувача.
    for start in range(0, len(users_all), BATCH_USERS):
        batch_users = sample_batch_users(users_all, start, BATCH_USERS)

        # відкидаємо користувачів без позитивів
        pos_items, keep_mask = sample_pos_for_users(batch_users)
        if keep_mask.sum() == 0:
            continue
        batch_users = batch_users[keep_mask]
        pos_items = pos_items[keep_mask]

        # негативи з resample
        neg_items = sample_neg_for_users(batch_users, user_pos_items)
```

Рисунок 8. Побудова міні-батчів для більш швидкого навчання

Градiente обчислюються лише для вибірки користувачів і пар (positive, negative), що дозволяє знизити споживання пам'яті.

Навчання триває кілька епох (у тестах – 5), після чого отримані ембедінги використовуються для генерації рекомендацій.

Під час тренування значення BPR-loss монотонно зменшувалося, що підтверджує збіжність моделі. Метрики Recall@10 та NDCG@10 зростали протягом епох, демонструючи покращення якості ранжування.

3.4. Оцінка моделі за метриками якості

Після навчання моделі GraphSAGE було проведено оцінку її якості за допомогою поширених метрик для імпліцитних рекомендацій (рис. 9):

- Recall@K показує, яка частка релевантних, тобто фактично прослуханих треків потрапила у топ-K рекомендацій;
- NDCG@K (Normalized Discounted Cumulative Gain) враховує порядок у списку рекомендацій: релевантні треки, які стоять ближче до початку, оцінюються вище;
- Precision@K є часткою правильних рекомендацій у топ-K списку.

```

from sklearn.metrics import ndcg_score

def evaluate_metrics(topk=10):
    recalls, precisions, ndcgs = [], [], []
    for u in list(user_pos_items.keys())[:100]:
        recs = recommend(u, topk=topk).numpy()
        actual = user_pos_items[u]
        hits = len(set(recs) & actual)
        recall = hits / len(actual) if actual else 0
        precision = hits / topk

        relevance = np.zeros(num_items)
        for i in actual:
            relevance[i] = 1
        scores = torch.matmul(item_embs, user_embs[u]).detach().numpy()
        ndcg = ndcg_score([relevance], [scores])

        recalls.append(recall)
        precisions.append(precision)
        ndcgs.append(ndcg)

    print(f"Recall@{topk}: {np.mean(recalls):.4f}")
    print(f"Precision@{topk}: {np.mean(precisions):.4f}")
    print(f"NDCG@{topk}: {np.mean(ndcgs):.4f}")

evaluate_metrics(topk=10)

```

Рисунок 9. Реалізація обчислення метрик

Метрики показали, що модель GraphSAGE на основі неявних вподобань формує релевантні рекомендації: Recall@10 у середньому становив 0.36-0.41 (тобто в ~40% випадків модель правильно вгадує хоча б 1 справді прослуханий трек у топ-10).

3.5 Візуалізація графової структури та ембедінгів

Для глибшого аналізу структури даних та ембедінгів, отриманих за допомогою графової нейронної мережі, реалізовано кілька типів візуалізацій.

Побудовано підграф на основі ребер між треками та їх жанровими тегами (рис. 10). Така візуалізація дозволяє виявити, як згруповані треки за жанрами, та які жанри є найбільш пов'язаними, допомагаючи зрозуміти жанрову кластеризацію.

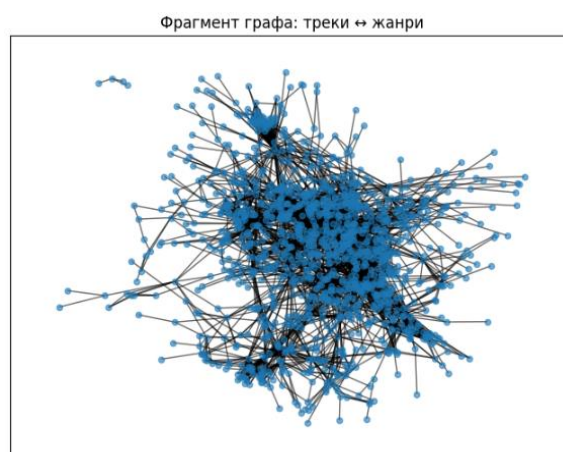


Рисунок 10. Фрагмент графа: треки ↔ жанри

Для візуалізації простору ембедінгів треків виконано зниження розмірності методом головних компонентів (Principal Component Analysis, PCA). Це дозволяє побачити, як модель "розподілила" треки у векторному просторі. Кластери на PCA-проекції відображають подібні жанри або групи треків, які часто слухаються разом.

В прикладі на рис. 11 проекція трек-ембедінгів містить один виражений кластер, що є наслідком того, що модель GraphSAGE тренувалася лише на зв'язках "користувач-артист" без жанрових ознак і тому демонструє слабку жанрову сегментацію у імпліцитних даних.

Застосування PCA до ембедінгів жанрових тегів дає змогу побачити, які жанри розміщуються поруч у векторному просторі, тобто мають подібний контекст використання (рис. 11). Такий аналіз корисний для виявлення близьких за стилем або популярністю жанрів.

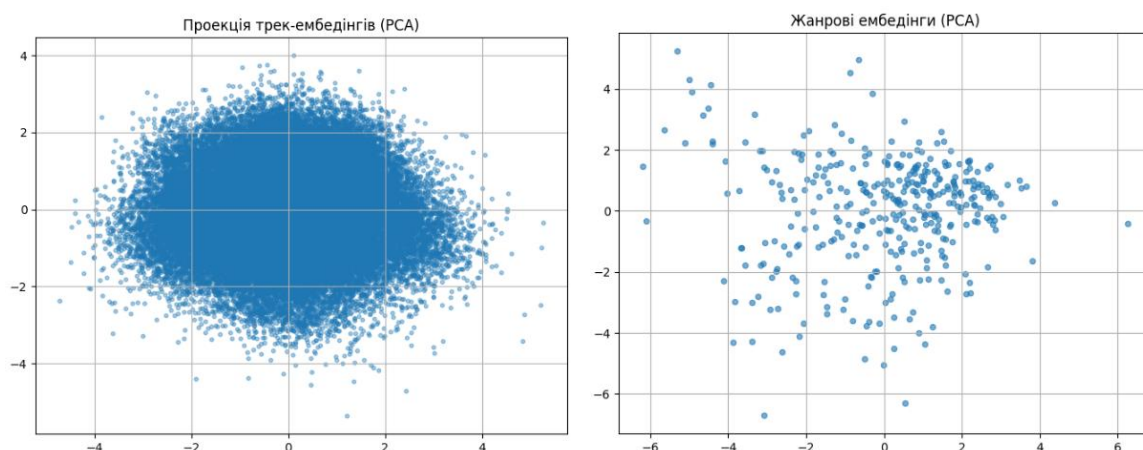


Рисунок 11. Проекції PCA

Для ілюстрації локальної взаємодії реалізовано функцію візуалізації персонального підграфа – графа взаємодій одного користувача з треками, які він слухав. Це дає уявлення про індивідуальні вподобання. На рис. 12 користувач user_0 з'єднаний з 10 треками, які він прослухав.

Даний підхід можна масштабувати, додавши зв'язки з жанрами або друзями для створення персоналізованої карти вподобань.



Рисунок 12. Жанрові ембедінги (PCA)

4. ВИСНОВКИ

У ході роботи було створено рекомендаційну систему для музичного контенту на основі моделі GraphSAGE графових нейронних мереж. Модель GraphSAGE та її навчання з BPR-loss забезпечили формування якісних ембедінгів користувачів і артистів та показали хороші результати: значення Recall@10 близько 0.4. Система працює на біпартитному графі «користувач–артист» і може бути легко розширена додатковими типами вузлів.

Розроблена модель є масштабованою та не потребує явних оцінок, що робить її придатною для реальних стримінгових сервісів. Водночас система має потенціал для вдосконалення: врахування послідовності прослуховувань, використання складніших негативних прикладів та інтеграція мультимодальних ознак. Отримані результати підтверджують перспективність використання графових нейронних мереж у музичних рекомендаціях та відкривають можливості для подальшого розвитку і практичного застосування системи.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Hamilton W., Ying R., Leskovec J. Inductive Representation Learning on Large Graphs. Proceedings of the 31st Conference on Neural Information Processing Systems (NIPS'17), 2017, Long Beach, CA, USA. DOI:10.48550/arXiv.1706.02216
2. Kipf T., Welling M. Semi-Supervised Classification with Graph Convolutional Networks. Proceedings of the International Conference on Learning Representations (ICLR'17), 2017. DOI: 10.48550/arXiv.1609.02907
3. He X. et al. LightGCN: Simplifying and Powering Graph Convolution Network for Recommendation. Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR'20), 2020, P. 639 - 648. DOI: 10.1145/3397271.3401063.
4. Wang X. et al. Neural Graph Collaborative Filtering. Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR'19), 2019, P. 165 - 174. DOI: 10.1145/3331184.3331267
5. Ying R. et al. Graph Convolutional Neural Networks for Web-Scale Recommender Systems. Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining (KDD'18), 2018, P. 974 - 983. DOI: 10.1145/3219819.3219890
6. lastfm Music Recommendation Dataset. URL: <https://www.last.fm/http://ocelma.net/MusicRecommendationDataset/lastfm-360K.html> (дата звернення 20.11.2025)
7. Rendle S. et al. BPR: Bayesian Personalized Ranking from implicit feedback. Proceedings of the 25th Conference on Uncertainty in Artificial Intelligence (UAI'09), 2009, P. 452 - 461. DOI: 10.48550/arXiv.1205.2618

ІНТЕЛЕКТУАЛЬНА СИСТЕМА АВТОМАТИЗОВАНОГО СКЛАДАННЯ РОЗКЛАДУ ЗАНЯТЬ З ВИКОРИСТАННЯМ АДАПТИВНИХ ГІБРИДНИХ МЕТАЕВРИСТИК

Гнідобор С.М.¹, Тимошук О.Л.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ hnidobor.serhii@lml.kpi.ua

У роботі розглядається задача автоматизованого складання розкладу занять у закладах вищої освіти, яка належить до класу NP-складних задач. Запропоновано інтелектуальну систему, що базується на гібридному метаевристичному підході – Adaptive Hybrid Scheduling Algorithm (AHSA). AHSA поєднує глобальний пошук на основі генетичних алгоритмів з локальною оптимізацією за допомогою імітації відпалу, доповненою адаптивним налаштуванням вагових коефіцієнтів у процесі оптимізації. Продemonстровано переваги такого підходу над класичними точними та евристичними методами на тестових прикладах.

Ключові слова: складання розкладу, метаевристики, генетичний алгоритм, імітація відпалу, адаптація параметрів, NP-складна задача.

1. ВСТУП

У сучасних умовах цифровізації та автоматизації освітніх процесів підвищуються вимоги до ефективності планування навчального навантаження. Складання розкладу занять у закладах вищої освіти є однією з найскладніших організаційних задач, оскільки вимагає одночасного врахування великої кількості обмежень: наявності аудиторій, обмежень викладачів, побажань студентів, навчального плану тощо. Задача належить до класу NP-складних, що унеможливає її точне розв'язання за прийнятний час при великій кількості вхідних параметрів.

Традиційні методи побудови розкладу – ручне складання або використання статичних евристик – дедалі частіше демонструють незадовільні результати в умовах складних і змінних вимог сучасних університетів. Це обумовлює потребу в інтелектуальних адаптивних системах, здатних гнучко реагувати на зміну вхідних даних і підтримувати стабільну якість розкладу.

Метою представленої роботи є розробка інтелектуального застосунку для автоматизованого складання розкладу занять, що базується на гібридному метаевристичному підході Adaptive Hybrid Scheduling Algorithm (AHSA). Запропонований підхід поєднує можливості глобального пошуку на основі генетичного алгоритму, локального вдосконалення через імітацію відпалу, а також модуль динамічного налаштування параметрів в процесі оптимізації. Кожен компонент спрямований на подолання окремого типу труднощів: локальні мінімуми, низька швидкодія, нестабільність параметрів.

Запропонована система демонструє здатність генерувати високоякісні розклади за короткий час, зберігаючи узгодженість з жорсткими обмеженнями й водночас оптимізуючи м'які критерії. Таким чином, AHSA може бути використаний як основа для впровадження автоматизованих модулів у сучасних системах управління вищою освітою.

2. ОПИС РОЗРОБЛЕНОГО АЛГОРИТМУ

Для вирішення задачі автоматизованого складання розкладу занять було розроблено власний алгоритм, оскільки існуючі методи, попри свою ефективність у певних випадках, мають низку суттєвих обмежень. Більшість класичних підходів або не забезпечують належної гнучкості при зміні вхідних параметрів, або демонструють погану масштабованість на великих обсягах даних. Це особливо критично для сучасних університетів, де кількість дисциплін, викладачів, студентських груп та обмежень постійно змінюється.

Запропонований адаптивний гібридний метаевристичний алгоритм (Adaptive Hybrid Scheduling Algorithm, AHSA) поєднує сильні сторони генетичних алгоритмів, імітації відпалу та адаптивної логіки регулювання параметрів. Основна ідея полягає в поступовій еволюції розкладу: від випадкових генерацій до вдосконалених, оптимізованих варіантів. На першому етапі створюється початкова популяція рішень, де кожне рішення (розклад) кодується як хромосома. За допомогою операторів відбору, кросоверу та мутацій створюються нові покоління. Потім для кожного отриманого варіанта застосовується процедура локального вдосконалення через механізм імітації відпалу, що дозволяє покращити рішення, уникнувши локальних мінімумів.

Ключовою особливістю AHSA є адаптивний модуль, який відстежує прогрес пошуку та автоматично змінює ваги обмежень і параметри мутації. Наприклад, якщо протягом кількох поколінь не спостерігається поліпшення, система посилює штрафи за м'які порушення або збільшує ймовірність випадкових змін у хромосомах. Це дозволяє алгоритму зберігати високу ефективність у динамічному середовищі, де параметри задачі можуть змінюватись у процесі її розв'язання.

Алгоритм AHSA функціонує як багаторівнева система з трьома взаємопов'язаними контурами (рис. 1):

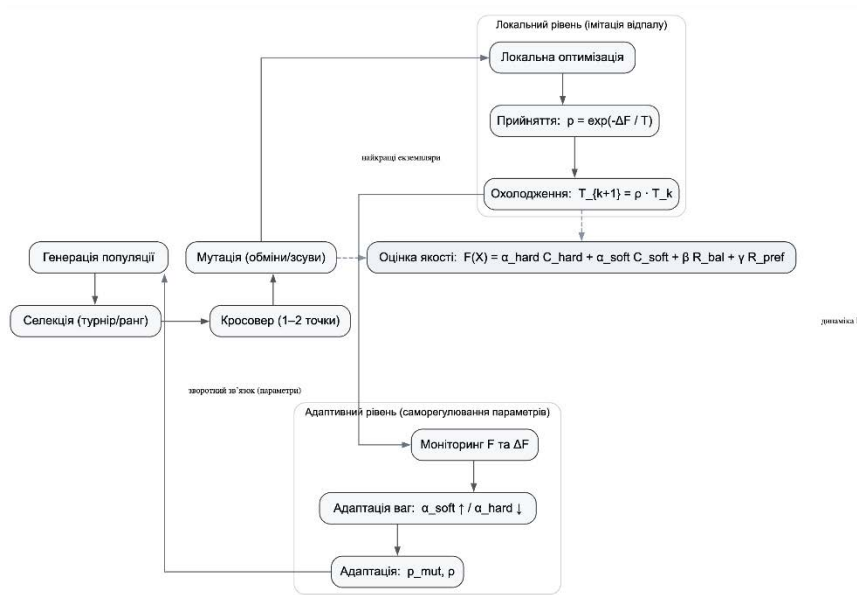


Рисунок 1. Загальна схема алгоритму

Алгоритм AHSA складається з трьох рівнів, кожен з яких виконує свою специфічну роль у процесі оптимізації розкладу.

На глобальному рівні відбувається формування та еволюція популяції рішень за допомогою механізмів генетичного алгоритму. Початкова популяція генерується випадково

або з використанням евристичних шаблонів. Подальші покоління створюються через оператори селекції (турнірного відбору), кросоверу (одноточкового або двоточкового) та мутації. Кожен індивід, тобто розклад, оцінюється за функцією якості, яка враховує кількість порушень жорстких обмежень, ступінь невідповідності м'яким критеріям та значення штрафних коефіцієнтів. Найкращі варіанти розкладів із глобального рівня передаються на локальний рівень, де вони вдосконалюються за допомогою методу імітації відпалу (Simulated Annealing). Тут відбуваються локальні перестановки елементів розкладу з ймовірністю прийняття гіршого рішення відповідно до температурного режиму. Цей механізм дозволяє уникати «застрягання» в локальних мінімумах і забезпечує більш глибоке дослідження обмежених областей рішень.

Третій - адаптивний рівень відслідковує темп покращення функції якості. Якщо протягом певної кількості ітерацій не спостерігається прогресу, активується адаптивний механізм: змінюється ймовірність мутацій, вагові коефіцієнти, параметри температури або охолодження. Це дозволяє зберігати ефективність пошуку навіть в умовах динамічної зміни вхідних даних, таких як кількість курсів, слотів чи обмежень.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для оцінки ефективності запропонованого алгоритму було проведено серію експериментів.

Для цього використовувався підготовлений набір даних UniTime University Course Timetabling Dataset, що містив 45 курсів, 28 викладачів, 12 аудиторій та 25 часових слотів.

Мета експерименту полягала у перевірці того, наскільки адаптивний гібридний підхід може покращити якість розкладу та зменшити обчислювальні витрати порівняно з традиційними методами.

Ефективність роботи алгоритмів оцінювалася за трьома основними показниками, що відображають як швидкодію, так і якість сформованого розкладу.

Першим критерієм виступав час обчислення, який визначався як кількість секунд, необхідних для досягнення стабільного стану алгоритму або завершення заданої кількості ітерацій. Цей показник характеризує продуктивність методу та дозволяє оцінити його придатність для інтерактивних освітніх систем, де швидкість побудови розкладу є критичною.

Другим критерієм була кількість конфліктів, тобто число порушень жорстких обмежень у розкладі. До таких порушень належали випадки, коли одна аудиторія використовувалася одночасно для кількох курсів, викладач був призначений на два заняття в один і той самий час або коли розклад студентської групи містив накладання занять. Зменшення кількості таких конфліктів свідчить про підвищення структурної узгодженості отриманого рішення.

Головним показником оцінки виступала інтегральна якість розкладу (Fitness Score), яка відображає загальну відповідність рішення як жорстким, так і м'яким обмеженням. Ця метрика виражається у відсотках і розраховується за формулою (1):

$$Q = 100 - (\alpha \cdot C_h + \beta \cdot C_s), \quad (1)$$

де Q – інтегральна якість розкладу (у відсотках, від 0 до 100); C_h – частка порушених жорстких обмежень; C_s – частка порушених м'яких обмежень; α ; β – вагові коефіцієнти, що визначають відносну важливість кожного типу обмежень.

Підсумкове значення Q подається у шкалі від 0 до 100 %, де більші значення свідчать про кращу якість сформованого розкладу. Результати порівняльного аналізу наведено в таблиці 1.

Таблиця 1. Результати порівняння алгоритмів

Алгоритм	Час обчислення, с	Кількість конфліктів	Якість розкладу Q , %
Цілочисельне лінійне програмування (ILP)	312.4	0	83
Жадібний алгоритм (Greedy)	17.9	6	64
Генетичний алгоритм (GA)	69.8	0	88
AHSA (розроблений алгоритм)	45.6	0	94

Як видно з таблиці, класичний метод ILP забезпечує безконфліктний розклад, проте має найбільшу обчислювальну складність. Час його виконання перевищує 300 секунд навіть для обмеженого набору даних, що робить цей підхід непридатним для використання у реальних умовах, де розклади мають сотні або тисячі комбінацій курсів. Крім того, отримане рішення має нижчу якість (83 %), оскільки ILP не враховує динамічних пріоритетів та зручності розподілу занять протягом тижня.

Жадібний алгоритм демонструє найвищу швидкість побудови (менше 20 секунд), однак формує розклад із шістьма конфліктами, тобто з порушенням жорстких обмежень. Через це інтегральна оцінка якості знижується до 64 %.

Генетичний алгоритм (GA) показує хороший компроміс між якістю та швидкодією: за 70 секунд він формує безконфліктний розклад із якістю 88 %. Проте алгоритм має схильність “застрягати” в локальних оптимумах і потребує додаткових ітерацій для стабілізації рішення.

Розроблений алгоритм AHSA продемонстрував найкраще співвідношення між якістю та часом виконання. Він формує безконфліктний розклад лише за 45,6 секунди, при цьому досягаючи інтегральної якості 94 %. Вищий результат пояснюється поєднанням глобального еволюційного пошуку з локальною оптимізацією методом імітації відпалу, що дозволяє алгоритму не лише уникати локальних мінімумів, а й ефективно “доводити” часткові рішення до більш збалансованих конфігурацій.

Крім того, AHSA автоматично адаптує параметри роботи – наприклад, змінює інтенсивність мутацій і швидкість охолодження залежно від темпу збіжності, що дозволяє досягати стабільного рішення швидше, ніж у класичному GA.

Таким чином, AHSA забезпечує підвищення якості розкладу на 6–10 % порівняно з класичним генетичним алгоритмом, при цьому зменшуючи середній час виконання приблизно на 35 %. У порівнянні з ILP швидкодія покращується майже у 7 разів, що робить AHSA придатним для інтеграції в автоматизовані системи формування навчальних розкладів вивсвітських ринків. Новини охоплювали широкий спектр тем, зокрема економічні звіти, аналітику фінансових ринків, макроекономічні прогнози, а також важливі політичні події, які можуть впливати на динаміку валютного курсу. Зібраний новинний контекст був попередньо оброблений для видалення шуму, таких як дублікат новин або несуттєва інформація, що дозволило покращити точність моделі.

3. ВИСНОВКИ

У роботі розглянуто задачу автоматизованого складання розкладу занять у закладах вищої освіти та запропоновано інтелектуальний підхід до її розв’язання на основі адаптивного гібридного метаевристичного алгоритму AHSA. На відміну від класичних

методів, що або гарантують оптимальність ціною значних обчислювальних витрат (цілочисельне лінійне програмування), або працюють швидко, але формують розклади з конфліктами та низькою якістю (жадібні алгоритми), АНСА поєднує глобальний пошук генетичного алгоритму, локальне вдосконалення методом імітації відпалу та адаптивне налаштування параметрів у єдиній багаторівневій схемі.

Проведене експериментальне дослідження на тестовому наборі даних показало, що розроблений алгоритм забезпечує безконфліктний розклад із найкращим значенням інтегральної якості $Q=94\%$ за часу обчислення 45,6 с. Порівняно з класичним генетичним алгоритмом це дає приріст якості на 6–10 % та скорочення середнього часу виконання приблизно на 35 %, а у порівнянні з ІЛР час розрахунку зменшується майже у сім разів. Отримані результати підтверджують ефективність поєднання еволюційного пошуку з локальною оптимізацією та динамічним налаштуванням параметрів.

Адаптивний модуль АНСА продемонстрував здатність автоматично реагувати на зміну динаміки збіжності: у фазах повільного покращення посилюються штрафи за порушення м'яких обмежень і зростає інтенсивність мутацій, тоді як у фазах стабільного прогресу параметри повертаються до базових значень. Це забезпечує стійкість алгоритму до локальних мінімумів та дозволяє підтримувати високу якість розкладу навіть за зміни вхідних даних (кількість курсів, груп, аудиторій).

Отже, АНСА можна розглядати як перспективну основу для впровадження у складі модулів автоматизованих систем управління навчальним процесом у закладах вищої освіти. Подальший розвиток роботи може бути пов'язаний із масштабуванням алгоритму на більші реальні датасети, інтеграцією з інформаційними системами університетів та розширенням набору врахованих м'яких обмежень (індивідуальні побажання викладачів, міжкампусні переміщення тощо).

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Pinedo M. *Scheduling: Theory, Algorithms, and Systems*. – 5th ed. – New York : Springer, 2016. – 670 p.
2. Burke E. K., Petrovic S. Recent research directions in automated timetabling // *European Journal of Operational Research*. – 2002. – Vol. 140, No. 2. – P. 266–280.
3. Lewis R. A survey of metaheuristic-based techniques for university timetabling problems // *OR Spectrum*. – 2008. – Vol. 30, No. 1. – P. 167–190.
4. Colomi A., Dorigo M., Maniezzo V. Metaheuristics for high-level timetabling // *European Journal of Operational Research*. – 1998. – Vol. 109, No. 1. – P. 115–124.
5. UniTime University Course Timetabling Dataset [Електронний ресурс]. – Режим доступу: https://www.unitime.org/uct_dataset.php – Дата звернення: 23.10.2025.

ПРОГНОЗУВАННЯ ФІНАНСОВИХ ПОКАЗНИКІВ ІЗ ВИКОРИСТАННЯМ ЗОВНІШНІХ ІНФОРМАЦІЙНИХ ФАКТОРІВ

Заваріхін В.О.¹, Тимошук О.Л.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ zavarikhin.v@gmail.com, ² tymoshchuk.oksana@lkl.kpi.ua [0000-0003-1863-3095]

Сучасні методи прогнозування фінансових показників, крім класичного аналізу часових рядів, часто спираються на різні зовнішні інформаційні чинники. Метою цієї роботи є дослідження впливу таких факторів, як новинний сентимент, рух ціни акцій компаній-конкурентів, інсайдерські угоди та передбачення галузевих експертів, на точність прогнозування ціни акції Apple Inc. (AAPL). Результатом роботи є розроблена модель на основі штучної нейронної мережі LSTM, яка демонструє значно кращі результати у порівнянні з підходом, який використовує виключно історичні дані.

Ключові слова: фінансові показники, прогнозування, нейронні мережі, LSTM, зовнішні інформаційні фактори.

1. ВСТУП

У сучасному світі прогнозування динаміки руху фінансових показників є одним з найбільш важливих та складних викликів для всіх учасників фінансового ринку. Точне передбачення ціни активів має вирішальне значення для мінімізації ризиків та максимізації прибутків як приватних інвесторів, так і великих корпорацій. Однак, класичні підходи, які засновані виключно на історичних даних, є недостатньо ефективними. Вони не враховують зовнішні чинники, що можуть значно впливати на ринкові тенденції. У цьому контексті інформація з відкритих джерел стає невід'ємною складовою для створення точних моделей прогнозування.

Метою даної роботи є дослідження можливостей покращення якості прогнозування за рахунок інтеграції зовнішніх інформаційних чинників. Для досягнення цього було проведено факторний аналіз впливу сентименту новин, стану акцій компаній-конкурентів, біржових індексів, інсайдерських угод та передбачень аналітиків на ціну акцій компанії Apple Inc протягом 2025 року. На основі цих даних було реалізовано рекурентну нейронну мережу типу LSTM, яка здійснює прогнозування на основі різних наборів вхідних даних, та порівняно їхні результати точності у порівнянні з базовою моделлю.

У результаті було запропоновано підхід, який комбінує історичні дані ціни активу з інформацією про сентимент новин навколо компанії та зміну цін акцій компаній-конкурентів. Експериментальні дослідження підтверджують, що використання цих зовнішніх даних у поєднанні з традиційним підходом на основі часових рядів значно покращує точність прогнозування. Запропонована реалізація моделі розширює можливості до прогнозування фінансових показників та створює підґрунтя для використання її для прийняття більш ефективних інвестиційних рішень.

2. МЕТОДИ ПРОГНОЗУВАННЯ ФІНАНСОВИХ ПОКАЗНИКІВ

Моделі на основі штучних нейронних мереж все частіше стають основою для прогнозування фінансових показників. Завдяки їхній здатності враховувати складні зв'язки та обробляти великі обсяги даних з багатьма різними факторами вони показують кращі результати роботи у порівнянні з класичними методами на основі аналізу часових рядів (наприклад, ARIMA або GARCH). Одним з найбільш перспективних та важливих моделей такого роду є рекурентні нейронні мережі (RNN). Однак, класичні реалізації цих моделей мають ряд проблем, таких як затухання градієнта, що зменшує їх ефективність за наявності довгострокових залежностей у часових рядах. Для вирішення цього недоліка використовують вдосконалені моделі, наприклад, LSTM, яку було покладено в основу даної роботи [1].

Модель довгої короткострокової пам'яті (Long Short-Term Memory, LSTM) була розроблена для подолання недоліків класичних RNN, яка дозволяє працювати з часовими рядами, які містять довгострокові залежності. Ця архітектура включає в себе спеціалізовані комірки пам'яті та керуючі ворота (вхідні, вихідні та ворота забування), які дозволяють зберігати та використовувати інформацію з попередніх кроків як додатковий контекст для прогнозування. Цей механізм дозволяє враховувати довгострокові тренди та циклічності у даних, що в свою чергу покращує загальну якість моделі [1].

Архітектура класичної моделі типу LSTM зображена на рисунку 1.

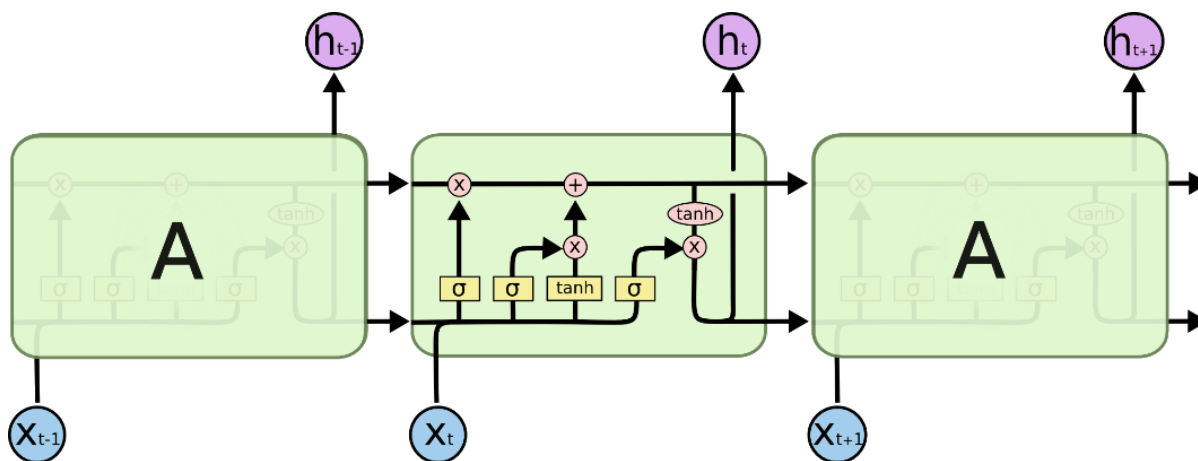


Рисунок 1. Архітектура LSTM [2]

Для доповнення даних історичних цін було проаналізовано вплив настрою новин, зміни ціни акцій компаній-конкурентів, ринкових індексів, даних інсайдерських транзакцій та очікувань від галузевих експертів.

Факторний аналіз проводився за допомогою таких методів як кореляція Пірсона, метод взаємної інформації та важливість у моделі випадкового лісу (Random Forest).

Кореляція Пірсона – метод, який оцінює рівень лінійного зв'язку між ознакою та цільовою змінною (коефіцієнт від -1 до 1 , 0 – відсутність лінійної кореляції) [3].

Показник взаємної інформації – це показник залежності між двома змінними, який показує, скільки інформації про одну змінну дає спостереження іншої [4].

Важливість у моделі випадкового лісу – показник, який обчислюється як середнє зменшення «нечистоти» у вузлах дерев [5].

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

3.1. Вхідні дані

За основу прогнозування було взято дані про ціни акцій компанії Apple Inc. (тікер AAPL). Ця компанія була обрана через її високу ліквідність, доступність якісних історичних даних, а також стабільний торговий календар без довгих пропусків. Для дослідження було проаналізовано проміжок з 1 січня по 7 жовтня 2025 року, які мають як періоди стрімкого росту, так і падіння, що дозволяє моделі навчатися на різносортних даних та враховувати більш складні залежності.

Історичні ціни акцій компанії Apple зображені на рисунку 2.



Рисунок 2. Динаміка ціни акцій Apple Inc.

Для доповнення моделі зовнішніми інформаційними факторами було використано інформацію про сентимент новин, які були пов'язані з компанією. Серед цих даних було виділено ряд авторитетних ресурсів, які мають більший вплив на настрій людей та ринок загалом. Прикладами таких джерел є Reuters, Bloomberg, AP news, CNBC, NY Times, прес-релізи Apple та інші. Для врахування більшого вкладу цих каналів у загальний сентимент, для кожного торгового дня було розраховано середнє значення тональності авторитетних джерел та решти новин.

Задля зменшення волатильності даних було використано згладження за допомогою ковзного середнього з вікном у 7 днів.

Також до аналізу було включено інформацію про ціни компаній-конкурентів (Google, Microsoft, Meta, Nvidia), біржових індексів (S&P500 та NASDAQ), інсайдерські транзакції та зміни прогнозів аналітиків.

Усі значення було приведено до єдиного числового формату та нормалізовано за допомогою Standard Scaler. Дані, які припадали на неторгові дні (вихідні або свята), було перенесено як властивість для прогнозування значення наступного торгового дня.

3.2. Факторний аналіз

Для всіх перелічених властивостей було проведено факторний аналіз на основі трьох методів: кореляції Пірсона, показника взаємної інформації та важливості у моделі випадкового лісі. Для визначення найважливіших ознак було розраховано середній ранг серед усіх тестів.

Отриманий результат показаний у таблиці 1.

Таблиця 1. Факторний аналіз змінних

Властивість	Кореляція	Взаємна інформація	Важливість Random Forest	Середній ранг
Ціна Google	0,756437	0,82949	0,518602	1,66
S&P500	0,630473	0,865253	0,060282	2,66
NASDAQ	0,625203	0,864397	0,124087	2,66
Згладжена тональність новин	0,553717	0,28872	0,142742	4,33
Ціна Microsoft	0,247558	0,658477	0,077615	5
Ціна Meta	0,485552	0,510957	0,034408	5,66
Ціна Nvidia	0,449697	0,540621	0,03036	6
Тональність новин	0,189739	0,023514	0,010099	8
Зміна рекомендацій аналітиків	0,117996	0,003966	0,001219	9
Інсайдерські транзакції	0,040424	0	0,000586	10

За результатами у таблиці 1 видно, що найбільш важливими виявились ціни акцій пов'язаних компаній, ринкових індексів та згладженої тональності новин. Натомість, зміна цінових цінкових таргетів та рекомендацій аналітиків показали слабкий вплив на цільову змінну. Частково така ситуація може пояснюватись доволі малою кількістю даних у цих змінних, а також слабким впливом на короткострокову динаміку цільової змінної.

Для полегшення роботи моделі та зменшення ризику до перенавчання було прийнято рішення не використовувати всі ознаки, а залишити тільки кілька найбільш впливових. Для цього було проаналізовано взаємну кореляцію властивостей, та виявлено, що ціни акцій компаній-конкурентів та показники біржових індексів мають сильну лінійну залежність.

Результат у вигляді кореляційної матриці зображено на рисунку 4.

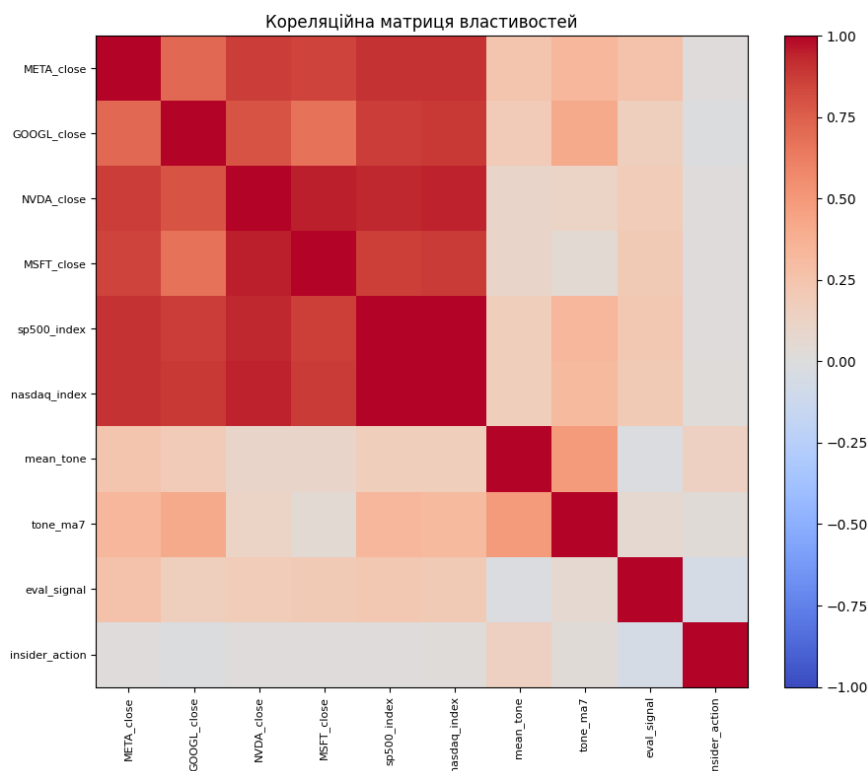


Рисунок 4. Кореляційна матриця властивостей

Для запобігання використанню надлишкової інформації було сформовано агреговану ознаку, яка узагальнює сигнал корельованих змінних – відсотковий приріст ціни за день по вибраних компаніях-конкурентах. Такий підхід дозволив моделі враховувати загальну ринкову динаміку в межах технологічного сектору, при цьому не перенавантажувавсь подібними сигналами.

3.3. Результат прогнозування

Для оцінки можливості покращення результатів прогнозування було реалізовано 5 моделей з різними наборами даних:

- базова модель, яка враховує тільки історичні ціни;
- модель із додаванням згладженої тональності новин;
- модель із додаванням змін цільової ціни та рекомендацій від аналітиків;
- модель із додаванням з агрегованої денної зміни цін пов'язаних компаній;
- модель, яка включає в себе і тональність, і агреговані дані пов'язаних компаній.

Дані інсайдерської торгівлі було виключено через занадто малий вплив на цільову змінну та дуже малу кількість даних (всього 7 днів мають дані про такі операції).

Результат прогнозування різних моделей показано на рисунку 5 та в таблиці 2.

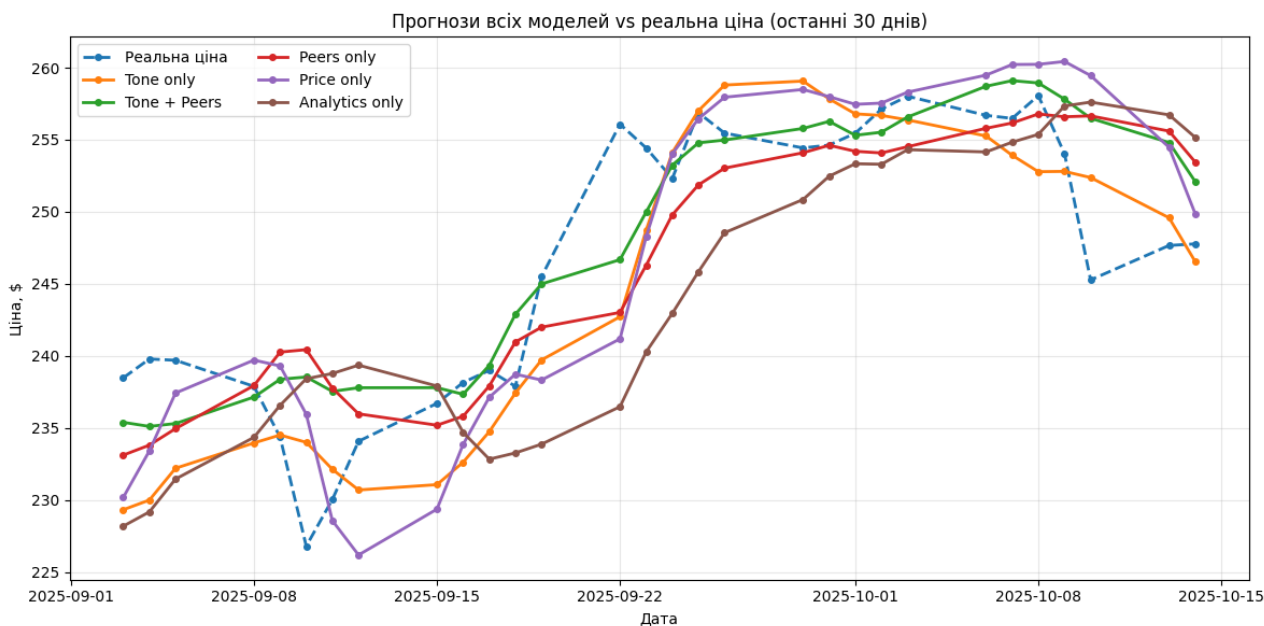


Рисунок 5. Графік порівняння реальних та прогнозованих цін

Таблиця 2. Характеристика побудованих моделей

Модель	RMSE	MAE	MAPE, %	R ²
Тональність та дані пов'язаних компаній	4,64	3,44	1,41	0,76
Тільки тональність	5,13	4,04	1,65	0,71
Тільки дані пов'язаних компаній	5,58	4,21	1,72	0,66
Базова модель	5,89	4,59	1,88	0,62
Тільки зміни рекомендацій аналітиків	8,08	6,77	2,75	0,28

Результати порівняльного аналізу в таблиці 2 свідчать про перевагу інтеграції зовнішніх інформаційних факторів над використанням виключно історичних даних. Моделі з додаванням додаткового контексту тональності новин та даних зміни цін акцій пов'язаних компаній мають кращі результати прогнозування, ніж базова LSTM. Найвища точність була досягнена при поєднанні цих даних, де модель показала найнижчі похибки (RMSE = 4,64; MAE = 3,44, MAPE = 1,41%) та найвище значення коефіцієнта детермінації ($R^2 = 0,76$). Водночас, модель, що була доповнена виключно змінами рекомендацій та цільових цін від галузевих аналітиків, продемонструвала найгірші показники навіть у порівнянні з базовим варіантом. Такий ефект може бути пов'язаний низькою частотою оновлень цих даних (за весь досліджуваний період відбулося всього 37 змін) та слабким впливом цих даних на короткострокову динаміку ціни, перетворюючи їх на інформаційний шум.

4. ВИСНОВКИ

Проведений факторний аналіз та результати експериментів підтверджують, що інтеграція зовнішніх інформаційних факторів може значно підвищувати якість прогнозу моделі LSTM у порівнянні з базовим підходом на основі виключно історичних цін. Навіть окреме додавання даних про тональність новин та динаміки зміни цін компаній-конкурентів покращило точність прогнозування, однак найкращою виявилась модель, яка комбінувала ці дані. Водночас, необхідно обережно підходити до вибору додаткових чинників, адже на прикладі моделі, яка підкріплювалась виключно змінами рекомендацій аналітиків, було виявлено тенденцію до погіршення якості прогнозу за наявності даних, які мали слабкий вплив на короткострокову динаміку.

Запропонований у роботі підхід показав свою ефективність як інструмент прогнозування фінансових показників, який дозволяє враховувати не лише внутрішні чинники, а і зовнішні фактори, які впливають на ринкову динаміку. Отримані результати слугують підґрунтям для подальших досліджень, таких як використання більш довгострокових факторів, наприклад, макроекономічних показників, розширення новинних джерел, вдосконалення методів обробки неструктурованих даних, а також застосування більш складних архітектур нейронних мереж, наприклад, на основі трансформерів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Dixon M., London J. (2021). Financial Forecasting With α -RNNs: A Time Series Modeling Approach. *Frontiers in Applied Mathematics and Statistics*. Vol 6, article 551138. URL: <https://doi.org/10.3389/fams.2020.551138>
2. Colah's blog. Understanding LSTM Networks. 2015. URL: <https://colah.github.io/posts/2015-08-Understanding-LSTMs/>
3. Swapnilbobe. Pearson's Correlation. 2021. URL: <https://medium.com/analytics-vidhya/pearsons-correlation-b6ea5cb0eb24>
4. Scikit-learn documentation. Mutual information URL: https://scikit-learn.org/stable/modules/generated/sklearn.feature_selection.mutual_info_classif.html
5. Scikit-learn examples. Feature importances with a forest of trees. URL: https://scikit-learn.org/stable/auto_examples/ensemble/plot_forest_importances.html

ІНТЕЛЕКТУАЛЬНІ СИСТЕМИ ДЛЯ ПРОГНОЗУВАННЯ У БАНКІВСЬКІЙ СФЕРІ

Міщенко А.С.¹, Гуськова В.Г.

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ mischenkoanton7@gmail.com

У роботі розглянуто підхід до формування та прогнозування клієнтського портфелю банку на основі методів машинного навчання та інтелектуальних систем підтримки прийняття рішень. Наведено етапи підготовки даних, побудови ознак, сегментації клієнтів за допомогою алгоритму KMeans та розробки моделей класифікації для прогнозування ймовірності відкриття депозитів і отримання кредитів. Отримані результати демонструють можливість підвищення точності прогнозів порівняно з традиційними підходами.

Ключові слова: банківська аналітика, прогнозування, машинне навчання, сегментація, клієнтський портфель, інтелектуальні системи.

1. ВСТУП

У сучасних умовах цифрової трансформації банківського сектору питання формування та управління клієнтським портфелем набуває особливої актуальності. Банки працюють з великими масивами даних про клієнтів, їх фінансову поведінку, історію взаємодії з продуктами та каналами обслуговування. Ефективний аналіз цих даних дозволяє не лише підвищувати прибутковість, але й зменшувати ризики, пов'язані з неповерненням кредитів, відтоком клієнтів або неефективним використанням ресурсів.

Традиційні статистичні методи часто виявляються недостатньо гнучкими для роботи з високимірними та гетерогенними даними, що характерні для клієнтських баз банків. У зв'язку з цим зростає роль інтелектуальних систем та алгоритмів машинного навчання, які здатні автоматично виявляти приховані закономірності, сегментувати клієнтів за поведінковими ознаками та прогнозувати їх подальші дії. Особливого значення набувають моделі класифікації, що дозволяють оцінювати ймовірність відкриття депозиту чи залучення кредиту в розрізі кожного клієнта.

У рамках даної роботи розглядається програмна реалізація повного циклу побудови системи прогнозування клієнтського портфелю – від підготовки вихідних даних до сегментації клієнтів, навчання моделей та інтерпретації отриманих результатів.

2. ПОСТАНОВКА ЗАДАЧІ ТА ВИХІДНІ ДАНІ

2.1. Мета та задачі дослідження

Метою дослідження є розробка та використання інтелектуального інструментарію для аналізу та прогнозування клієнтського портфелю банку на основі методів машинного навчання. Для досягнення поставленої мети необхідно вирішити такі основні задачі: визначити релевантні вхідні ознаки, забезпечити коректну підготовку даних, побудувати сегментаційну модель клієнтської бази та розробити класифікаційні моделі для прогнозування цільових показників.

У роботі розглядаються дві ключові задачі прогнозування. Перша пов'язана з оцінюванням ймовірності відкриття клієнтом строкового депозиту, друга – з прогнозуванням ймовірності оформлення кредиту. Обидві задачі мають вигляд бінарної класифікації, де результатом є належність клієнта до однієї з двох груп: «буде користуватися продуктом» або «не буде». Такі прогнози є важливими для планування ресурсів банку, таргетованих маркетингових кампаній та управління ризиками.

2.2. Характеристика вихідних даних

Вихідні дані являють собою табличну вибірку, що містить інформацію про соціально-демографічні характеристики клієнтів, параметри банківських продуктів та історію транзакційної активності. До набору ознак входять, зокрема, вік, сімейний стан, рівень доходів, наявність поточних кредитів, тип платіжної картки, рівень використання кредитного ліміту, частота та обсяги операцій. Окремі стовпці відповідають цільовим змінним – факту наявності строкового депозиту та факту користування кредитними продуктами.

Оскільки в реальних банківських системах особлива увага приділяється захисту персональних даних, при підготовці вибірки було закладено принципи анонімізації клієнтських ідентифікаторів. Для цього створено спеціальне поле з маскованим ідентифікатором клієнта, в якому зберігаються лише перші та останні символи, а проміжні замінюються на «*». Такий підхід дозволяє з одного боку зберегти можливість відстежувати записи одного клієнта, а з іншого – унеможливити ідентифікацію конкретної особи.

3. МЕТОДИ ТА ПРОГРАМНА РЕАЛІЗАЦІЯ

3.1. Попередня обробка та побудова ознак

Попередня обробка даних включає заповнення пропусків, кодування категоріальних змінних та масштабування числових ознак. Для числових полів застосовано медіанне заповнення пропусків та стандартизацію, що забезпечує однаковий масштаб змінних і сприяє стабільнішій роботі моделей. Категоріальні ознаки було перетворено за допомогою one-hot кодування, при цьому використано обмеження на максимальну кількість категорій, що дозволяє уникнути надмірного розширення простору ознак.

Особливу увагу приділено виключенню змінних, які можуть призводити до витoku інформації про цільову змінну. Зокрема, з набору ознак було вилучено технічні поля, що напряду або опосередковано містять підказку щодо факту відкриття продукту або результатів минулих маркетингових кампаній. Такий підхід дає змогу отримати більш реалістичну оцінку якості моделей та підвищити їх придатність до використання у виробничих умовах.

3.2. Сегментація клієнтів методом KMeans

Для сегментації клієнтів використано алгоритм KMeans, який дозволяє розбити вибірку на однорідні кластери за сукупністю поведінкових та соціально-демографічних ознак (рис. 1). Перед сегментацією дані було перетворено за допомогою побудованого препроцесора, що забезпечило коректний масштаб ознак та усунення впливу різниці в одиницях виміру. Оптимальну кількість кластерів визначено на основі аналізу коефіцієнта силуету та графіка «лікоть», які відображають баланс між компактністю кластерів та відстанню між ними.

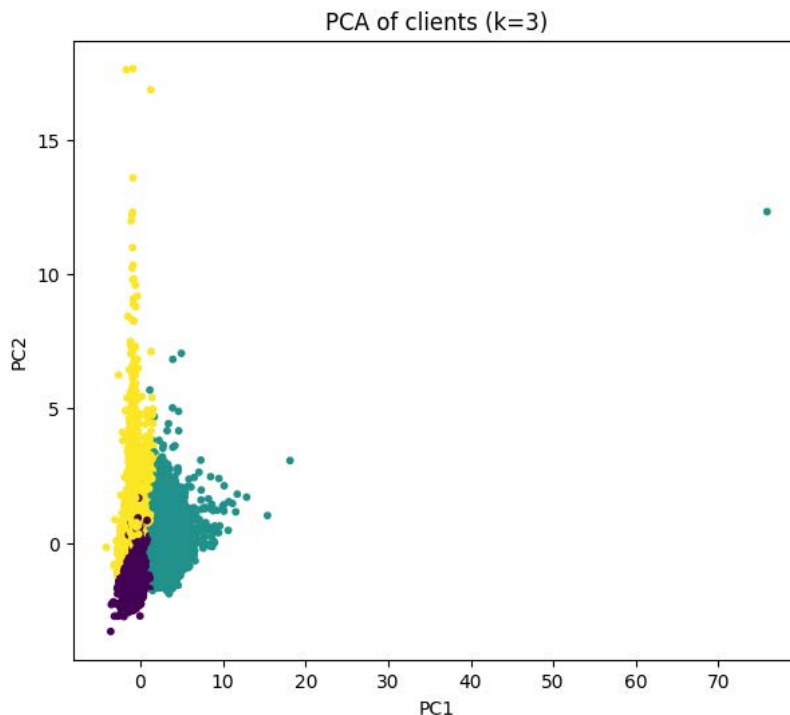


Рисунок 1. Візуальне відображення сегментації клієнтів

У результаті було отримано декілька сегментів, які відрізняються за середнім рівнем доходів, віковою структурою, активністю користування продуктами та схильністю до депозитних чи кредитних операцій. Наприклад, один із сегментів характеризується молодшою аудиторією з невисокими доходами, але високою транзакційною активністю, тоді як інший сегмент включає клієнтів середнього віку з вищими доходами та інтенсивним використанням кредитних продуктів (табл. 1). Таке профілювання дозволяє банку більш точно налаштовувати маркетингові стратегії.

Таблиця 1. Сегментація клієнтів

cluster	age	balance	day	campaign
0	34.0606	941.32	16.09	2.82
1	39.4215	1324.66	13.93	2.13
2	51.4592	1923.72	16.22	3.06

3.3. Побудова моделей прогнозування

Для розв'язання задач прогнозування було обрано декілька моделей машинного навчання, серед яких логістична регресія, випадковий ліс (Random Forest) та градієнтний бустинг. Навчання моделей здійснювалося за допомогою конвеєрів (Pipeline), що поєднують у собі препроцесор ознак та безпосередньо алгоритм класифікації. Такий підхід спрощує повторне навчання на нових даних і зменшує ризик помилок при ручному виконанні окремих етапів.

Якість моделей оцінювалася за стандартними метриками бінарної класифікації: точністю, повнотою, F1-мірою та площею під ROC-кривою (табл. 2). Результати експериментів показали, що ансамблеві методи, зокрема випадковий ліс та градієнтний бустинг, забезпечують кращий компроміс між чутливістю та специфічністю порівняно з базовою логістичною регресією. Отримані значення F1-міри та ROC-AUC свідчать про можливість практичного використання моделей для підтримки прийняття рішень у банку.

Таблиця 2. Порівняння показників точності моделей

	task	model	accuracy	F1	precision	recall	Roc_auc	y_test_mean
0	B_term_deposit	LogReg	0.8993	0.3469	0.4364	0.2878	0.7217	0.0928
1	B_term_deposit	RandomForest	0.9104	0.1450	0.6377	0.0818	0.7372	0.0928
2	B_term_deposit	GradBoost	0.9112	0.1858	0.6279	0.1090	0.7472	0.0928
3	C_take_credit	LogReg	0.7720	0.8218	0.8251	0.8186	0.8270	0.6424
4	C_take_credit	RandomForest	0.7852	0.8398	0.8060	0.8767	0.8400	0.6424
5	C_take_credit	GradBoost	0.7883	0.8433	0.8041	0.8865	0.8492	0.6424

4. РЕЗУЛЬТАТИ СЕГМЕНТАЦІЇ ТА ПРОГНОЗУВАННЯ

4.1. Характеристика отриманих сегментів

На основі результатів роботи алгоритму KMeans було виділено кілька стійких сегментів клієнтів, для яких розраховано середні значення ключових числових показників та визначено домінуючі категорії за якісними ознаками. Це дозволило сформувати узагальнені «портрети» клієнтів для кожного сегменту. Зокрема, один сегмент об'єднує молодь з відносно невисокими доходами, яка активно використовує платіжні картки для повсякденних операцій, але рідше користується кредитами. Інший сегмент представлений клієнтами середнього віку з вищими доходами, високою кредитною активністю та схильністю до відкриття депозитів.

Ще один сегмент формується за рахунок більш старших клієнтів, які мають стабільні доходи та значні залишки на рахунках, але демонструють низьку транзакційну активність. Для цієї групи більш релевантними є продукти, пов'язані зі збереженням та нарощенням капіталу. Таким чином, сегментація дозволяє ідентифікувати групи з різною чутливістю до депозитних та кредитних пропозицій, що є важливим для таргетованого маркетингу.

4.2. Інтерпретація прогнозних результатів

Після навчання моделей для кожного клієнта було розраховано ймовірність відкриття депозиту та отримання кредиту. На основі показників точності кожної з моделей, які були використані, була обрана з найкращими результатами. Додатково було проаналізовано розподіл прогнозів у розрізі сегментів, що дозволило визначити, які кластери є найбільш перспективними для просування певних продуктів.

4.3. Результати роботи

Отриманим результатом роботи є сформована сегментована база клієнтів із зазначенням ймовірності того, що кожен клієнт відкриє депозит чи оформить кредит (табл. 3). Для кожного сегменту система автоматично генерує перелік клієнтів із найвищими прогнозними значеннями, що дозволяє банку оперативно визначати найбільш перспективні цільові групи. Такий підхід спрощує аналіз великих вибірок і забезпечує можливість подальшої персоналізації пропозицій для кожної групи. На основі отриманих прогнозів банк може ефективно планувати маркетингові кампанії та підвищувати результативність комунікації з клієнтами.

Таблиця 3. Приклад результату роботи програми

	Client id masked	Proba deposit
34017	217****5931	0.9475
31283	385****0760	0.9350
33953	798****0249	0.9275
34047	653****6346	0.9250
39763	363****5578	0.9225

5. ВИСНОВКИ

У роботі представлено підхід до побудови інтелектуальної системи аналізу та прогнозування клієнтського портфелю банку, що поєднує методи попередньої обробки даних, сегментації та машинного навчання. Запропоновано структуру програмної реалізації, яка включає модулі побудови ознак, сегментації клієнтів методом KMeans та навчання моделей класифікації для прогнозування ймовірності відкриття депозитів і отримання кредитів.

Результати експериментальних досліджень показали, що використання ансамблевих алгоритмів, таких як випадковий ліс та градієнтний бустинг, забезпечує високу якість класифікації та може бути ефективно застосоване у банківській практиці. Сегментація клієнтів дозволяє виділити групи з різним рівнем схильності до використання певних продуктів, що створює додаткові можливості для таргетованого маркетингу та оптимізації ризиків.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Breiman L. Random Forests. Machine Learning. 2001. Vol. 45, No. 1. P. 5–32.
2. Friedman J.H. Greedy Function Approximation: A Gradient Boosting Machine. Annals of Statistics. 2001. Vol. 29, No. 5. P. 1189–1232.
3. Pedregosa F. et al. Scikit-learn: Machine Learning in Python. Journal of Machine Learning Research. 2011. Vol. 12. P. 2825–2830.
4. Hosmer D.W., Lemeshow S., Sturdivant R.X. Applied Logistic Regression. 3rd ed. Wiley, 2013.
5. Kaufman L., Rousseeuw P.J. Finding Groups in Data: An Introduction to Cluster Analysis. Wiley, 2005.
6. Rousseeuw P.J. Silhouettes: A graphical aid to the interpretation and validation of cluster analysis. Journal of Computational and Applied Mathematics. 1987. Vol. 20. P. 53–65.
7. Hair J.F. et al. Multivariate Data Analysis. 8th ed. Cengage Learning, 2019.
8. Box G.E.P. et al. Time Series Analysis: Forecasting and Control. 5th ed. Wiley, 2016.
9. Gujarati D.N., Porter D.C. Basic Econometrics. 5th ed. McGraw-Hill, 2009.
10. Hand D.J., Mannila H., Smyth P. Principles of Data Mining. MIT Press, 2001.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ АВТОМАТИЗОВАНОЇ ТОРГІВЛІ НА ФІНАНСОВИХ РИНКАХ НА ПРИКЛАДІ КРИПТОВАЛЮТ

Мороз В.В.¹, Канцедаль Г.О.²

Національний технічний університет України
«Київський політехнічний інститут ім. Ігоря Сікорського», Київ, Україна

¹ moroz.vlada@lil.kpi.ua, ² g.kantsedal@protonmail.com

У роботі досліджено підходи до створення інтелектуальної системи підтримки прийняття рішень для автоматизованої торгівлі на фінансових ринках із фокусом на ринку криптовалют. Розроблена система поєднує сучасні алгоритми машинного навчання для прогнозування цінних коливань криптовалют. У дослідженні реалізовано інтеграцію технічних індикаторів та проведено оцінювання точності моделей. Запропонована система забезпечує підвищення точності прогнозування, сприяє ефективнішому прийняттю торгових рішень і може бути використана як аналітичний інструмент у сфері фінансових технологій.

Ключові слова: система підтримки прийняття рішень, машинне навчання, криптовалюти, прогнозування, фінансові ринки.

1. ВСТУП

Робота присвячена розробленню інтелектуальної системи підтримки прийняття рішень (СППР), орієнтованої на автоматизовану торгівлю на фінансових ринках із фокусом на криптовалютний сегмент. Актуальність теми визначається високою динамічністю криптовалютних активів, їхньою нестабільністю та складністю прогнозування [1]. Традиційні аналітичні методи все частіше виявляються недостатніми для забезпечення своєчасного й обґрунтованого прийняття торгових рішень, що зумовило активний розвиток алгоритмічних торгових систем на основі машинного навчання та глибинного прогнозування. У такому середовищі особливу цінність мають моделі, здатні працювати з великими масивами даних, формувати прогнози короткострокової динаміки та надавати користувачу інтерпретовані результати.

2. ПОСТАНОВКА ЗАДАЧІ ДОСЛІДЖЕННЯ

Метою дослідження є створення інтелектуальної системи підтримки прийняття рішень для автоматизованої торгівлі на фінансових ринках, зокрема криптовалют. Результатом роботи є формування рекомендацій щодо купівлі або продажу активів, що дозволяє мінімізувати ризики та підвищити прибутковість торгових стратегій. Об'єктом дослідження є процес прийняття рішень у крипторгівлі, а предметом – методи машинного навчання для побудови інтелектуальних систем підтримки рішень.

3. ОСНОВНІ МОДЕЛІ

Для побудови прогнозної системи було обрано сучасні моделі та архітектури, здатні забезпечити високу точність аналізу часових рядів і гнучкість у динамічних ринкових умовах. Основу підходу сформували нейромережеві моделі N-BEATS [2] та N-HiTS [3], які належать до класу глибинних архітектур прямого поширення та спеціалізуються на задачах

прогнозування. Їх ключовою перевагою є здатність працювати без ручного конструювання ознак, автоматично виокремлюючи тренди, періодичності та інші приховані закономірності.

Побудована система приймає рішення на основі ймовірного прогнозу, де моделі формують три квантили зміни ціни, за якими обчислюються медіанний апсайд, медіанний даунсайд, ширина сигналу та витрати на угоду. Далі перевіряється, чи є сигнал достатньо сильним і асиметричним: якщо ширина діапазону перевищує мінімальний поріг, а очікуване зростання після врахування витрат суттєво переважає потенційне падіння, система відкриває позицію LONG. Якщо ж за тих самих умов очікуване падіння є більшим за апсайд, відкривається позиція SHORT. У всіх випадках, коли сигнал слабкий, симетричний або не перекидає витрати, система обирає режим FLAT і утримується від будь-яких угод.

Унікальність запропонованої СППР полягає у використанні квантильних смуг (q10–q90) замість точкових прогнозів, що дозволяє приймати рішення з урахуванням ширини та асиметрії очікуваного діапазону змін ціни після врахування витрат. Система використовує окремі моделі для верхньої та нижньої меж, актив-специфічні ознаки без витоку майбутнього та модульну архітектуру під BTC/ETH/ADA. Для підвищення надійності застосовано захист від артефактів і прозорі візуалізації прогнозів, що робить рішення інтерпретованим та відтворюваним.

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для оцінювання ефективності СППР було проведено симуляцію, у межах якої виконано бектест на тестовому відрізку з кроком 24 години. Використано ринкові та новинні дані для криптоактивів BTC, ETH та ADA. Первинні ринкові ряди мають хвилинну роздільність та містять стандартні поля OHLCV, а також кількість транзакцій. На кожному кроці система обирала одну з трьох дій – LONG, SHORT або NEUTRAL – після чого результат порівнювався з фактичною 24-годинною дохідністю активу. Всі показники розраховувалися з урахуванням транзакційних витрат у розмірі 3 б.п. за повний цикл входу/виходу. Для порівняння запропонованої системи було використано дві базові стратегії: Buy&Hold, яка постійно утримує актив, та Random, у якій вибір дії випадковий, але частота угод підібрана так, щоб відповідати активності СППР. Для кожної стратегії обчислено такі метрики: частота торгів (trade_rate), частка виграшних кроків (win_rate), денний Sharpe-подібний індикатор, максимальна втрата (maxDD) та сукупна дохідність (total_return). Порівняльні результати наведено в таблиці 1.

Таблиця 1. Порівняння результатів моделей класифікації

Актив	Стратегія	trade_rate	win_rate	Sharpe	maxDD	total_return
BTC	СППР	0.8241	0.4824	2.8495	0.3077	2.4838
	Buy&Hold	1.0000	0.5101	1.1661	0.2828	0.6189
	Random	0.8166	0.4171	0.3724	0.3454	0.0759
ETH	СППР	0.8542	0.4394	0.9241	0.6214	0.7146
	Buy&Hold	1.0000	0.5257	0.0724	0.6387	-0.2454
	Random	0.8378	0.4559	2.3082	0.3299	4.9680
ADA	СППР	0.8293	0.4521	1.3924	0.4883	1.4353
	Buy&Hold	1.0000	0.4910	1.3755	0.5682	1.4672
	Random	0.8024	0.3683	-0.7494	0.7687	-0.7320

Отримані результати демонструють, що для BTC запропонована СППР суттєво перевищує альтернативи: модель забезпечує найвищий Sharpe-показник (2.85), майже вчетверо більшу сукупну дохідність порівняно з Buy&Hold (2.48 проти 0.62) та зберігає

прийнятний рівень максимальної втрати. Це свідчить про здатність системи коректно реагувати на зміни ринку та уникати небажаних просідань. Для ETH результати більш змішані: СППР демонструє кращий Sharpe (0.92 проти 0.07 у Buy&Hold) і позитивну сумарну дохідність, тоді як стратегія утримання за цей період є збитковою. Водночас випадкова стратегія з рівною частотою угод показала нетипово високий результат, що підкреслює специфічно високу волатильність активу на досліджуваному інтервалі та можливу чутливість моделі до вибору порогів. Для ADA СППР демонструє порівнювану ефективність із Buy&Hold і значно випереджає Random, забезпечуючи стабільний профіль ризику-прибутковості. Це показує, що для активів зі спокійнішими трендами модель залишається конкурентною та здатною генерувати послідовні результати.

5. ВИСНОВКИ

У роботі було розроблено та досліджено систему підтримки прийняття рішень для прогнозування короткострокової динаміки криптовалют та формування торгових дій на основі ймовірнісних моделей. Запропонована СППР поєднує квантильне прогнозування, модульну архітектуру під окремі активи та механізм рішень, який враховує як асиметрію можливих змін ціни, так і транзакційні витрати. На відміну від традиційних підходів, модель оперує не точковими оцінками, а повними діапазонами майбутніх значень, що дозволяє підвищити стійкість до волатильності та шуму.

У роботі було проаналізовано сучасні моделі прогнозування часових рядів, зокрема архітектури N-BEATS та N-HiTS, а також систематизовано технічні індикатори, які відіграють важливу роль у виявленні ринкових закономірностей. Створено симулятор для проведення бектесту з дискретним прогнозним кроком у 24 години, розраховано ключові метрики ефективності та виконано порівняння з базовими стратегіями Buy&Hold і Random.

Результати симуляції показали, що СППР демонструє високу ефективність на ринку BTC, забезпечуючи найкращий баланс між ризиком і дохідністю, конкурентність на ETH та стабільність результатів на ADA. Це підтверджує працездатність запропонованого підходу та його здатність адаптуватися до різних ринкових режимів. Унікальність системи проявляється у використанні квантильних смуг, роздільного моделювання верхньої та нижньої меж розподілу, захисті від артефактів часових рядів і прозорій інтерпретації прогнозів.

Попри отримані результати, дослідження відкриває можливості для подальшого розвитку. До потенційних напрямів удосконалення належать оптимізація параметрів для окремих прогнозних горизонтів, удосконалення моделі настроєння на основі новинних даних та вивчення міжринкових зв'язків шляхом включення інших криптовалют як додаткових ознак.

Таким чином, запропонована система підтримки прийняття рішень є ефективним, модульним і розширюваним інструментом, що поєднує сучасні методи прогнозування та практичні механізми формування торгових рішень. Вона може слугувати основою для подальших досліджень та вдосконалення алгоритмічних стратегій на фінансових ринках.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Köchling, G., & Schmidtke, P. (2022). The impact of Bitcoin futures trading on volatility. *Finance Research Letters*, 47, 102599.
2. Matos R., Roque L., Cerqueira V. N-BEATS-MOE: N-BEATS with a Mixture-of-Experts Layer for Heterogeneous Time Series Forecasting [Electronic resource]. – 2025.
3. Challu, C., Olivares, K. G., Oreshkin, B. N., Ramirez, F., Canseco, M. M., & Dubrawski, A. (2022). N-HiTS: Neural Hierarchical Interpolation for Time Series Forecasting. *International Conference on Machine Learning (ICML)*.

СИСТЕМА ПРОГНОЗУВАННЯ І АНАЛІЗУ СЕЗОННИХ ЧАСОВИХ РЯДІВ У МАРКЕТИНГОВИХ ДОСЛІДЖЕННЯХ

Пащенко А.І.¹, Тимошук О.Л.², Стусь О.В.³

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ skullbasher696@gmail.com, ² oxana.tim@gmail.com [0000-0003-1863-3095], ³ stus@kpi.ua

Метою роботи є дослідження проблеми прогнозування і аналізу сезонних часових рядів. Особливо у випадках, коли сигнал нестационарний або містить циклічні коливання, то задача прогнозування постає перед дослідником цікавою та складною проблемою, у тому числі у контексті маркетингових досліджень для здійснення прогнозів у питаннях бізнесу.

Наведено огляд деяких відомих методів прогнозування часових рядів, а також розглянуто підходи до виявлення сезонних коливань у часовому ряді. Запропоновано алгоритм прогнозування на основі методу спектрального сингулярного аналізу. Представлено результати прогнозування вибраних часових рядів за допомогою методу спектрального сингулярного аналізу, а також відомих методів на основі сезонних моделей. Виконано порівняльний аналіз методів, використовуючи створену програму у вигляді веб-додатку.

Ключові слова: маркетингові дослідження, стратегія, часові ряди, тренд, сезонність, оцінка, прогнозування, модель.

1. ВСТУП

На сьогодні маркетингові дослідження є невід'ємною складовою будь-якого бізнес-плану. Кінцева мета будь-якого маркетингового дослідження полягає у формуванні ефективної стратегії та тактики компанії з урахуванням як реальних і потенційних ринкових умов і факторів, так і її поточної позиції та перспектив розвитку.

У діяльності підприємства важливе місце займають процеси планування та контролю, які є взаємодоповнюючими. Обидва вони починаються з прогнозування, яке охоплює всі ключові змінні діяльності підприємства – обсяг продажів, попит, рівень запасів, чисельність персоналу. Вважається, що починати необхідно з прогнозування показника обсягу продажів. На основі прогнозу обсягу продажів планують обсяг запасів, виробництво, чисельність персоналу, рівень накладних витрат, план потоку грошових коштів. При цьому попит на більшість товарів та послуг має сезонний характер. Тому постає задача моделювання та прогнозування певних показників з врахуванням сезонної циклічності.

Маркетинговий аналіз забезпечує цінну та своєчасну інформацію про стан ринку, ефективність просування продукції компанії, а також сприяє вибору оптимальної стратегії просування та визначенню перспективних напрямків розвитку бізнесу.

Тому прогнози відіграють важливу роль у прийнятті управлінських рішень щодо інвестування в нові ризиковані проєкти, а також у виборі ринкової стратегії та тактики. У процесі прогнозних досліджень використовуються різноманітні підходи та методи аналізу.

Таким чином, метою роботи є дослідження методів прогнозування сезонних часових рядів, розробка алгоритму, який би міг універсально застосовуватися до сезонних часових рядів різного типу.

Створений програмний продукт у вигляді веб-застосунку призначатиметься для прогнозування реальних сезонних часових рядів з порівнянням результатів за відомими методами, зокрема Уінтерса, Тейла-Вейджа та спектрального сингулярного аналізу.

2. МАРКЕТИНГОВІ ДОСЛІДЖЕННЯ ТА СЕЗОННІ КОЛИВАННЯ В ЕКОНОМІЦІ

Проведення маркетингових досліджень – найважливіша складова аналітичної функції маркетингу. Відсутність подібних досліджень загрожує вельми несприятливими наслідками для будь-якої фірми.

Головна увага у маркетингових дослідженнях приділяється ринковим аспектам: оцінка стану та тенденцій (кон'юнктури) розвитку ринку, дослідження поведінки споживачів, аналіз діяльності конкурентів, постачальників, посередників, вивчення комплексу маркетингу, що включає управління товарним асортиментом, ціноутворення та розробку стратегії цін, формування каналів збуту продукції та спрямоване застосування засобів стимулювання.

Проведення маркетингових досліджень та прийняття на основі їх результатів продуманих маркетингових рішень передбачає необхідність виділення макро- та мікросередовища маркетингу як об'єкта дослідження.

Варіантів проведення маркетингових досліджень існує безліч, все залежить від цілей і від того, як використовувати результати аналітики. Незважаючи на різноманіття видів проведення маркетингових досліджень, в їх основі лежить загальна методологія, що визначає певний порядок виконання [1]. Її концепція може бути представлена у вигляді рисунку 1.



Рисунок 1. Етапи маркетингового дослідження

Прогнозування – це найбільш складний вид діяльності у системі маркетингових досліджень. Воно є основним і завершальним етапом досліджень, головні результати якого товаровиробники закладають у основу програм своєї діяльності. Відмінності у горизонті прогностичної роботи знаходять відображення у характері планування кожної окремої фірми: стратегічне планування – на базі середньо- та довгострокових прогнозів, а поточне планування господарських операцій – на основі короткострокових прогнозів.

Бізнес-діяльність в умовах періодичної змінюваності сезонів супроводжується специфічними змінами динаміки соціально-економічних процесів. Більш-менш стійкі внутрішньорічні коливання рівнів економічних показників спостерігаються з різним ступенем інтенсивності в усіх сферах життя суспільства: виробництві, обігу, споживання. Це проявляється у внутрішньорічних спадах і підйомах випуску продукції, обсягах її реалізації, споживання сировини та енергії, коливанні фінансових результатів та інші показники.

Одним із важливих завдань статистики є дослідження змін показників у часових інтервалах. Соціально-економічні явища суспільного життя перебувають у постійному розвитку: від кількох секунд на біржі до років у макроекономіці. У статистиці процес змін і розвитку соціально-економічних явищ у часі прийнято називати динамікою.

Основна мета статистичного вивчення динаміки комерційної діяльності полягає у виявленні та вимірюванні закономірностей їх розвитку у часі. Ці закономірності можна вивчати, якщо мати дані з певного кола показників на ряд моментів часу або за ряд проміжків часу, що слідує одне за одним.

Ряд розташований у хронологічній послідовності значень статистичних показників, являє собою часовий ряд (іноді його ще називають рядом динаміки). У часовому ряді процес економічного розвитку зображується у вигляді сукупності перерв безперервного, що дозволяють детально проаналізувати особливості розвитку за допомогою характеристик, які відображають зміну параметрів економічної системи в часі. У кожному часовому ряді є два основні елементи: показники часу t ; відповідні їм рівні розвитку досліджуваного явища у [2].

Як показники часу в рядах динаміки виступають або певні дати (моменти) часу, або окремі періоди (роки, квартали, місяці, добу). Рівні рядів динаміки відображають кількісну оцінку (міру) розвитку в часі досліджуваного явища.

3. МАТЕМАТИЧНІ МЕТОДИ ПРОГНОЗУВАННЯ СОЦІАЛЬНО-ЕКОНОМІЧНИХ ПРОЦЕСІВ

Прогнозування являє собою процес оцінки майбутнього стану об'єкта або явища шляхом аналізу його минулого та поточного розвитку, а також систематизації інформації про його якісні й кількісні характеристики у перспективі. Підсумком такого процесу є прогноз – уявлення про ймовірний розвиток існуючих тенденцій у майбутньому.

У будь-якій фінансовій, комерційній активності виникає потреба у передбаченні майбутнього. Прогнозування подій у соціально-економічній сфері є невід'ємною частиною

3.1. Сезонні моделі часових рядів

Як уже згадувалося, в економіці чимало процесів мають періодичні сезонні ефекти. Відповідно, часові ряди, що їх відображають, несуть регулярні сезонні коливання. Такі ряди та їхні коливання можна подати як результат дії моделей двох основних типів: з мультиплікативними або з адитивними коефіцієнтами сезонності [3].

Моделі першого типу мають вигляд:

$$\begin{aligned}x_t &= \xi_t + \varepsilon_t \\ \xi_t &= a_{1,t} f_t,\end{aligned}$$

де $a_{1,t}$ – динаміка величини, що характеризує тенденцію розвитку процесу;

$f_t, f_{t-1}, \dots, f_{t-l+1}$ – коефіцієнти сезонності;

l – кількість фаз в повному сезонному циклі;

ε_t – неавтокорельований шум з нульовим математичним сподіванням.

Моделі другого типу записуються як

$$\begin{aligned}x_t &= \xi_t + \varepsilon_t \\ \xi_t &= a_{1,t} + g_t,\end{aligned}$$

де $a_{1,t}$ – величина, що описує тенденцію розвитку процесу;

$g_t, g_{t-1}, \dots, g_{t-l+1}$ – адитивні коефіцієнти сезонності;

l – кількість фаз в повному сезонному циклі;

ε_t – неавтокорельований шум з нульовим математичним сподіванням.

Тепер подивимося на адаптивну модель з мультиплікативною сезонністю, описану П. Р. Уінтерсом, а також адитивну модель, запропоновану Г. Тейлом і С. Вейджем [3]. Ці моделі рекомендовано застосовувати для прогнозування об'ємів сезонних продажів із використанням комп'ютерного аналізу.

Модель Уінтерса базується на аналізі ізольованих часових рядів про продажі. Єдиною використовуваною інформацією є передісторія продажів даного товару. Модель Уінтерса є моделлю експоненціального типу.

Тому доцільно в прогностичних моделях враховувати конкретний характер тенденції та сезонних коливань. Це й зробив Уінтерс за допомогою експоненційної схеми. Модель при цьому стає складнішою, проте й точність прогнозів для більшості товарів істотно зростає [4].

Тоді повна сезонна модель Уінтерса з лінійним зростанням має наступний вигляд:

$$\begin{aligned}\hat{a}_{1,t} &= \alpha_1 \frac{x_t}{\hat{f}_{t-l}} + (1 - \alpha_1)(\hat{a}_{1,t-1} + \hat{a}_{2,t-1}) & 0 < \alpha_1 < 1; \\ \hat{f}_t &= \alpha_2 \frac{x_t}{\hat{a}_{1,t}} + (1 - \alpha_2)\hat{f}_{t-1} & 0 < \alpha_2 < 1; \\ \hat{a}_{2,t} &= \alpha_3(\hat{a}_{1,t} - \hat{a}_{1,t-1}) + (1 - \alpha_3)\hat{a}_{2,t-1} & 0 < \alpha_3 < 1.\end{aligned}$$

У моделі $\hat{a}_{1,t}$ є зваженою сумою поточної оцінки $\frac{x_t}{\hat{f}_{t-l}}$, отриманої шляхом очищення від сезонних коливань фактичних даних x_t та попередньої оцінки $\hat{a}_{1,t-1}$ з врахуванням найпізнішої оцінки адитивного фактору зростання $\hat{a}_{2,t-1}$. В якості коефіцієнта сезонності f_t береться його найпізніша оцінка, зроблена для аналогічної фази циклу. Потім величина $\hat{a}_{1,t}$, отримана за першим рівнянням, використовується для визначення нової оцінки коефіцієнта сезонності за другим рівнянням. Тоді прогноз на τ кроків: $\hat{x}_\tau(t) = (\hat{a}_{1,t} + \tau\hat{a}_{2,t})\hat{f}_{t-l+\tau}$. А оптимальні значення $\alpha_1, \alpha_2, \alpha_3$ знаходяться експериментально.

Далі розглянемо адитивну модель сезонних явищ з лінійним зростанням, запропоновану Г. Тейлом і С. Вейджем:

$$\begin{aligned}x_t &= a_{1,t} + g_t + \varepsilon_t \\ a_{1,t} &= a_{1,t-1} + a_{2,t},\end{aligned}$$

де $a_{1,t}$ – величина рівня процесу після елімінування сезонних коливань;

$a_{2,t}$ – адитивний коефіцієнт зростання;

g_t – адитивний коефіцієнт сезонності;

ε_t – білий шум.

Розглянемо адаптивну процедуру оновлення значень $\hat{a}_{1,t}$. В момент часу t маємо спостереження x_t , про яке відомо, що $x_t = a_{1,t} + g_t + \varepsilon_t$.

Про шум та сезонний фактор g_t нічого не відомо. Величину ε_t замінюємо нулем, а в якості g_t візьмемо останню оцінку сезонного фактора g_{t-l} , де l – період сезонного циклу.

Величина $x_t - \hat{g}_{t-l}$ розглядається як нове «фактичне» значення $a_{1,t}$. Останньою оцінкою рівня a_1 є $\hat{a}_{1,t-1}$, але вона відповідає моменту часу $t-1$, а не t . Тому до $\hat{a}_{1,t-1}$ додається $\hat{a}_{2,t}$. Але оскільки оцінку $\hat{a}_{2,t}$ ми ще не можемо отримати, то замість неї беремо $\hat{a}_{2,t-1}$, одержану на попередньому кроці. Це призводить до наступної процедури адаптації:

$$\hat{a}_{1,t} = \alpha_1 (x_t - \hat{g}_{t-l}) + (1 - \alpha_1) (\hat{a}_{1,t-1} + \hat{a}_{2,t-1}).$$

При цьому: $\hat{a}_{2,t} = \alpha_2 (\hat{a}_{1,t} - \hat{a}_{1,t-1}) + (1 - \alpha_2) \hat{a}_{2,t-1}$,

де α_2 та $(1 - \alpha_2)$ – ваги джерел інформації. Процедура дає одержання оцінки g_t :

$$\hat{g}_t = \alpha_3 (x_t - \hat{a}_{1,t}) + (1 - \alpha_3) \hat{g}_{t-l}.$$

Параметри згладжування задовольняють наступну умову:

$$0 < \alpha_1, \alpha_2, \alpha_3 < 1.$$

Адаптивне прогнозування проводиться наступним чином. Нехай t – поточний момент часу. Отже, $\hat{a}_{1,t}$, $\hat{a}_{2,t}$, $\hat{g}_t$, $\hat{g}_{t-1}$, ... відомі. Тоді прогноз на τ кроків вперед виглядатиме як

$$\hat{x}_\tau(t) = \hat{a}_{1,t} + \tau \hat{a}_{2,t} + \hat{g}_{t-l+\tau}.$$

Для деяких продуктів, що характеризуються стабільною інтенсивністю продажу та малими сезонними коливаннями, вже проста експоненційна модель є цілком задовільною. Для певних продуктів характерні істотні сезонні зміни рівня продажів. Багато продуктів, однак, мають помітну тенденцію зростання або падіння продажів, особливо коли вони тільки виходитимуть на ринок вперше або коли з'явилися нові конкуруючі товари.

3.2. Спектральний сингулярний аналіз

Першою ідеєю, що лежить в основі методу, є створення повторюваності шляхом переходу від часового ряду до послідовності векторів, що складаються з відрізків часового ряду вибраної довжини. Другою ідеєю є аналіз отриманої багатовимірної вибірки (траєкторної матриці) за допомогою її сингулярного розкладу або, використовуючи статистичні аналоги, аналізу головних компонент. Так, виходить розкладання вихідного часового ряду (його траєкторної матриці) за базисом, породжуваного ним самим [5–8].

Базовий алгоритм складається з двох взаємопов'язаних етапів, розкладу та відновлення. Отож, починаючи з етапу розкладу.

Розглянемо часовий ряд $F = (f_0, f_1, \dots, f_{N-1})$ дійсних значень довжини N , нехай $N > 2$. Будемо припускати, що ряд F – ненульовий, тобто існує, принаймні, одне i , таке що $f_i \neq 0$. Зазвичай вважається, що $f_i = f(i\Delta)$ для деякої функції $f(t)$, де t – час, а Δ – певний часовий інтервал. Тож, числа $0, \dots, N-1$ можуть бути інтерпретовані не тільки як дискретні моменти часу, але як і певні мітки, які мають лінійно-впорядковану структуру.

Крок 1. Вкладення:

Процедура вкладення переводить вихідний часовий ряд у послідовності багатовимірних векторів.

Будуємо траєкторну матрицю – це матриця вигляду:

$$X = (x_{ij})_{i,j=1}^{L,K} = \begin{pmatrix} f_0 & f_1 & \dots & f_{K-1} \\ f_1 & f_2 & \dots & f_K \\ f_2 & f_3 & \dots & f_{K+1} \\ \dots & \dots & \dots & \dots \\ f_{L-1} & f_L & \dots & f_{N-1} \end{pmatrix}.$$

Крок 2. Сингулярний розклад:

Результатом цього кроку є сингулярний розклад (SVD) траєкторної матриці ряду.

Позначимо при E , що $\lambda_1, \lambda_2, \dots, \lambda_L$, як власні числа матриці S , взяті у порядку не спадання ($\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_L \geq 0$), та $U_1, U_2, \dots, U_L$ – ортонормовану систему власних векторів матриці S , відповідних власним числам.

Якщо при $d = \max \{i : \lambda_i > 0\}$ позначити $V_i = \frac{X^T U_i}{\sqrt{\lambda_i}}$, $i = 1, \dots, d$, то сингулярний розклад

матриці X може бути записаний як

$$X = X_1 + \dots + X_d, \quad (1)$$

де $X_i = \sqrt{\lambda_i} U_i V_i^T$.

Крок 3. Групування:

На основі розкладу (1) процедура групування поділяє всі множини індексів $(1, \dots, d)$ на m підмножин $I_1, I_2, \dots, I_m$, які попарно не перетинаються.

Нехай $I = \{i_1, \dots, i_p\}$. Тоді матриця X_I , відповідна групі I , визначається як

$$X_I = X_{i_1} + X_{i_2} + \dots + X_{i_p}.$$

Такі матриці обчислюються для $I_1, I_2, \dots, I_m$, тим самим розклад (1) може бути записано в згрупованому вигляді:

$$X = X_{I_1} + X_{I_2} + \dots + X_{I_m}. \quad (2)$$

Крок 4. Діагональне усереднення:

На останньому кроці базового алгоритму кожна матриця згрупованого розкладу (2) переводиться в новий ряд довжини N .

Нехай Y – деяка $L \times K$ матриця з елементами y_{ij} , де $1 < i < L$, $1 < j < K$.

Покладемо $L^* = \min(L, K)$, $K^* = \max(L, K)$ та $N = L + K - 1$.

Нехай $y_{ij}^* = y_{ij}$, якщо $L < K$, та $y_{ij}^* = y_{ji}$ інакше.

Діагональне усереднення переводить матрицю Y в ряд $g_0, g_1, \dots, g_{N-1}$ за формулою

$$g_k = \begin{cases} \frac{1}{k+1} \sum_{m=1}^{k+1} y_{m, k-m+2}^*, & 0 \leq k < L^*, \\ \frac{1}{L^*} \sum_{m=1}^{L^*} y_{m, k-m+2}^*, & L^* - 1 \leq k < K^*, \\ \frac{1}{N-k} \sum_{m=k-K^*+2}^{N-K^*+1} y_{m, k-m+2}^*, & K^* \leq k < N. \end{cases}$$

Застосовуючи діагональне усереднення до результуючих матриць X_k , ми отримуємо ряди $\hat{F}^{(k)} = (\hat{f}_0^{(k)}, \dots, \hat{f}_{N-1}^{(k)})$, а отже, вихідний ряд $(f_0, f_1, \dots, f_{N-1})$ розкладається у суму m рядів:

$$f_n = \sum_{k=1}^m f_n^{(k)}.$$

4. ПОРІВНЯЛЬНИЙ АНАЛІЗ МЕТОДІВ

Для подальшої роботи було обрано датасет, використовуючи репозиторій сайту Kaggle. Розглядатиметься набір з чотирьох параметрів на 36 точок даних зі сфери роздрібної торгівлі. Періодом дискретизації виступає 1 місяцем, для аналізу та прогнозування можна буде використовувати три інші виміри, а саме, «Відвідувачі» в одиницях вимірювання, «Температура» у градусах °C і «Продажі» в одиницях валюти. Нехай працюватимемо у подальшому для всіх моделей з параметром «Продажі».

Нижче подано результати роботи по моделі Уінтерса (рис. 2), Тейла-Вейджа (рис. 3) та по спектральному сингулярному аналізу (рис. 4, 5).

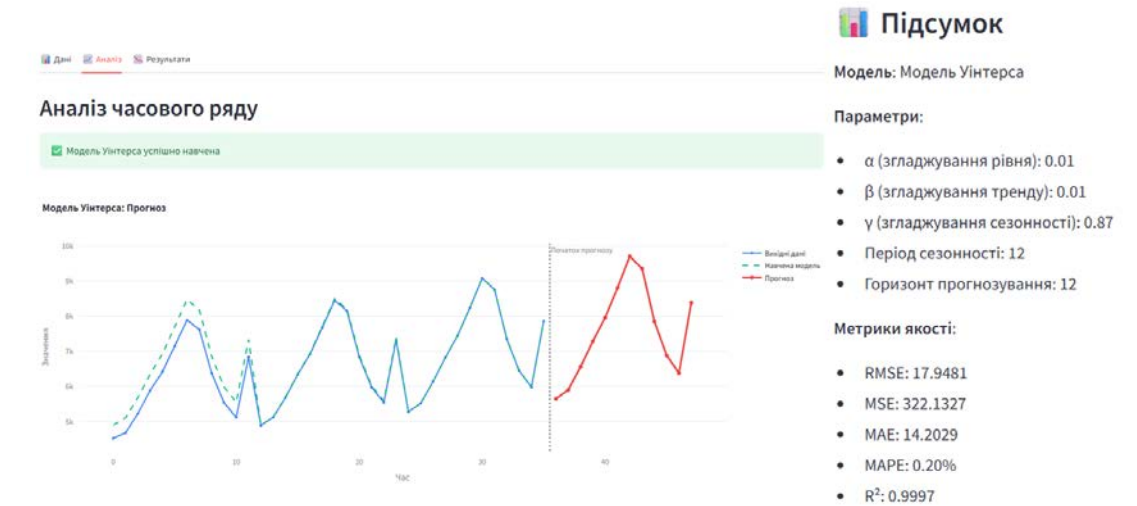


Рисунок 2. Результати роботи по моделі Уінтерса

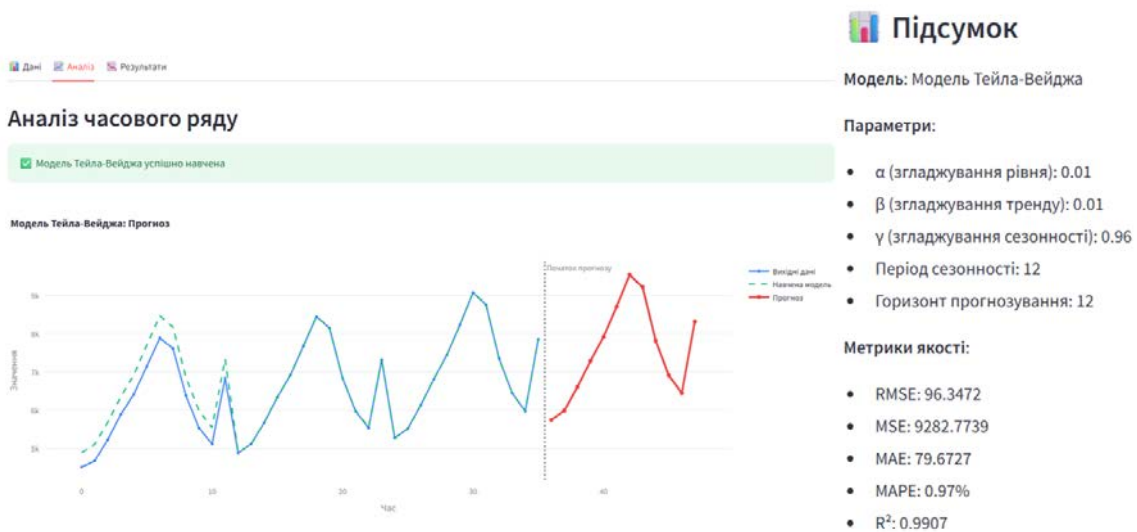
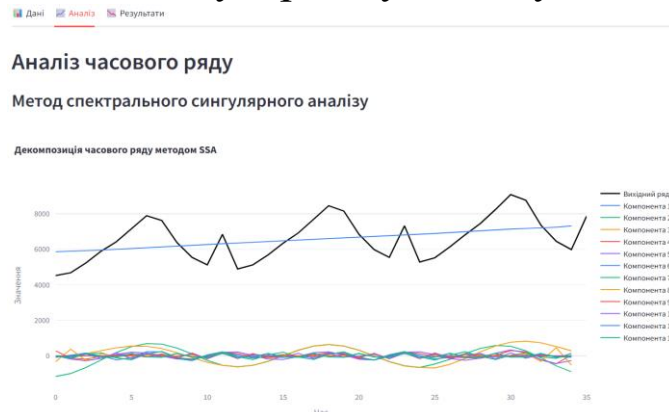


Рисунок 3. Результати роботи по моделі Тейла-Вейджа

Результати роботи по спектральному сингулярному аналізу



14

Рисунок 4. Результати роботи по спектральному аналізу



Рисунок 5. Значення сингулярного розкладу

Розглянутий процес добре піддався моделюванню, особливо для методів Уінтерса та Тейла-Вейджа. Ці моделі показали гарні результати за метриками якості при досить схожих вхідних параметрах, але кращою з них двох усе ж таки стала модель Уінтерса. При цьому для кожного окремого процесу кращими можуть виявитись різні моделі залежно від його природи та характеру. Так, можна виділити орієнтацію методу спектрального сингулярного аналізу на широкий клас процесів. Це універсальний метод дослідження сезонних часових рядів, що є актуальним у випадку, коли відсутня попередня інформація про характер процесу та немає апріорі відомої моделі, яка б добре описувала процеси зі спорідненого класу, або у випадку коли просто важко підібрати модель. Якості моделі можуть стати корисними у роботі з різноманітними потребами у сфері маркетингових досліджень.

5. ВИСНОВКИ

Отримані результати мають як теоретичну цінність для подальшого розвитку методів аналізу сезонності, так і практичне значення – для підвищення ефективності управлінських рішень у сфері економіки та бізнесу. Сезонні та інші явища у часових рядах мають значний вплив на сферу бізнесу, що підкреслює необхідність тактичного підходу в контексті маркетингових досліджень. Використання прогностичних моделей дозволяє точніше оцінити можливі наслідки на ринку та розробити ефективні стратегії реагування, що забезпечить стійкість і стабільність зі сторони маркетингу в бізнесі.

Маркетингові дослідження виступають важливим засобом для виявлення майбутніх тенденцій і прийняття обґрунтованих управлінських рішень. Розглянуті методи та моделі

забезпечують ефективний аналіз і прогнозування економічної динаміки, що є ключовим для досягнення цілей маркетингової діяльності.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Данченко О. Б., Дзюба Т. В. Маркетингові дослідження у проєктах: навч. посіб. Київ: ВНЗ «Університет економіки та права «КРОК», 2021. 224 с.
2. Бідюк П. І., Романенко В. Д., Тимошук О. Л. Аналіз часових рядів: навч. посіб. Київ: Політехніка, 2013. 600 с.
3. Клебанова Т. С. Оцінка, аналіз і попередження кризового стану підприємств житлово-комунального господарства: монографія / за ред.: Т. С. Клебанової, О. В. Димченко, О. О. Рудаченко. Харківський національний університет міського господарства ім. О. М. Бекетова. Харків: ХНУМГ ім. О. М. Бекетова, 2016. 213 с.
4. Боднар Г. Й., Гембара Т. В. Математичне моделювання структурних стохастичних динамічних процесів в освітньому середовищі ВНЗ. Львівський державний університет безпеки життєдіяльності, 2015. URL: <https://sci.ldubgd.edu.ua/bitstream/123456789/2167/1/%D0%BD%D0%BE%D0%B2%20%D0%A1%D0%A2%D0%90%D0%A2%D0%A2%D0%AF%20%D0%9A%D0%9E%D0%9D%D0%A4%201.pdf> (дата звернення: 05.09.2025).
5. Jenkins G. M., Watts D. G. Spectral Analysis and Its Applications. Original from, the University of Michigan. Holden-Day, 1969. 525 p.
6. Golyandina N., Nekrutkin V., Zhigljavsky A. Analysis of Time Series Structure: SSA and related techniques. Chapman&Hall/CRC, 2001.320 p.
7. Основи та цілі спектрального сингулярного аналізу. Academic Dictionaries and Encyclopedias. URL: <https://en-academic.com/dic.nsf/enwiki/8828599> (дата звернення: 03.09.2025).
8. Vautard R., Yiou P., Ghil M. Singular-spectrum analysis: A toolkit for short, noisy chaotic signal. Physica D. 1992. Vol. 58. P. 95-126.

МОДЕЛЬ ДЛЯ АНАЛІЗУ ВПЛИВУ ПОГОДНИХ ФАКТОРІВ НА ФОРМУВАННЯ ЦІН ЗЕРНОВИХ КУЛЬТУР В УКРАЇНІ

Тихолоз А.А.¹, Тимошук О.Л.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ tikholoz01@gmail.com

Сільське господарство України функціонує в умовах зростання кліматичних ризиків, що створює загрози для стабільності аграрного ринку. Метою роботи є розробка моделі для аналізу та кількісної оцінки впливу метеорологічних факторів на динаміку цін зернових культур. Результатом дослідження є обґрунтування використання панельного регресійного аналізу для виявлення залежностей між погодними аномаліями та ціновими коливаннями, що дозволить удосконалити прогнозування ринкового становища.

Ключові слова: ринок зернових, ціноутворення, погодні ризики, регресійний аналіз, продовольча безпека.

1. ВСТУП

Сільське господарство є стратегічним сектором економіки України, що забезпечує продовольчу безпеку та значну частину експортних надходжень. Україна посідає місце одного з провідних світових виробників і експортерів зернових культур, тому будь-які коливання врожайності мають вплив не лише на національну економіку, а й на глобальну продовольчу стабільність.

Проте сучасний агросектор стикається з серйозними викликами, серед яких, звичайно, найбільш непередбачуваним є вплив погодно-кліматичних умов. Останні десятиліття в нашій країні характеризуються зростанням кількості екстремальних явищ. Наприклад, у 2025 році аномальна посуха призвела до значних втрат посівів та зниження врожайності на 30–50 % на значних площах, що спричинило дестабілізацію ринку [1].

Довгострокові тенденції, такі як підвищення температури та дефіцит опадів, особливо у південних регіонах, призводять до зміщення агрокліматичних зон та зниження потенціалу врожайності [2]. В таких умовах традиційні підходи до прогнозування цін виявляються недостатньо ефективними. Виникає гостра потреба у моделюванні, розробці та впровадженні новітніх аналітичних інструментів, які будуть здатні враховувати комплексний вплив усіх метеорологічних факторів на формування ринкових цін.

Метою даного дослідження є створення математичної моделі, яка дозволить оцінити вплив погодних аномалій на цінову динаміку ключових зернових культур в Україні, що є важливим кроком для покращення управління ризиками в агропромисловому комплексі.

2. АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ

Формування цін на зернові культури в Україні має виражену сезонність, яка суттєво модифікується під дією агрокліматичних шоків. Емпіричні дослідження підтверджують вирішальну роль температури та опадів: зокрема, високі температури (понад +25 °C) та дефіцит вологи у критичні фази періоду зросту призводять до значного недобору врожаю [3]. Через ринкові механізми скорочення пропозиції трансформується у зростання цін, особливо у періоди посух, таких як у 2010 або 2020 роках.

У науковій літературі [4] виділяють кілька підходів до оцінки впливу клімату на аграрні ринки. Класичні агрономічні моделі (crop simulation models), такі як WOFOST або DSSAT, фокусуються на біофізичних процесах росту рослин, проте вони часто ігнорують економічні чинники формування ціни. З іншого боку, суто економетричні моделі часових рядів (ARIMA, GARCH) ефективно прогнозують цінові тренди на коротких проміжках часу, але погано враховують екзогенні шоки, спричинені аномальною погодою. Сучасні дослідження [5] свідчать про нелінійний характер впливу температури, а саме зростання температури до певного порогу (наприклад, 29°C для кукурудзи) сприяє росту, тоді як подальше підвищення призводить до різкого падіння врожайності. Саме тому в даній роботі пропонується гібридний підхід – використання панельної регресії, яка поєднує гнучкість економетрики з урахуванням фізичних кліматичних параметрів.

Для побудови адекватної моделі було здійснено збір та попередню обробку різнорідних даних, які об'єднано у єдину панельну структуру (Panel Data) з розділенням по областях України.

Джерела даних про ціни: Як залежну змінну обрано середньомісячні ціни реалізації зернових культур (пшениця, кукурудза, ячмінь). Основними джерелами є офіційні дані Державної служби статистики України [6], що забезпечують регіональний розріз, та база даних WFP (World Food Programme) для часової деталізації внутрішніх ринкових цін [7].

Джерела метеорологічних даних: Для формування незалежних змінних використано комбінацію наземних та супутникових спостережень:

- **ClimUAm** – ґридований кліматичний датасет УкрГМІ з щомісячною температурою та опадами (1946–2020) [8].
- **ERA5 (Copernicus/ECMWF)** – глобальний атмосферний реаналіз для отримання актуальних даних про температуру, опади та вологість ґрунту за останні роки [9].
- **Супутникові індекси** - дані NDVI (MODIS) та індекси посухи SPEI для оцінки стану рослинності та зволоження.

Критичним етапом підготовки даних стало просторове узгодження: накладання координатної сітки кліматичних даних на адміністративну карту областей України та агрегація показників (обчислення середніх зважених значень) для кожного регіону. Це дозволило сформувати узгоджений часовий ряд для подальшого статистичного аналізу та математичного моделювання.

3. СТРУКТУРА ВХІДНИХ ДАНИХ

Основою дослідження є інтегрований масив даних панельного типу, що охоплює період 2005 – останній доступний рік. Одиницею спостереження обрано «область–місяць», що дозволяє врахувати локальні кліматичні аномалії, які нівелюються при усередненні показників по країні.

Масив даних сформовано з трьох блоків змінних:

1. **Цінові показники (Y):** середньомісячні ціни реалізації зернових (пшениця, кукурудза, ячмінь) за даними Держстату та WFP.
2. **Кліматичні предиктори (X):** температура повітря та опади (дані ClimUAm та ERA5).
3. **Контрольні змінні:** бінарна змінна (dummy) для періоду воєнного стану (з 2022 р.), що враховує шоки пропозиції через логістичні обмеження.

Оскільки вихідні кліматичні дані отримано у форматі регулярної сітки (Gridded Data), виконано їх просторову агрегацію до меж адміністративних областей України. Також сформовано лагові змінні для врахування відкладеного впливу погоди на врожайність.

Попередня обробка та підготовка даних (Data Preprocessing). Перед етапом моделювання сирі дані проходять кілька етапів попередньої обробки для забезпечення їх якості та співставності:

1. **Дефлювання цін:** оскільки досліджуваний період (2005 – останній доступний рік) характеризується значними інфляційними процесами, номінальні ціни було приведено до реальних (у цінах базового 2021 року) за допомогою Індексу споживчих цін (ІСЦ) [10]. Це дозволяє виділити фундаментальні зміни ціни, не пов'язані з знеціненням національної валюти.
2. **Перевірка на стаціонарність:** часові ряди цін та кліматичних показників перевіряються на наявність одиничного кореня за допомогою тесту Дікі-Фуллера (ADF test). У разі нестационарності використовується перехід до перших різниць або логарифмування змінних.
3. **Обробка викидів:** аномальні значення в даних ERA5, які можуть виникати через помилки вимірювання або локальні збої, ідентифікуються за методом міжквартильного розмаху (IQR) та замінюються інтерпольованими значеннями.
4. **Синхронізація часових кроків:** кліматичні дані (щоденні або погодинні) агрегуються до місячних значень шляхом розрахунку середнього (для температури) або суми (для опадів).

У таблиці 1 наведено фрагмент підготовленого набору даних, що ілюструє різницю між кліматичними умовами у посушливий 2020 рік (критичний дефіцит вологи на півдні) та сприятливий 2021 рік.

Таблиця 1. Фрагмент панельних даних для моделювання (ілюстративний)

Region (Область)	Date (Дата)	Crop (Культура)	Price, UAH/t (Ціна)	Temp, °C	Precip, mm (Опади)	War_Dummy (Війна)	Примітка
Одеська	2020-04	Пшениця	5 950	12.4	8.2	0	Посуха
Одеська	2020-08	Пшениця	7 200	24.1	15	0	Сплеск ціни
Полтавська	2021-06	Кукурудза	6 800	20.5	65.3	0	Норма
Полтавська	2021-11	Кукурудза	5 400	4.2	42	0	Сезонність
Вінницька	2022-04	Пшениця	8 100	10.1	48	1	Вплив війни

Як видно з таблиці, модель враховує не лише прямі показники температури, але й структурні зрушення ринку, спричинені некліматичними факторами. Така деталізація дозволяє підвищити точність оцінки коефіцієнтів еластичності в економетричній моделі.

4. МЕТОДОЛОГІЯ МОДЕЛЮВАННЯ

Для аналізу впливу погодних факторів було розглянуто три класи моделей:

1. **Моделі часових рядів (ARIMA/SARIMAX):** ефективні для прогнозування трендів, але ігнорують просторову варіативність регіонів.
2. **Моделі машинного навчання (Gradient Boosting, Random Forest):** здатні виявляти нелінійні зв'язки, проте мають низьку прозорість та зрозумілість результатів («чорна скринька»).
3. **Економетричні моделі панельних даних (Panel Data Regression):** дозволяють контролювати індивідуальні ефекти регіонів (наприклад, родючість ґрунтів) та виділяти чистий вплив погодних змінних.

Враховуючи структуру даних та нелінійний характер впливу температур на вегетацію рослин та наявність часових лагів у реакції ринку, базову специфікацію моделі було розширено. Запропоноване рівняння регресії має вигляд:

$$Price_{it} = \alpha_i + \beta_1 Temp_{it} + \beta_2 Temp_{it}^2 + \beta_3 Prec_{it} + \beta_4 Prec_{it-1} + \gamma WarDummy_t + \delta_m Month_{mt} + \epsilon_{it},$$

де:

- $Price_{it}$ – реальна ціна культури в області i у місяць t ;
- α_i – фіксований ефект області (враховує ґрунтово-кліматичні особливості регіону);
- $Temp_{it}$, $Temp_{it}^2$ – середня температура та її квадрат (для перевірки гіпотези про нелінійний U-подібний вплив тепла);
- $Prec_{it}$, $Prec_{it-1}$ – рівень опадів у поточному та попередньому місяцях (лаговий ефект);
- $WarDummy_t$ – змінна контролю воєнного стану;
- $Month_{mt}$ – набір фіктивних змінних для контролю сезонності (місячні ефекти);
- ϵ_{it} – випадкова похибка.

Використання цього підходу дозволить кількісно оцінити еластичність цін відносно погодних аномалій та побудувати сценарні прогнози ринкової ситуації.

5. ВИСНОВКИ

В роботі проведено системний аналіз проблеми ціноутворення на ринку зернових культур України в умовах кліматичних змін.

Сформовано інформаційну базу дослідження. Створено інтегрований масив панельних даних за період 2005 – останній доступний рік (орієнтовно до 2023/2024 рр.), залежно від перетину цінових та кліматичних рядів. Масив поєднує економічні показники (ціни) з високоточними кліматичними даними (ERA5, супутникові індекси). Виконано просторову агрегацію даних, що дозволило узгодити ґридовані метеорологічні поля з адміністративним поділом України.

Обґрунтовано методологічний підхід. Доведено, що для оцінки впливу погодних факторів найбільш доцільним є використання регресійних моделей панельних даних з фіксованими ефектами. На відміну від класичного аналізу часових рядів, цей підхід дозволяє врахувати регіональну гетерогенність умов вирощування та відокремити кліматичний вплив від шоків, спричинених війною.

Реалізація запропонованої моделі дозволяє отримати кількісні оцінки еластичності цін відносно температурних аномалій та дефіциту опадів. Це створює підґрунтя для побудови сценаріїв розвитку ринку та розробки інструментів хеджування погодних ризиків для агровиробників і трейдерів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Україна втратила врожай через посуху [Електронний ресурс] // Landlord.ua. – 2025. – URL: <https://landlord.ua/news/roslinnistvo/ukrayina-vtratyla-vrozhaj-cherez-posuhu-ale-zbere-76-mln-tonn-zernovyh-i-zberezheeksportnyj-potenczial-vysoczkyj/>
2. Balabukh V. Yield shortfall of cereals in Ukraine caused by the change in air temperature and precipitation amount [Електронний ресурс] / V. Balabukh // Agricultural Science and Practice. – 2023. – URL: https://agrisp.com/index.php/agrisp/article/view/2023_01_04
3. The effects of weather extremes on wheat prices in Russia: The role of inputs and Russia's war in Ukraine [Electronic resource] // Food Security. – 2025. – URL: <https://link.springer.com/article/10.1007/s12571-025-01550-8>

4. Wooldridge J. M. *Econometric Analysis of Cross Section and Panel Data*. 2nd ed. Cambridge: MIT Press, 2010. 1064 p.
5. Schlenker W., Roberts M. J. Nonlinear temperature effects indicate severe damages to U.S. crop yields under climate change. *Proceedings of the National Academy of Sciences*. 2009. Vol. 106, No. 37. P. 15594–15598.
6. Середні ціни продукції сільського господарства, реалізованої підприємствами [Електронний ресурс] // Портал відкритих даних data.gov.ua. – URL: <https://data.gov.ua/dataset/776c9346-1c46-49ba-98e0-e23255760584>
7. WFP Food Prices [Electronic resource] // Humanitarian Data Exchange. – URL: https://data.humdata.org/dataset?dataseries_name=WFP+-+Food+Prices
8. Dataset of gridded time series of monthly air temperature and atmospheric precipitation for Ukraine covering the period of 1946-2020 [Електронний ресурс] // ResearchGate. – URL: <https://www.researchgate.net/publication/362894263>
9. ERA5 monthly averaged data on single levels from 1940 to present [Електронний ресурс] // Copernicus Climate Data Store. – URL: <https://cds.climate.copernicus.eu/datasets/reanalysis-era5-single-levels-monthly-means>
10. Індeksi споживчих цін за регіонами [Електронний ресурс] // Портал відкритих даних data.gov.ua. – URL: <https://data.gov.ua/dataset/776c9346-1c46-49ba-98e0-e23255760584>

ІНФОРМАЦІЙНА СИСТЕМА ДЛЯ ОЦІНЮВАННЯ КРЕДИТНИХ РИЗИКІВ НА ОСНОВІ АНСАМБЛЕВИХ МОДЕЛЕЙ

Халімончук Р.А.¹, Гуськова В.Г.

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ halimonchukrostik@gmail.com

У роботі представлено підхід до створення інформаційної системи для оцінювання кредитних ризиків позичальників із використанням методів машинного навчання та ансамблевих моделей. Система реалізує повний цикл обробки даних і дозволяє порівнювати ефективність різних алгоритмів, серед яких логістична регресія, LightGBM, MLP та комбіновані ансамблі. Додатково розроблено зручний веб-інтерфейс, що забезпечує завантаження даних, вибір моделей, перегляд метрик та графічних візуалізацій, спрощуючи проведення аналізу. Отримані експериментальні результати підтверджують, що ансамблеві методи підвищують точність прогнозування дефолтів, а запропонована система може бути корисним інструментом у практичних задачах кредитного скорингу.

Ключові слова: кредитний ризик, скоринг, машинне навчання, LightGBM, MLP, логістична регресія, стекінг, блендинг, гібридна модель.

1. ВСТУП

У сучасній фінансовій сфері задачі оцінювання кредитних ризиків потребують швидких, точних та адаптивних рішень. Традиційні статистичні методи часто не враховують складні взаємозв'язки між ознаками та не забезпечують достатньої точності при великих обсягах даних. У зв'язку з цим дедалі більшої популярності набувають інтелектуальні системи, побудовані на методах машинного навчання та ансамблевого моделювання.

Ансамблеві алгоритми дозволяють поєднати результати різних моделей та підвищити стабільність прогнозування. Це робить їх ефективним інструментом у задачах кредитного скорингу, де помилки класифікації можуть призвести до фінансових втрат або необґрунтованих відмов у кредитуванні.

У роботі створено інформаційну систему, яка забезпечує автоматизоване завантаження даних, налаштування моделей, їх навчання та формування прогнозів, що робить підхід придатним до використання у банківських установах.

2. ПОСТАНОВКА ЗАДАЧІ ТА ВИХІДНІ ДАНІ

Метою дослідження є створення інформаційної системи, що забезпечує автоматизовану оцінку ймовірності дефолту позичальників на основі ансамблевих моделей машинного навчання. Задача передбачає поєднання кількох компонентів: підготовку структурованого набору ознак, розроблення базових моделей, побудову ансамблів, отримання прогнозів та презентацію результатів у зручному для аналізу вигляді.

У роботі значну увагу приділено попередній обробці даних. Було виконано очищення вибірки від пропусків, трансформацію категоріальних ознак, масштабування числових змінних, а також обробку змінних із вкрай нерівномірним розподілом. Для уникнення зміщення моделей під час навчання застосовано стратифікований поділ на тренувальну та тестову вибірки, що дозволяє зберегти пропорцію дефолтних клієнтів.

3. МЕТОДИ ТА ПРОГРАМНА РЕАЛІЗАЦІЯ

3.1. Побудова базових моделей

На першому етапі побудовано три базові моделі. Логістична регресія використана як інтерпретований статистичний метод, що дозволяє оцінити значущість ознак та забезпечує базову точність. LightGBM, який реалізує градієнтний бустинг на деревовидних структурах, продемонстрував високу здатність моделювати складні залежності та взаємодії між ознаками. Нейронна мережа MLP забезпечила можливість аналізу нелінійних закономірностей та виявила чутливість до параметрів навчання.

3.2. Реалізація ансамблевих методів

Для підвищення точності класифікації побудовано ансамблеві моделі.

Стекінг реалізовано як дворівневу структуру, де базові моделі формують прогнозні ймовірності, які надалі виступають входними ознаками для метамоделі на основі логістичної регресії.

Блендинг забезпечує об'єднання моделей на окремій валідаційній вибірці, що дозволяє спростити процедуру тренування при збереженні високої ефективності.

Гібридна модель GBDT+LR поєднує нелінійність дерев LightGBM із інтерпретованістю лінійної моделі, що досягається шляхом перетворення листків дерев у набір бінарних ознак.

3.3. Інформаційна система та інтерфейс

Розроблена інформаційна система оснащена веб-інтерфейсом, який забезпечує послідовну роботу з даними та моделями. Користувач може завантажити власний CSV-файл або використати демонстраційний датасет, після чого вибрати цільову змінну для моделювання. Інтерфейс дозволяє гнучко обирати одну або кілька базових моделей, а також визначати ансамблевий метод.

У налаштуваннях передбачено вибір параметрів тестового поділу та випадкового генератора, що забезпечує відтворюваність результатів. Після запуску моделювання система автоматично формує таблицю метрик (AUC та Ассигасу) та відображає графічні візуалізації. Користувачу доступні стовпчикові діаграми для порівняння моделей та графік розподілу прогнозних ймовірностей, який демонструє поведінку різних алгоритмів на тестовій вибірці. Така функціональність забезпечує швидкий аналіз результатів і дозволяє ефективно порівнювати якість окремих моделей та ансамблю.

На рис. 1 наведено головне вікно інтерфейсу, яке демонструє основні елементи взаємодії з даними та моделями. На рис. 2 показано результати роботи системи: таблицю метрик для обраних алгоритмів, а також графічні візуалізації – порівняння AUC та Ассигасу та розподіл прогнозних ймовірностей для кожної моделі.

Представлені візуалізації свідчать про те, що система не лише автоматизує процес моделювання, а й забезпечує наочне порівняння результатів, що є критично важливим для прийняття рішень у задачах кредитного скорингу. Така інтеграція аналітичних інструментів робить інтерфейс практичним та зручним для подальшого застосування у фінансовій сфері.

Ensemble Model Web Interface

1. Load Data

Upload CSV file

☒ Use demo dataset if no file is uploaded

2. Select Target Column

Target column

If empty, "target" or last column will be used.

3. Select Models & Ensemble Method

Base models (choose one or more)
☒ Logistic Regression
☒ Random Forest
☐ Gradient Boosting
☐ MLP (Neural Net)
☒ LightGBM
You can tick any combination of models.

Ensemble method (choose one)

4. Settings

Test size

Random state

Model Metrics

Рисунок 1. Інтерфейс користувача

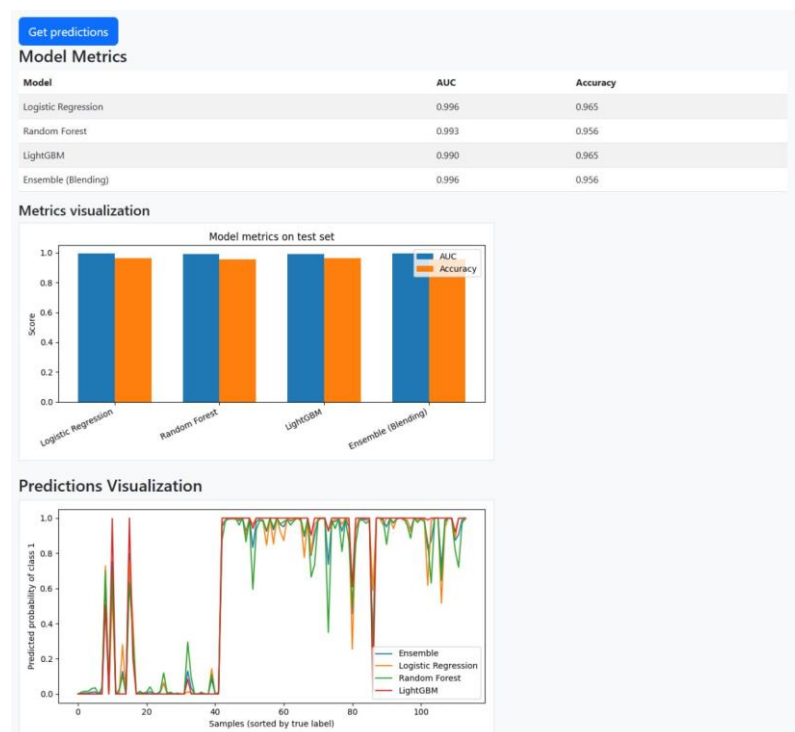


Рисунок 2. Таблиця метрик і графічна візуалізація в інтерфейсі

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

4.1. Логістична модель

Оцінювання якості моделей проводилося на тестовій вибірці, яка була повністю відокремлена від даних, використаних у процесі навчання. Для зменшення впливу випадкових

коливань у розподілі даних додатково застосовано перехресну валідацію. Такий підхід забезпечив отримання усереднених результатів, що коректно відображають загальну надійність і стабільність моделей.

Логістична регресія виступила базовою моделлю для порівняння з іншими алгоритмами (табл. 1). Її лінійний характер забезпечив інтерпретованість та стабільну класифікацію більшості «надійних» клієнтів. Водночас модель показала обмежену здатність розпізнавати складні нелінійні взаємозв'язки, що призвело до пропуску частини дефолтних випадків. Це типова особливість лінійних методів, які не враховують взаємодії ознак і тому слугують радше початковим орієнтиром для оцінювання інших підходів.

Таблиця 1. Метрики логістичної регресії

	Precision	Recall	F1-score	Support
Class 0	0.92	1.00	0.96	70687
Class 1	0.32	0.00	0.00	6191
Accuracy	-	-	0.92	76878
Macro avg	0.62	0.50	0.48	76878
Weighted avg	0.87	0.92	0.88	76878

4.2. LightGBM

Модель LightGBM продемонструвала значно вищу точність завдяки можливості моделювати складні структури та взаємодії між ознаками. Завдяки деревоподібній архітектурі й механізму градієнтного бустингу вона ефективно розпізнавала ризикових позичальників, зменшила кількість хибних негативних прогнозів і показала стабільність навіть за нерівномірних або зашумлених даних. LightGBM став одним із найрезультативніших методів серед протестованих (табл. 2).

Таблиця 2. Метрики LightGBM

	Precision	Recall	F1-score	Support
Class 0	0.92	1.00	0.96	70687
Class 1	0.57	0.02	0.04	6191
Accuracy	-	-	0.92	76878
Macro avg	0.74	0.51	0.50	76878
Weighted avg	0.89	0.92	0.88	76878

4.3. Нейронна мережа (MLP)

Багатошаровий перцептрон добре впорався з нелінійними зв'язками в даних та продемонстрував високу здатність розпізнавати дефолтних клієнтів. Однак модель виявилась чутливою до параметрів навчання, що в окремих випадках призводило до збільшення хибнопозитивних рішень. У загальному результаті MLP показав сильні сторони у задачах із складною структурою ознак, але вимагає точного налаштування гіперпараметрів (табл. 3).

Таблиця 3. Метрики нейронної мережі

	Precision	Recall	F1-score	Support
Class 0	0.93	0.94	0.93	70687
Class 1	0.17	0.15	0.16	6191
Accuracy	-	-	0.87	76878
Macro avg	0.55	0.54	0.55	76878
Weighted avg	0.87	0.87	0.87	76878

4.4. Стекінг

Стекінг, як дворівневий ансамблевий підхід, став однією з найефективніших моделей дослідження (табл. 4). Метамодель об'єднала прогнози логістичної регресії, LightGBM та MLP, навчившись використовувати їх сильні сторони. Це дозволило суттєво знизити кількість помилок класифікації та досягти найвищих значень AUC серед усіх алгоритмів. Стекінг продемонстрував найкраще узагальнення на тестовій вибірці.

Таблиця 4. Метрики стекінгу

	Precision	Recall	F1-score	Support
Class 0	0.92	0.99	0.96	70687
Class 1	0.48	0.06	0.10	6191
Accuracy	-	-	0.92	76878
Macro avg	0.70	0.53	0.53	76878
Weighted avg	0.89	0.92	0.89	76878

4.5. Блендинг

Блендинг показав результати, близькі до стекінгу, використовуючи простішу схему об'єднання прогнозів на окремій валідаційній частині даних (табл. 5). Метод виявився ефективним і менш вимогливим до ресурсів, забезпечивши хороший баланс між точністю та обчислювальною простотою. Він успішно усунув частину недоліків окремих моделей і показав стабільну роботу.

Таблиця 5. Метрики блендингу

	Precision	Recall	F1-score	Support
Class 0	0.93	0.94	0.93	70687
Class 1	0.17	0.15	0.16	6191
Accuracy	-	-	0.87	76878
Macro avg	0.55	0.54	0.55	76878
Weighted avg	0.87	0.87	0.87	76878

4.6. Гібридна модель (GBDT + LR)

Гібридна модель поєднала деревоподібний підхід LightGBM із лінійною логістичною регресією. Використання leaf-фіч дозволило побудувати новий простір ознак, що відображає структуру рішень дерев. Це забезпечило чіткіше розмежування класів та високу точність при збереженні інтерпретованості моделі. За результатами дослідження, гібридний підхід став одним із найстабільніших і найзбалансованіших методів (табл. 6).

Таблиця 6. Метрики гібридної моделі

	Precision	Recall	F1-score	Support
Class 0	0.92	0.99	0.95	70687
Class 1	0.34	0.07	0.12	6191
Accuracy	-	-	0.91	76878
Macro avg	0.63	0.53	0.54	76878
Weighted avg	0.88	0.91	0.89	76878

5. ВИСНОВКИ

У роботі створено інформаційну систему для оцінювання кредитних ризиків позичальників із використанням сучасних методів машинного навчання. Проведені дослідження показали, що окремі моделі забезпечують різний рівень точності: логістична регресія слугує інтерпретованою базою, LightGBM демонструє високу якість за рахунок деревоподібної структури, а нейронна мережа добре працює з нелінійними залежностями.

Найкращих результатів вдалося досягти за допомогою ансамблевих методів. Стекінг та гібридна модель GBDT+LR показали найвищі значення AUC і найменшу кількість помилок класифікації, що підтверджує ефективність комбінування різних алгоритмів. Блендинг також продемонстрував стабільні результати при нижчій складності реалізації.

Розроблений інтерфейс системи забезпечив зручну взаємодію з моделями, можливість налаштовувати параметри та аналізувати результати без необхідності втручання в програмний код.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Géron, A. Hands-On Machine Learning with Scikit-Learn, Keras and TensorFlow. / A. Géron; O'Reilly Media, USA, 2023. 850 p.
2. Brownlee, J. Ensemble Learning Algorithms with Python. / J. Brownlee; Machine Learning Mastery, Australia, 2022. 290 p.
3. Zhang, C.; Ma, Y. Ensemble Machine Learning: Methods and Applications. / C. Zhang, Y. Ma; Springer, London, 2012. 360 p.
4. Friedman, J.; Hastie, T.; Tibshirani, R. The Elements of Statistical Learning. / Springer, New York, 2017. 745 p.
5. Chen, T.; Guestrin, C. XGBoost: A Scalable Tree Boosting System. / Proceedings of the 22nd ACM SIGKDD Conference; San Francisco, USA, 2016. 13 p.
6. Ke, G.; Meng, Q.; Finley, T. LightGBM: A Highly Efficient Gradient Boosting Decision Tree. / NIPS Proceedings; California, USA, 2017. 15 p.
7. Pedregosa, F. et al. Scikit-learn: Machine Learning in Python. / Journal of Machine Learning Research, Vol. 12, 2011. p. 2825-2830.

8. Thomas, L. C.; Edelman, D. B.; Crook, J. N. Credit Scoring and Its Applications. / SIAM, Philadelphia, USA, 2017. 384 p.
9. Abdou, H. A.; Pointon, J. Credit Scoring, Statistical Techniques and Evaluation Criteria: A Review. / Intelligent Systems in Accounting, Finance and Management; Vol. 18, 2011. p. 59–88.
10. Lessmann, S.; Baesens, B.; Seow, H. V.; Thomas, L. C. Benchmarking State-of-the-Art Classification Algorithms for Credit Scoring. / European Journal of Operational Research, Vol. 247, 2015. p. 124–136.
11. Marqués, A. I.; García, V.; Sánchez, J. S. Exploring the Behaviour of Base Classifiers in Credit Scoring Ensembles. / Expert Systems with Applications, Vol. 39, 2012. p. 10244–10250.
12. Jadhav, S.; Channe, H. Comparative Study of KNN, Naive Bayes and Decision Tree Classification Techniques. / International Journal of Science and Research, Vol. 5 (1), 2016. p. 1842–1845.

INTELLIGENT DECISION SUPPORT SYSTEMS FOR FINANCIAL RISK MODELING AND ANALYSIS

Tymoshchuk O.L.¹, Bidyuk P.I.², Levenchuk L.B.³, Guskova V.G.⁴

National Technical University of Ukraine
“Igor Sikorsky Kyiv Polytechnic Institute”, Kyiv, Ukraine

¹ oxana.tim@gmail.com [0000-0003-1863-3095], ² pbidyuke_00@ukr.net [0000-0002-7421-3565], ³ lusi.levenchuk@gmail.com [0000-0002-8600-0890], ⁴ guskovavera2009@gmail.com [0000-0001-7637-201X]

The article examines approaches to financial risk analysis based on mathematical models and intelligent decision support systems (IDSS). The study describes key types of financial risks and the properties of economic processes that complicate modeling, such as non-stationarity, nonlinearity, and heteroscedasticity. It outlines the main stages of model construction – from data preprocessing and structure identification to parameter estimation and adequacy assessment. Classical risk management models, including BIA, LDA, and IMA, are analyzed, and the potential of Bayesian networks and fuzzy logic methods for handling uncertainties is highlighted. The article demonstrates that specialized IDSS, designed according to the principles of systems analysis, improve the quality of risk assessment and provide users with recommendations on model selection, uncertainty mitigation, and evaluation of forecasting performance. It is emphasized that the use of such systems reduces the labor intensity of analysis and increases the accuracy of forecasting potential financial losses.

Keywords: financial risks; risk analysis; non-stationary processes; nonlinear models; heteroscedasticity; Bayesian networks; fuzzy logic; Loss Distribution Approach; Basic Indicator Approach; Internal Measurements Approach; risk modeling; forecasting; decision support system (IDSS).

1. INTRODUCTION

Analysis of financial risks is in the focus of attention for many researchers and practitioners worldwide, as effective risk management plays a crucial role in ensuring the stability and resilience of financial institutions and economic systems. Among the most frequently encountered categories of financial risks are market, credit, and operational risks, each of which may significantly influence the performance and sustainability of financial activities. The analysis of such risks typically involves several fundamental components, including the identification of the underlying financial processes, appropriate data preprocessing, construction and validation of mathematical models, and estimation of potential losses together with their associated probabilities. Equally important is the task of forecasting future losses, which allows decision-makers to anticipate adverse events and implement preventive or mitigating strategies.

In modern financial environments characterized by increasing volatility, structural changes, and expanding volumes of heterogeneous data, the development of reliable methods for risk analysis has become increasingly complex. This complexity necessitates the application of advanced modeling techniques capable of capturing non-linear relationships, dynamic variance, non-stationary behavior,

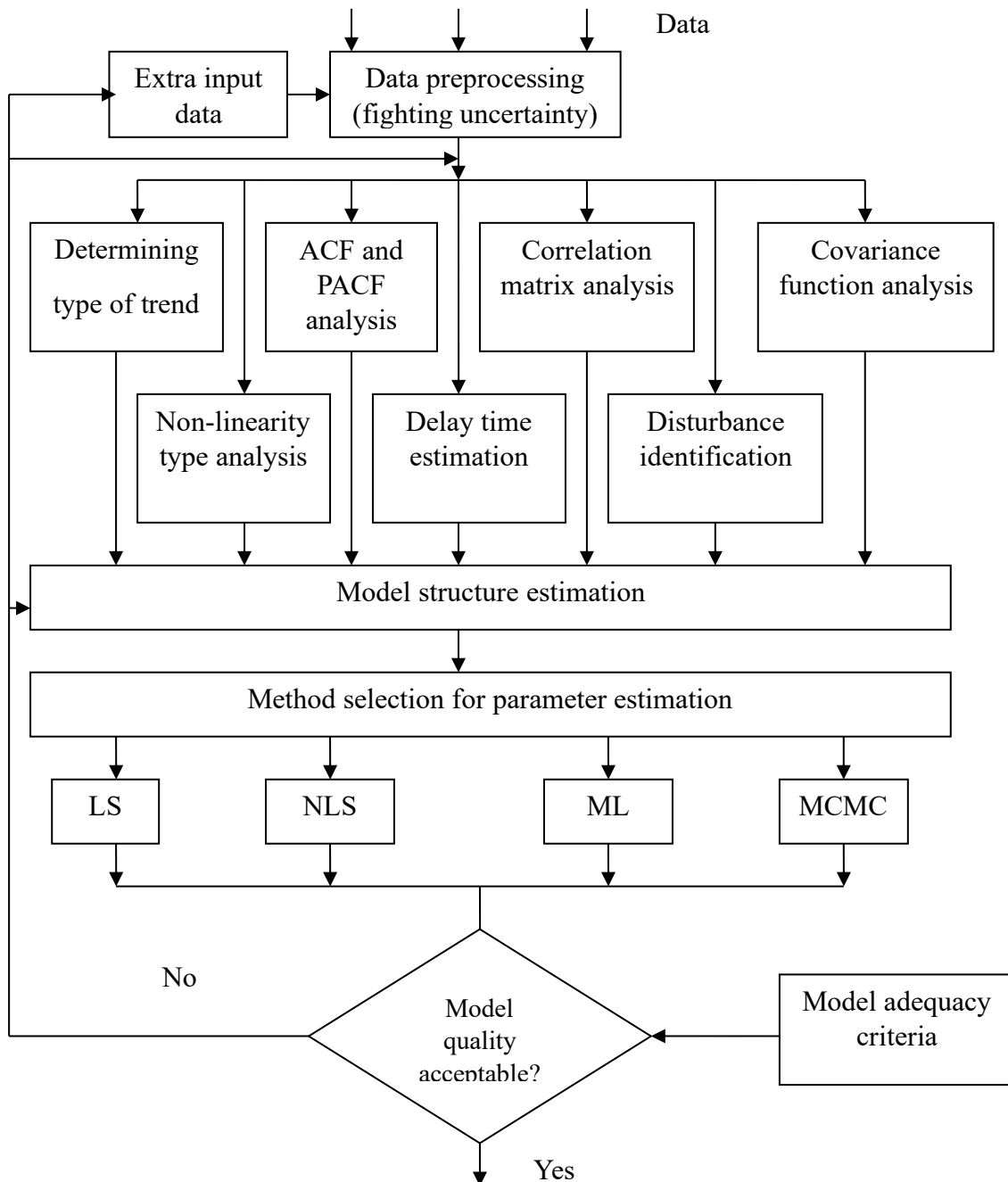


Figure 2. Approximate functional chart of intellectual decision support system

After parameter estimation model adequacy is analyzed using appropriate statistical parameters (statistics). Besides application of these parameters each model constructed can be tested on the quality of forecasting immediately after constructing. The forecasting capability of the models is important from the point of view of predicting possible future loss. Especially important is possibility for high quality forecasting of dynamic variance with heteroscedastic process models because the variance is basic parameter for estimating possible loss in many cases of financial risk estimation, for example market risk.

4. INTELLIGENT DECISION SUPPORT SYSTEMS (IDSS)

Of particular importance for applications in financial risk analysis are probabilistic models like Bayesian networks, decision trees, various distributions etc. The models provide the possibility for taking into consideration probabilistic and some other types of uncertainties that accompany financial risk management. Especially flexible are Bayesian networks that exhibit many advantages of modeling complex multivariable processes [4]. The models dimensionality can be very high, they can include discrete and continuous variables, and they are compatible for constructing combined models [5]. Besides, there exist multiple methods for the model structure and parameter (conditional probabilities) estimation, and for computing final result – marginal probabilities.

To analyze financial process and estimate financial risks it is recommendable to construct specialized intellectual decision support systems (IDSS) that already found wide applications in many applied areas [6]. IDSS designed on the principles of system analysis is very convenient instrument for researchers, engineers and students; such systems provide substantial help in modeling, forecasting and decision making.

Among substantial advantages of IDSS are following:

- correctly organized functioning of IDSS provides a user with recommendations regarding selection of data preprocessing methods, especially the methods for fighting statistical uncertainties often available in collected data (external disturbances, missing measurements etc.);
- statistical uncertainties of data are practically always available, and the means to fight them can be incorporated into IDSS itself what makes easier usage and modification of the methods if necessary;
- next the recommendations can be generated regarding possible types of mathematical models to be constructed in each specific case and algorithms for their constructing; this is especially convenient for the users who still do not have enough experience of analyzing non-stationary and non-linear processes;
- the IDSS is usually directed towards solution of specific problems – for example, constructing models for risk estimation can include necessary (available and designed for specific tasks) statistical quality criteria providing the possibility for thorough analysis of model adequacy and forecast quality;
- the system provides a user with necessary recommendations in cases when a model constructed does not exhibit necessary adequacy and quality of forecasting is not satisfactory; these recommendations indicate what other model type can be constructed or what additional data processing like fighting uncertainties can be performed;
- besides, specialized IDSS may include all necessary additional information (like special data and model tests, indications to processes simulation; constants, etc.).

The experience of practical application of specialized IDSS shows that quality of risk analysis is substantially higher with constructing the specialized system. Also necessary time for carrying out the analysis is usually shorter.

5. CLASSIFICATION OF FINANCIAL RISK MODELS

Consider some types of mathematical models that are used in financial risk analysis. The models can be classified by the use of the distinguishing features (criteria) considered below. Generally financial risks can be classified from the point of view of input factors, and from the point of view of consequences of the events linked to risks. Mathematical models used for describing risks can also be divided into two classes: models related to analysis of consequences (or on *up down analysis*), and models based upon analysis of risk factors (or on *bottom up analysis*).

On the other side, as distinguishing feature of a model can be considered the problems that could be solved with the model. For example, the models directed on function construction for mean

distribution of loss, can be used for estimating mean and maximum loss at given level of significance (Table 1).

Table 1. Model classification for risk management

Class of models	Data distribution identification and risk classification into critical groups	Model application for mean loss estimation	Possibility of model application for estimation of maximum loss	Possibility of model application for risk factor identification
Models related to analysis of consequences (<i>Top Down</i>)				
BIA			+	
IMA			+	
LDA	+	+	+	
Models based upon analysis of risk factors (<i>Bottom Up</i>).				
Sb-AMA	+	+	+	
Method of functional correlations	+	+	+	
Regression analysis	+	+	+	+
Bayesian networks	+	+	+	+
Fuzzy logic	+			+

The models based upon analysis of consequences are mostly constructed using the techniques of mathematical statistics. This group of methods uses collecting and analysis of data related to loss during previous periods of time with subsequent extrapolation of values on future periods. The most widely used are the following models:

BIA – Basic Indicator Approach;

LDA – Loss Distribution Approach;

IMA (Internal Measurements Approach).

The BIA model describes requirements regarding necessary volume of capital for covering loss due to risks. It was proposed for solving the problem of estimation maximum possible loss at given level of significance (99%). In application to analysis of operational risk the model is based upon initial idea that operational loss is estimated as residual of general loss relevant to other types of risks. The quintile of interest of random value x , that characterizes volume of loss is determined via the formula:

$$\hat{Q}_{99\%} = \overline{GI} \cdot \alpha, \quad (1)$$

where $\hat{Q}_{99\%}$ is an estimate of quintile at the level 99% of the random value that characterizes the loss (actually this requirement to the capital necessary to cover possible loss); $\overline{GI}$ is mean net income for three previous years that showed positive net annual income; $\alpha = 15\%$ is coefficient, determined by the Basel Committee on the bases of bank statistics.

Practical application of the model is based upon general estimate of risk (possible loss) that is principally possible. The general estimate of possible loss is based upon reports about loss in former time periods with correction on business scale. The coefficient α should be corrected to cover the risk of fraud. The drawback of the model is that estimate of fraud is dependent on business volume, and does not depend on the quality of monitoring procedures that should identify the cases of fraud.

The models based upon LDA (Loss Distribution Approach). The model constructing is based upon assumption that random value x , that characterizes the loss within time period t , can be defined as follows:

$$x = \sum_{i=1}^{n(t)} L_i, \quad (2)$$

where $n(t)$ is a random number that characterizes number of cases when loss was observed of the same type within time period t ; L_i is a set of random values that characterize separate loss for each observation. It is supposed here that the values of, L_i , are independent and equally distributed for selected types of loss.

To construct the model an analysis is performed for selected period of time for each pair “business line”/“type of loss”. Using the data regarding frequency of loss and volume of loss (observed in former periods) a mean frequency of loss is calculated (risk events), $(E(n(t)))$, and mean value of loss for each risk event, $(E(L))$. The next task is estimation of distribution type for the random values, $n(t)$ and L . In practice, the standard distributions are used, for example, Poisson or negative binomial, and empirical ones using historical data. Parameters of standard distributions are estimated with statistical data.

The distribution functions constructed for random values, $n(t)$ and L , allow for constructing distribution function for x . To fulfill the task as a rule Monte-Carlo method is applied using expression (2) that characterizes general loss for specific risk type. Next the point estimate of loss is possible, and/or quintile of given level of significance that allows (for example) estimation of OpVaR (Operational Value-at-Risk). Graphical representation of the LDA method is given in Fig. 3.

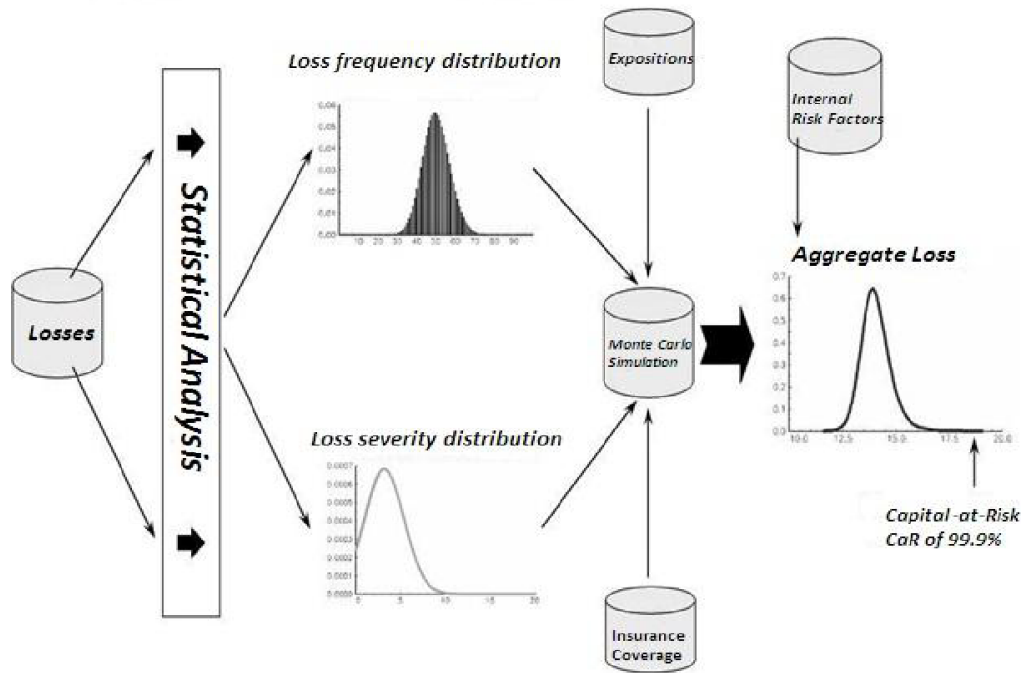


Figure 3. Graphical representation of LDA method

There exist various approaches to constructing model according to the LDA method; they can be divided into two groups.

The first group of models is based upon analysis risk events and respective loss without taking into consideration risk factors and cause and effect links. These models include estimation of distribution characteristics for risky events frequency and usage of the extreme value theory (EVT).

The second group of models takes into consideration wider set of parameters; besides risk

events risk factors are considered. Thus, in such models are used as parameters of distribution in the form of risk factors frequencies; such approach provides the possibility for taking into consideration interaction between risks.

The LDA class model can be hired for analysis of risks related to fraud (RF). In such cases it is necessary for each type of RF risk estimate distribution of its frequency, and distribution of loss using historical data. However, a distinguished feature of the fraud risk is that model is usually directed towards identification only a part of attempts to perform fraud. Another part of cases of general set available are not selected and analyzed. It means that empirical data should be supplied with other models to complete the picture.

The models built on the bases of Internal Measurements Approach (IMA).

This method provides a possibility for estimating maximum possible loss without constructing distribution for random value x , characterizing volume of loss. The approach to modeling is based upon assumption that all loss can be divided into expected (that is close to mathematical expectation of the loss within selected period of time), and non-predictable (that exceed mean value and belong to the “tail” of distribution. In simple case such model is represented by the linear equation:

$$\hat{Q}_{99\%} = EI \cdot PE \cdot LGE = \gamma \cdot EL,$$

where $\hat{Q}_{99\%}$ is estimate of 99-th quintile of the distribution for possible loss (volume of capital necessary for covering risk) related to specific type of events; PE is probability of negative event within the time period considered; EI is scale coefficient; LGE is mean value of specific loss under condition that negative event has occurred; γ is coefficient that allows to estimate requirements to capital via the estimate of expected loss EL . An integrated (total) estimate of maximum possible loss is determined as the sum of all risk estimates determined for possible risk events and business lines. The distinguished feature of this method comparing to LDA is that no assumption is made regarding distributions of $n(t)$ and L , that are not considered here. Instead an assumption is made regarding the volume of expected and unexpected loss. The use of the assumption requires application of additional coefficients (γ), and results in degrading the estimates to be found. On the other hand, IMA provides the possibility for estimating maximum loss without solving the problem of constructing distribution for X , characterizing total possible loss.

6. DISCUSSION

The short description of IDSS given above provides indications how to create practically operating system, and what functions it should perform to be effective assistant for a user. The system should be constructed on the use of system analysis principles – this is highly effective way for creating the best possible working instrument. Certainly, there are many useful ideas regarding the system functionality, and convenience to work with it, but the approach to design and implementation of it should be realistic.

7. CONCLUSIONS

The article demonstrated that specialized IDSS, designed according to the principles of systems analysis, helps to improve the quality of financial risk assessment and provide users with recommendations on model selection, uncertainty mitigation, and evaluation of loss forecasting performance. It is emphasized that the use of such system reduces the time and labor intensity for risk analysis and increases accuracy of forecasting potential financial loss. Thanks to modular construction of the system there exists a possibility for its modification: adding extra functions for risk analysis, criteria for testing intermediate and final results of computations, improving interaction with system etc.

REFERENCES

1. Bidyuk P. I., Tymoshchuk O. L., Huskova V.G., Prosiankina-Zharova T.I., Levenchuk L.B. Analysis of Non-stationary and Non-linear Process in Economy and Finances. – Kyiv: NTUU “Igor Sikorsky KPI”, 2025. – 343 p.
2. Bidyuk P. I., Romanenko V. D., Tymoshchuk O. L. Time Series Analysis. – Kyiv: Polytechnika, 2013. – 600 p.
3. Tsay R. S. Analysis of Financial Time Series. – New York: J. Wiley & Sons, Inc., 2010. – 715 p.
4. Zgurovsky M. Z., Bidyuk P. I., Terentiev O. M., Prosiankina-Zharova T. I. Bayesian Networks in Decision Support Systems. – Kyiv: «Edelweis», 2015. – 300 p.
5. Jorion Ph. Financial risk-management: Second edition. New Jersey: John Wiley & Sons, 2003. – 708 p.
6. Gupta J. N. D., Forgionne G. A., Mora M. Intelligent Decision-making Support Systems. – London: Springer-Verlag, 2006. – 522 p.

РЕКОМЕНДАЦІЙНІ СИСТЕМИ З ВИКОРИСТАННЯМ ТЕКСТІВ ОГЛЯДІВ КОРИСТУВАЧІВ ТА АНСАМБЛЮВАННЯ НА ОСНОВІ СТЕКІНГУ

Артеменко Є.В.¹, Недашківська Н.І.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ artemenko.yevhenii@lil.kpi.ua, ² nedashkovskaya.nadezhda@lil.kpi.ua [0000-0002-8277-3095]

У роботі досліджено гібридну рекомендаційну систему, що поєднує матричну факторизацію, семантичні векторні моделі та стекінг для прогнозування рейтингів товарів Amazon. Стекінгові моделі продемонстрували здатність ефективно інтегрувати текстові ознаки, подібності та прогнози окремих алгоритмів, стабільно покращуючи MAE порівняно з окремими компонентами.

Ключові слова: рекомендаційна система, матрична факторизація, текстові вкладки, сентимент-аналіз, стекінг.

1. ВСТУП

У сучасних електронних платформах обсяг доступної інформації стрімко зростає, що ускладнює вибір товарів та послуг для користувачів. Тому рекомендаційні системи стали ключовим інструментом персоналізації, підвищуючи зручність пошуку та покращуючи користувацький досвід. Традиційні моделі рекомендацій ґрунтуються переважно на колаборативній фільтрації або факторизації матриць, але такі підходи часто страждають від проблеми розрідженості даних та холодного старту [1]. Задля вирішення цих проблем, дослідники активно інтегрують у моделі тексти оглядів, які містять багатий семантичний опис переваг, досвіду та контексту використання товарів.

Використання текстових оглядів користувачів дозволяє покращити точність рекомендацій завдяки їх емоційним, тематичним і латентним ознакам, які не представлені у звичайних числових рейтингах [2]. Крім того, сучасні методи обробки природної мови (трансформери та контекстуальні векторні подання) значно розширили можливості аналізу текстів у задачах рекомендації [3]. Комбінування текстової інформації з класичними моделями є одним із найефективніших напрямів розвитку рекомендаційних систем [4].

2. МАТЕРІАЛИ ТА МЕТОДИ

2.1. Колаборативна фільтрація та матрична факторизація

SVD++ – матрична факторизація, яка враховує як явні, так і неявні переваги користувача. На відміну від класичного SVD, ця модель використовує не лише оцінені товари, але й ті, з якими користувач взаємодівав (перегляди, кліки, історію). Це дозволяє точніше моделювати латентні фактори та пом'якшує проблему розрідженості даних і холодного старту.

CoClustering (спільне кластеризаційне групування) одночасно кластеризує користувачів і товари, формуючи комірки з подібними шаблонами оцінювання. Замість моделювання кожного користувача окремо, система працює на рівні кластерів, що зменшує шум у даних. Метод добре працює на структурованих оцінках, але може бути менш точним у випадку дуже різномірних груп користувачів.

SlopeOne – проста й швидка модель прогнозування, яка базується на різниці між середніми оцінками пар товарів. Передбачається, що якщо товар А зазвичай отримує оцінку на 1 бал більше за товар В, то для нового користувача ця залежність також буде збережена. Модель майже не має параметрів і добре працює на великих наборах рейтингів, але зазвичай поступається сучасним нелінійним методам точністю.

KNNBaseline поєднує метод найближчих сусідів із базовою моделлю, яка компенсує систематичні зміщення користувачів та товарів. Спочатку обчислюються базові упередження (наприклад, користувач ставить у середньому вищі оцінки), після чого виконується пошук схожих товарів або користувачів за обраною метрикою. Це дає більш стабільні результати порівняно зі звичайним KNN та покращує роботу на розріджених даних.

2.2. Контентна фільтрація та векторизація текстів

TF-IDF – це класичний підхід до векторизації тексту, що оцінює інформативність слова у межах документа та всього корпусу. Вага зростає пропорційно частоті появи слова в конкретному тексті та зменшується, якщо це слово часто трапляється в інших документах. Така модель дозволяє ефективно порівнювати схожість між відгуками та описами товарів, ґрунтуючись на ключових термінах.

Word2Vec – нейронний метод побудови векторів слів, який відображає семантичні зв'язки між ними. У процесі навчання модель формує простір, у якому близькі за змістом слова розташовуються поруч. Вектор документа отримується шляхом усереднення векторів слів, що дає можливість зіставляти текстові описи та відгуки на основі їхньої латентної семантики.

Sentence-BERT – це модель створення векторних представлень речень, оптимізована для задач обчислення семантичної близькості. На відміну від традиційних підходів, вона формує щільні контекстуальні вектори, що відображають загальний зміст тексту, а не лише окремі слова. Застосування цього методу забезпечує значно точніше порівняння складних текстів, таких як огляди користувачів.

Подібність текстів визначається за допомогою косинусної подібності.

2.3. Сентимент-аналіз

Сентимент-аналіз відіграє важливу роль у підсиленні рекомендаційних моделей, оскільки дозволяє врахувати емоційне забарвлення користувацьких відгуків. Навіть за наявності числових рейтингів текст містить додаткову інформацію про ступінь задоволеності, контекст використання та ключові проблеми товару, які не завжди відображаються у фінальній оцінці. Тому інтеграція тональності тексту додає моделі цінні поведінкові сигнали.

У роботі використано два підходи до аналізу тональності:

1. **VADER** – лексично-орієнтований інструмент, спеціально оптимізований для коротких англійських текстів, таких як огляди або коментарі. Він визначає частки негативної, позитивної, нейтральної тональності та загальний компаунд-показник. Перевага VADER полягає в здатності коректно обробляти сленг, підсилювачі, емодзі та інші маркери емоційності.
2. **TextBlob** – інструмент, що оцінює суб'єктивність і полярність тексту. Суб'єктивність відображає, наскільки текст містить суб'єктивні судження, тоді як полярність показує напрям емоційної оцінки. Це дозволяє моделі розрізняти нейтральні інформаційні огляди та емоційно насичені.

2.4. Методи стекінгу

Стекінг (stacking) – метод ансамблювання, що поєднує результати різних моделей у єдиному мета-регресорі. У роботі використано:

1. **RandomForestRegressor**.

2. XGBoost.

3. Linear Regression.

Комбінування прогнозів моделей CF та текстових подібностей дозволяє отримати точніші оцінки ніж використання окремих алгоритмів.

3. ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ АЛГОРИТМІВ І ПОРІВНЯННЯ РЕЗУЛЬТАТІВ

У процесі реалізації продукту були використані набори даних Amazon Reviews та Amazon Meta у сфері Office Products та Automovie. Мета дослідження полягала у зменшенні похибки у задачі rating prediction.

Для обох наборів даних, для покращення точності прогнозування рейтингу до моделей факторизації матриць було інтегровано додаткові текстові та sentimentні ознаки. Перш за все було розраховано косинусну схожість між текстом відгуку та описом товару за трьома підходами: Word2Vec, TF-IDF та Sentence-BERT. Ці показники дають змогу оцінити, наскільки зміст відгуку узгоджується з інформацією, поданою у характеристиці товару, та дають семантичні сигнали.

Другим блоком було додано sentimentні характеристики, отримані за допомогою VADER і TextBlob. Вони відображають пропорцію негативної, позитивної чи нейтральної тональності, загальний склад емоційності, а також ступінь суб'єктивності висловлювання. Ці ознаки дозволяють моделі урахувати емоційний контекст, який не завжди повністю проявляється у числовому рейтингу, але суттєво впливає на сприйняття товару користувачем.

Крім того, у модель було включено прогнози базових алгоритмів факторизації матриць, таких як SVDpp, CoClustering, SlopeOne і KNNBaseline. Використання цих прогнозів дає змогу поєднати переваги різних підходів у межах стекингу та сформувати більш насичений простір ознак. У сукупності ці текстові, емоційні та модельні сигнали створюють розширений набір ознак, який покращує навчання метамоделей та дозволяє ефективніше описувати поведінку користувачів.

На рис. 1 та рис. 2 відображено кореляцію нових ознак з рейтингом для обох наборів.

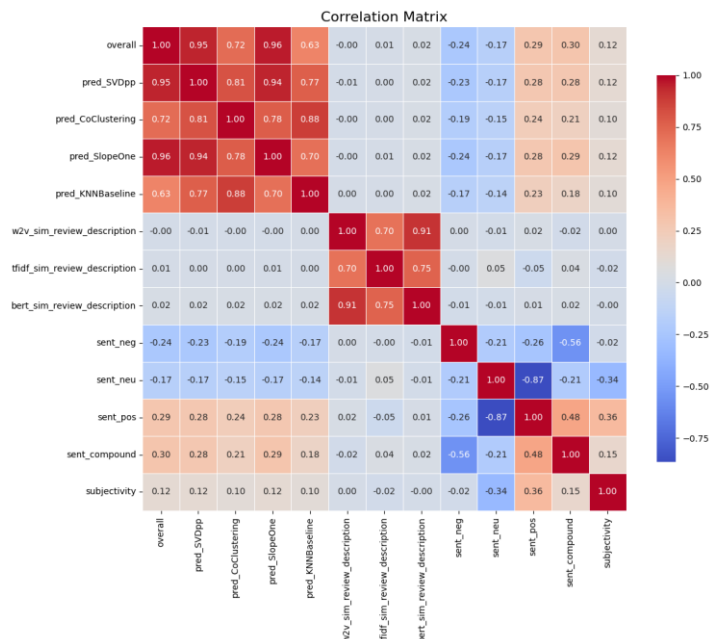


Рисунок 1. Кореляція нових ознак з рейтингом для набору Office Products

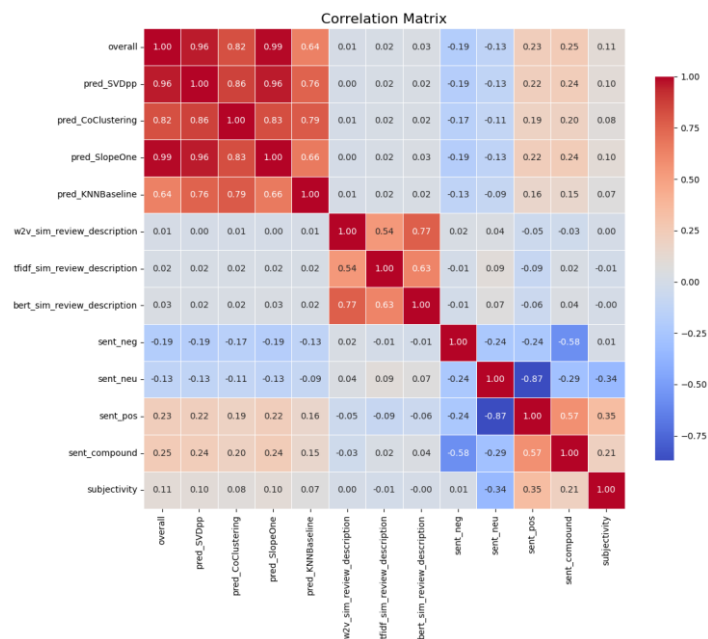


Рисунок 2. Кореляція нових ознак з рейтингом для набору Automotive

Моделі факторизації показують найвищу кореляцію з рейтингом: SVDpp (0.95) та SlopeOne, що робить їх найсильнішими предикторами. Між самими MF-моделями також спостерігається висока узгодженість. Подібність текстових ознак оглядів та описів товарів майже не корелюють із рейтингом, що означає слабку пряму залежність між семантикою та оцінками. Сентиментні ознаки показують помірний зв'язок: позитивна тональність і компаунд корелюють додатно, тоді як негативна – від'ємно. Ознака «subjectivity» впливає мінімально. До того ж значення кореляції ознак подібні для обох наборів.

Результати тестувань моделей колаборативної фільтрації та мета-моделей для даних Office Products та Automotive наведені відповідно у табл. 1 та табл. 2.

Таблиця 1. Результуючі значення метрик для моделей колаборативної фільтрації та мета моделей на основі стекінгу для набору Office Products

	MSE	RMSE	MAE
SVDpp	0.757879	0.870562	0.635171
CoClustering	0.832792	0.912574	0.642022
SlopeOne	0.904019	0.950799	0.665588
KNNBaseline	0.805138	0.897295	0.640438
STACK_RF	0.988260	0.994113	0.663077
STACK_XGB	0.928965	0.963828	0.654766
STACK_LR	0.845425	0.919470	0.661447

Таблиця 2. Результуючі значення метрик для моделей колаборативної фільтрації та мета моделей на основі стекінгу для набору Automotive

	MSE	RMSE	MAE
SVDpp	0.808834	0.899352	0.630363
CoClustering	1.077065	1.037817	0.656875
SlopeOne	1.080132	1.039294	0.684450
KNNBaseline	0.929706	0.964213	0.645797
STACK_RF	1.135888	1.065780	0.666643
STACK_XGB	0.997551	0.998775	0.634821
STACK_LR	0.975901	0.987877	0.670561

У обох наборах даних найнижчі помилки демонструє SVDpp, однак стек-моделі показують конкурентні результати. Зокрема, STACK_XGB стабільно зменшує MAE порівняно з іншими стек-підходами та наближається до точності базових MF-моделей, що свідчить про ефективність поєднання текстових, сентиментних і модельних ознак.

4. ВИСНОВКИ

У дослідженні виявлено, що поєднання колаборативної фільтрації з текстовими та сентиментними ознаками дозволяє покращити якість прогнозування рейтингу. Найсильнішим предиктором у обох наборах даних залишається SVDpp, що підтверджено кореляційним аналізом та метриками. Текстові й емоційні ознаки мають слабший, але додатковий вплив, допомагаючи мета-моделям. Серед стек-підходів найефективнішим виявився STACK_XGB, який зменшує MAE та наближається до базових MF-моделей. Гібридний підхід підтверджує перспективність комбінування семантичних і колаборативних методів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Zhang, S., Yao, L., Sun, A., Tay, Y. "Deep Learning Based Recommender System: A Survey and New Perspectives." ACM Computing Surveys, Vol.52, No.1, P.1-38, 2020.
2. Chen, C., Min, Z., Liu, Y., Ma, S. "Neural Attentional Rating Regression with Review-level Explanations." WWW Conference, 2018.
3. Zhao, J., Wang, X., Li, Y., Chen, H. Exploiting Deep Transformer Models in Review-Based Recommender Systems. Expert Systems with Applications, 2023. URL: <https://www.sciencedirect.com/science/article/pii/S0957417423016226>
4. Hasan, E. et. al. Review-based Recommender Systems: A Survey of Approaches, Challenges and Future Perspectives. ACM Computing Surveys, Vol.58, No.1, P. 1-41, 2025.
5. Reimers, N., Gurevych, I. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. Proceedings of EMNLP-IJCNLP Conference, 2019. DOI: 10.18653/v1/D19-1410

ГІБРИДНА РЕКОМЕНДАЦІЙНА СИСТЕМА ДЛЯ ПІДБІРКИ ВАКАНСІЙ

Бабанов Д.Є.¹, Тимошук О.Л.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ dimonbabanov@gmail.com

Дослідження спрямоване на розробку гібридної рекомендаційної системи для підбору вакансій, що поєднує контентно-орієнтовані методи (TF-IDF, BERT) та колаборативну фільтрацію (User-User, ALS). На основі синтетичних датасетів із регульованою розмірністю проведено порівняльний аналіз ефективності алгоритмів за метриками точності, повноти та NDCG. Наукова новизна полягає в оптимізації гібридної архітектури з динамічним зважуванням методів, що забезпечує стабільну якість рекомендацій навіть за умов розріджених даних. Практична значимість – проведені експерименти розв'язують задачу вибору ефективного метода для побудови моделі рекомендацій.

Ключові слова: гібридна система, колаборативна фільтрація, контентна фільтрація, BERT, TF-IDF, ALS.

1. ВСТУП

Сучасний ринок праці характеризується стрімким зростанням кількості вакансій та людей, що шукають роботу, що призводить до інформаційного перевантаження та ускладнює процес ефективного підбору персоналу. В цих умовах рекомендаційні системи стають інструментом для автоматизації та оптимізації пошуку релевантних кандидатур, забезпечуючи персоналізований підхід до кожного користувача.

Мета даної роботи полягає у розробці гібридної рекомендаційної системи для підбору вакансій, яка поєднує переваги контентно-орієнтованих методів та колаборативної фільтрації. Дослідження зосереджене на порівняльному аналізі різних алгоритмів машинного навчання, оцінці їх ефективності на синтетичних датасетах з регульованою розмірністю та оптимізації гібридної архітектури для досягнення максимальної точності рекомендацій.

Практичне значення отриманих результатів полягає у можливості автоматизації процесу підбору кандидатів, зменшення часу пошуку релевантних вакансій, підвищення якості персоналізованих рекомендацій, скорочення кількості непотрібних відгуків, оптимізації взаємодії між роботодавцями та соціальними працівниками.

2. МАТЕМАТИЧНА МОДЕЛЬ

2.1. Математичне забезпечення моделі

Контентно-орієнтовані методи рекомендацій базуються на аналізі характеристик об'єктів та користувачів. Основною перевагою цих методів є здатність обробляти нові об'єкти (проблема холодного старту) та надавати пояснювальні рекомендації. Однак вони часто страждають від надмірної спеціалізації та не враховують зміни в уподобаннях користувачів з часом. Для аналізу текстових даних вакансій та резюме застосовуються два підходи. Метод TF-IDF використовує статистичну міру важливості слів у документах. Цей метод ґрунтується на

двох основних концепціях: частоті термінів (TF) та оберненій частоті документа (IDF), як показано у (1) та (2).

$$TF = \frac{\text{number of times the term appears in the document}}{\text{total number of terms in the document}} \quad (1)$$

$$IDF = \log \left(\frac{\text{number of the documents in the corpus}}{\text{number of documents in the corpus contain the term}} \right) \quad (2)$$

Своєю чергою для обчислення схожості між документами була використана косинусна подібність, формула якої продемонстрована у (3).

$$\text{Cosine}(x, y) = \frac{\sum_{i=1}^n x_i y_i}{\sqrt{\sum_{i=1}^n x_i^2} \sqrt{\sum_{i=1}^n y_i^2}}, \quad (3)$$

де x та y – це вектори, що описують два об'єкти у векторному просторі.

Для врахування семантичного змісту та контекстної інформації використовується BERT. Ця модель формує контекстно-залежні векторні представлення. BERT використовує архітектуру трансформера, що дозволяє враховувати залежності між всіма словами в реченні одночасно.

Колаборативні методи використовують історичні дані про взаємодії користувачів з системою для прогнозування їх уподобань. Ці методи можуть виявляти складні нелінійні залежності між користувачами та об'єктами, але вони чутливі до розрідженості даних та не можуть обробляти нові об'єкти без додаткових механізмів. Для аналізу взаємодій між користувачем та вакансією застосовується метод ALS (Alternating Least Squares), який є варіантом матричної факторизації. ALS мінімізує функцію втрат з регулюванням їх. ALS постійно мінімізує помилку шляхом чергування оптимізації U при фіксованому V , а потім оптимізації V при фіксованому U :

– зафіксувати V та обчислити U шляхом розв'язання рівняння (4):

$$(\min_u |R - UV^t|^2 + \lambda |U|^2); \quad (4)$$

– зафіксувати U та обчислити V шляхом розв'язання рівняння (5):

$$(\min_v |R - UV^t|^2 + \lambda |V|^2). \quad (5)$$

Повторюючи ці кроки після тренування матриць U та V , ми можемо прогнозувати рейтинг для будь-якого користувача i для продукту j .

Додатково використовується User-User косинусна схожість, яка показана у рівнянні (6):

$$\text{sim}(u, v) = \frac{R_u * R_v}{||R_u|| * ||R_v||}, \quad (6)$$

де R_u та R_v – вектори рейтингів користувачів.

Гібридні системи поєднують різні методи рекомендацій для компенсації їх недоліків. Найпоширенішими підходами є зважене об'єднання, перемикання між методами та пошарове комбінування. Кожен з цих підходів має свої переваги та обмеження, що визначають область їх застосування. Запропонована система реалізує каскадну обробку даних з послідовним застосуванням модулів. Гібридна модель використовує зважене об'єднання методів, що дозволяє компенсувати недоліки окремих підходів та підвищити загальну якість рекомендацій.

Фінальна оцінка для вакансії j обчислюється як сума усіх контентних і колаборативних методів фільтрації перемножених на вагу, тобто їх коефіцієнт, який був обраний експериментальним шляхом підбору.

2.2. Алгоритм побудови моделі

Процес формування рекомендацій включає наступні етапи:

На цьому етапі проводиться векторизація текстів вакансій та профілів кандидатів. Для TF-IDF створюється словник термінів та обчислюються IDF-ваги. Для BERT генеруються ембедінги для кожного текстового опису. Для колаборативної фільтрації будується матриця взаємодій між користувачем та елементом.

Система одночасно обчислює рекомендації за допомогою трьох методів: TF-IDF, ALS, BERT.

TF-IDF: обчислення косинусної схожості між векторизованими описами;

BERT: обчислення семантичної схожості між ембедінгами;

ALS: прогнозування рейтингів на основі матричної факторизації;

Оцінки від різних методів приводяться до єдиної шкали [0,1] за допомогою мінімаксної нормалізації проілюстрованої у рівнянні (7).

$$score_{norm} = (score - \min_score) / (\max_score - \min_score) \quad (7)$$

Виконується лінійна комбінація нормалізованих оцінок з оптимальними вагами. Цей етап дозволяє збалансувати внесок кожного методу у фінальну рекомендацію.

Вакансії сортуються за фінальним рахунком і користувачеві надається персоналізований список рекомендацій. Система також може надавати пояснення рекомендацій, вказуючи внесок кожного методу.

Для оцінки ефективності запропонованої системи було проведено серію експериментів на синтетичних датасетах з регульованою розмірністю. Як метрики якості використовувалися Precision@K, Recall@K та NDCG@K.

Результати експериментів показали, що гібридна система значно перевершує окремі методи за всіма метриками. Зокрема, спостерігалось підвищення Precision@10 на 15–20% порівняно з найкращим окремим методом.

3. РЕЗУЛЬТАТИ РОБОТИ

Результати розрахунків наведені в табл. 1–4.

Таблиця 1. Середні значення метрик для різних методів

Метод	Precision	Recall	NDCG	Час обробки (с)
TF-IDF	0.324	0.287	0.412	1.2
BERT	0.381	0.342	0.478	8.7
User-User	0.295	0.253	0.368	3.4
ALS	0.356	0.318	0.441	4.1
Гібридна	0.447	0.398	0.532	9.8

Таблиця 2. Залежність якості від кількості навичок

Кількість навичок	Precision	Recall	NDCG
20	0.412	0.361	0.487
50	0.447	0.398	0.532
100	0.439	0.391	0.521

Таблиця 3. Вплив вагових коефіцієнтів на якість

Ваги (TF-IDF/BERT/ALS)	Precision	Recall	NDCG
0.4/0.4/0.2	0.428	0.382	0.509
0.3/0.3/0.4	0.447	0.398	0.532
0.2/0.2/0.6	0.431	0.385	0.514
0.5/0.3/0.2	0.439	0.391	0.523

Таблиця 4. Ефективність для нових користувачів

Метод	Precision@10	Recall@10
TF-IDF	0.315	0.278
BERT	0.372	0.334
User-User	0.124	0.098
ALS	0.153	0.121
Гібридна	0.389	0.347

4. ВИСНОВКИ

Запропонована гібридна система демонструє високу ефективність у задачі підбору вакансій, перевершуючи окремі методи за всіма метриками якості. Математичний апарат, що включає TF-IDF, BERT та ALS, забезпечує стабільну роботу навіть за умов обмежених даних.

Основними перевагами системи є:

- висока точність рекомендацій через комбінування методів;
- здатність обробляти нових користувачів та вакансії;
- масштабованість та ефективність обробки даних;
- адаптивність до різних сценаріїв використання.

Перспективи подальших досліджень включають впровадження глибинних нейронних мереж для покращення якості рекомендацій, розробку механізмів онлайн-навчання для адаптації до змін уподобань користувачів, а також інтеграцію додаткових джерел даних, таких як соціальні мережі та професійні профілі.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Eda Kavlakoglu Jacob Murel Ph.D. What is content-based filtering? URL: <https://www.ibm.com/think/topics/content-based-filtering>
2. Hael Al-bashiria,, Mansoor Abdullateef Abdulgabberra, Awanis Romlia, Fadhl Hujainah Collaborative Filtering Recommender System: Overview and Challenges. URL: https://www.researchgate.net/publication/320761149_Collaborative_Filtering_Recommender_System_Overview_and_Challenges#fullTextFileContent September 2017
3. Stephen E. Robertson Understanding Inverse Document Frequency: On Theoretical Arguments for IDF URL: https://www.researchgate.net/publication/238123710_Understanding_Inverse_Document_Frequency_On_Theoretical_Arguments_for_IDF October 2004 pp..503-520
4. Jennifer Scott Miroslav Tuma Introduction to Matrix Factorizations URL: https://www.researchgate.net/publication/370405640_Introduction_to_Matrix_Factorizations January 2023 pp.. 31-51
5. Jacob Devlin Ming-Wei Chang Kenton Lee Kristina Toutanova BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding URL: https://www.researchgate.net/publication/328230984_BERT_Pre-training_of_Deep_Bidirectional_Transformers_for_Language_Understanding October 2018
6. Sanna Persson Paper summary — BERT: Bidirectional Transformers for Language Understanding. URL: <https://medium.com/analytics-vidhya/paper-summary-bert-pre-training-of-deep-bidirectional-transformers-for-language-understanding-861456fed1f9> May 16, 2021

МОДЕЛЮВАННЯ ЗМІНИ ЦІНИ КРИПТОАКТИВІВ В ЗАЛЕЖНОСТІ ВІД ЇХНЬОЇ ЕМІСІЇ

Биль К.І.¹, Терентьев О.М.²

¹ Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

² Інститут телекомунікацій і глобального інформаційного простору
Національної академії наук України, Київ, Україна

¹ kirilbyl@gmail.com [0009-0006-6030-1720],

² o.terentiev@gmail.com [0000-0002-4288-1753],

Мета роботи полягає в аналізі статистичних даних, щодо впливу зміни емісії монет різних криптоактивів на зміну ціни. Для чого використовуються методи статистичного аналізу та побудова моделі логістичної регресії, для прогнозування зміни напрямку ціни криптоактиву. Для вирішення цієї задачі в роботі було використано ринкові дані цін з біржи Binance та новини про розблокування криптоактивів з агрегатору інформації сайту CoinMarketCap.

Ключові слова: математичне моделювання, логістична регресія, ринок криптовалют, емісія криптоактивів.

1. ВСТУП

Емісія криптоактивів – це механізм появи нових одиниць цифрових валют чи токенів, який забезпечує роботу будь-якої блокчейн-платформи. Подібно до створення грошей у класичній фінансовій системі, емісія криптовалют визначає обсяг їхнього обігу, ступінь дефіцитності та впливає на їхню вартість [1].

2. ЦИРКУЛЯЦІЯ КІЛЬКОСТІ МОНЕТ КРИПТОАКТИВУ, ТА ЗВ'ЯЗОК ІЗ ЦІНОЮ

На основі даних CoinMarketCap [2] щодо описів криптоактивів, строків та типів розблокування, а також даних щодо історії цін з сайту біржи Binance [3] було сформовано аналітичну таблицю та проведено моделювання історії торгів на основі запропонованого алгоритму. Загалом історія торгів містить 168 угод, які було оброблено алгоритмом, за період часу з 5 листопада 2020 року по 16 травня 2025 року.

Таблиця історії торгів містить наступні стовпчики:

- 1) Short_token – тікер криптоактиву, який аналізується, по якому алгоритм виконує угоду.
- 2) Unlock_date – дата рзблокування криптоактиву, тобто факт емісії криптоактиву у зазначений час та дату.
- 3) Percentage – процент розблокованих криптоактивів.
- 4) Entry_date – дата входу в угоду по даному криптоактиву.
- 5) Exit_date – дата виходу з угоди по даному криптоактиву.
- 6) Entry_price_short – ціна криптоактиву на момент входу в угоду.
- 7) Exit_price_short – ціна криптоактиву на момент виходу з угоди.
- 8) Short_profit – дохідність по відповідній угоді, в біпсах.

Нижче на рисунку 1 наведено перші 7 з 168 угод, які було оброблено алгоритмом на реальних даних.

Номер угоди	short_token (тікер)	unlock_date (дата розблокування)	percentage (процент розблокування)	entry_date (дата входу в угоду)	exit_date (дата виходу з угоди)	entry_price_short (ціна входу в угоду)	exit_price_short (ціна виходу з угоди)	short_profit (дохідність в біпсах)
1	BAND	2020-11-05	3,14	2020-10-29	2020-11-10	4,629	6,2902	-35,8868006
2	SNX	2020-12-15	3,88	2020-12-08	2020-12-20	4,669	5,967	-27,80038552
3	AVAX	2020-12-20	6,45	2020-12-13	2020-12-25	3,3445	2,9534	11,69382568
4	CTK	2021-01-24	6,41	2021-01-17	2021-01-29	1,114	0,90536	18,72890485
5	BAND	2021-02-17	3,97	2021-02-10	2021-02-22	13,1091	16,9522	-29,31627648
6	BAND	2021-03-03	3,14	2021-02-24	2021-03-08	13,2246	14,6601	-10,85477065
7	RSR	2021-03-05	3,81	2021-02-26	2021-03-10	0,045913	0,067837	-47,75118158

Рисунок 1. Фрагмент таблиці історії торгів на основі інформації про криптоактиви та цін, наведено перші 7 з 168 угод

На основі інформації отриманої з файлу історичних даних розблокування криптоактивів було проведено аналіз, щодо розподілу типів розблокування криптоактивів, який представлено у вигляді таблиці 1. Як можна побачити з таблиці 1, лише у 0,61% випадків змінна “Тип розблокування” містить значення пропуску (тобто відсутня інформація), при цьому у більшості випадків, а саме у 92,25% розблокування має лінійний характер, інші три види розблокування мають досить невеликі значення від 1,72% ситуація для “дефляційного” виду, 1,78% для “інфляційного” та 3,64% для “скелястого”.

Таблиця 1. Розподіл типів розблокування (змінна аналізу alloc_vestingType)

alloc_vestingType	Опис типу розблокування	Кількість ситуацій	Відсоток ситуацій (%)
Missing	пусте значення	400	0,61
cliff	скелястий	2 368	3,64
deflationary	дефляційний	1 118	1,72
inflationary	інфляційний	1 161	1,78
linear	лінійний	60 079	92,25

Наведемо стислий опис різних типів розблокування криптоактивів:

- 1) **Cliff** (скелястий) – одноразове розблокування після заданого періоду (наприклад, команда не може продавати токени перші 12 місяців, потім все або частина стає доступною).
- 2) **Deflationary** (дефляційний) – жорстка формула емісії (наприклад, зменшення емісії з часом) чи гнучка.
- 3) **Inflationary** (інфляційний) – нові токени регулярно випускаються як винагорода стейкерам/валідаторам – це форма постійного розблокування через майнінг/стейкінг.
- 4) **Linear** (лінійний) – це механізм поступового розблокування tokenів криптопроекту протягом певного періоду години рівномірними частинами.

Окрім цього на основі отриманих історичних даних, щодо розблокування криптоактивів було проведено статистичний аналіз розподілу значень класів розподілу монет, що розблоковуються, який наведено у вигляді таблиці 2.

Як можна побачити з таблиці 2, значень частот для класів розподілу монет, у 22,48% випадків це – “Казначейський фонд / Екосистемний фонд”, у 16% випадків це – “Команда та консультанти”, у 14,2% випадках – це “Винагорода за стейкінг / майнінг”, у 13,6% випадках – це “Приватні/початкові інвестори”, всі інші 11 класів розподілу займають дрібні частини від 3,27% та менше, але сумарно це приблизно 33% випадків.

Таблиця 2. Розподіл значень класів розподілу монет, що розблоковуються (змінна аналізу alloc_class)

№	alloc_class (клас розподілу)	Опис типу класу розподілу	Кількість ситуацій	Відсоток ситуацій
1	Other	інше	17 517	26,90
2	Treasury / Ecosystem Fund	Казначейський фонд / Екосистемний фонд	14 642	22,48
3	Team & Advisors	Команда та консультанти	10 481	16,09
4	Staking / Mining Rewards	Винагороди за стейкінг / майнінг	9 133	14,02
5	Private/Seed Investors	Приватні/початкові інвестори	8 910	13,68
6	Airdrop & Community Incentives	Ейрдроп та стимули для громади	2 129	3,27
7	Marketing / Operations	Маркетинг / Операції	1 702	2,61
8	Burn / Locked / Vesting	Запуск / Заблоковано / Надання права власності	432	0,66
9	Public Sale / ICO	Публічний продаж / ICO	180	0,28

3. МОДЕЛЬ ЛОГІСТИЧНОЇ РЕГРЕСІЇ ДЛЯ ПРОГНОЗУВАННЯ НАПРЯМКУ ЗМІНИ ЦІНИ

На основі отриманих даних була побудована модель логістичної регресії для прогнозування напрямку руху ціни криптоактиву, в залежності від таких змінних:

- тип розблокування (номінальна змінна, що має 5 станів);
- клас розподілу монет (номінальна змінна, що має 9 станів).

Код програми для побудови моделі логістичної регресії, на мові програмування SAS Base має наступний вигляд:

```
ods graphics on;
proc logistic data=fd.ABT_v_1 alpha=.05
    plots(only)=(effect oddsratio);
    CLASS AVT_TYPE alloc_class;
    model sign_short_profit=AVT_TYPE alloc_class / clodds=pl;
    OUTPUT OUT=WORK.PREDLogRegPredictions
        PREDPROBS=INDIVIDUAL;
run;
```

Статистичні характеристики побудованої моделі логістичної регресії (таблиця 3):

- True-Negative = 117
- False-Negative = 76
- False-Positive = 10
- True-Positive = 362

Загалом на ретроспективних даних отримана модель логістичної регресії дозволяє вірно прогнозувати напрямок зміни ціни у 84% ситуацій.

Таблиця 3. Значення оцінок коефіцієнтів моделі логістичної регресії

Параметр моделі	Значення	Значення оцінки коефіцієнта моделі	Стандартна похибка	Хі-квадрат статистика Вольда	р-значення для статистики хі-квадрат
Intercept		-2,8013	21,9275	0,0163	0,8983
AVT_TYPE	CLIFF	-0,5901	0,2137	7,6248	0,0058
AVT_TYPE	CLIFF / INEAR	0,1164	0,2246	0,2684	0,6044
alloc_class	Airdrop & Community Incentives	2,9068	21,9337	0,0176	0,8946
alloc_class	Burn / Locked / Vesting	-6,9552	157,2	0,002	0,9647
alloc_class	Marketing / Operations	2,3338	21,9323	0,0113	0,9153
alloc_class	Other	1,6769	21,9284	0,0058	0,939
alloc_class	Private/Seed Investors	1,4295	21,9282	0,0042	0,948
alloc_class	Public Sale / ICO	-8,0191	80,863	0,0098	0,921
alloc_class	Staking / Mining Rewards	2,6989	21,933	0,0151	0,9021
alloc_class	Team & Advisors	1,9807	21,9279	0,0082	0,928

4. ОПИС РОБОТИ АЛГОРИТМУ ТОРГІВЛІ

1. Вивантажуються дані з сайтів-агрегаторів щодо майбутніх розблокувань криптоактивів.

2. Перевіряється капіталізація крипто-проекту за відкритими джерелами даних, і якщо капіталізація достатньо велика, наприклад входить до ТОП-4000 крипто-проектів, за цим показником, то проект цікавий і аналізується далі.

3. Перевіряється значення розміру розблокування кількості монет крипто-проекту, якщо розмір розблокування досить суттєвий, наприклад більший за 2% від поточної кількості монет в обігу, то проект аналізується далі.

4. За 10–14 діб до розблокування монет крипто-проекту, виставляється ордер на продаж, по максимально вигідній ціні.

5. Після розблокування монет крипто-проекту, через 5–14 діб виставляється ордер на покупку, по максимально вигідній ціні.

5. РЕЗУЛЬТАТИ РОБОТИ АЛГОРИТМУ ТОРГІВЛІ

На рисунку 2 наведена загальна кумулятивна дохідність роботи торгівельного алгоритму за період з 5 листопада 2020 року по 16 травня 2025 року, загалом було здійснено 168 угод, які принесли загалом 633 біпси доходу, що становить майже 6,3% дохідності.

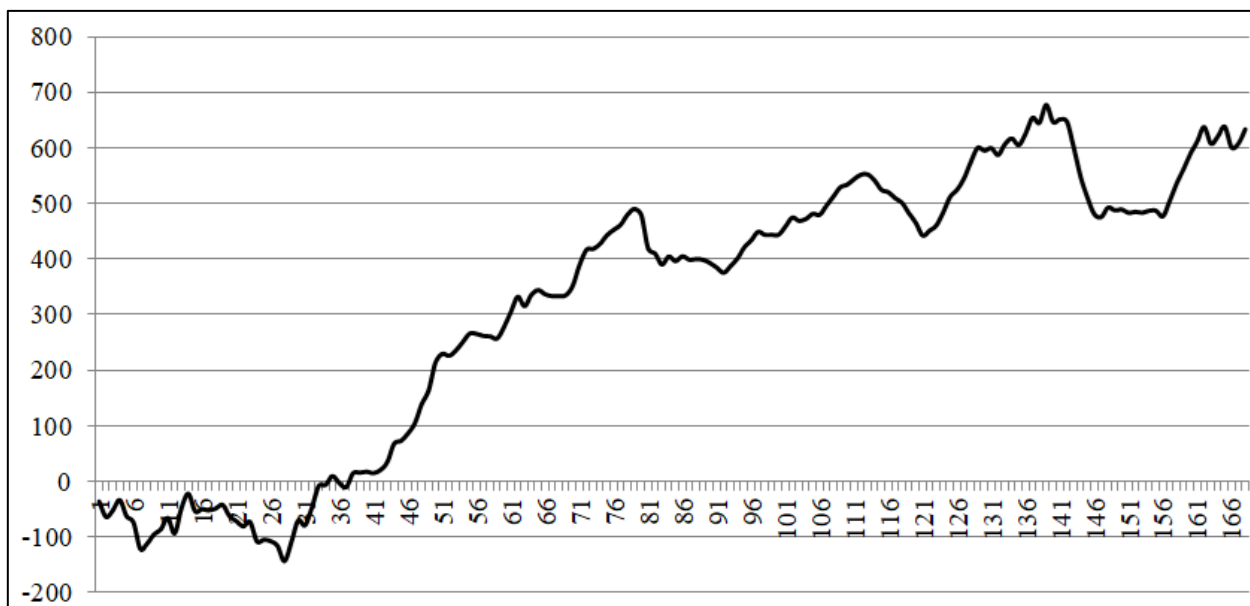


Рисунок 2. Ось абсцис – номер угоди, ось ординат – загальна кумулятивна дохідність роботи торгівельного алгоритму за період з 5 листопада 2020 року по 16 травня 2025 року, загалом 168 угод

На рисунку 3 наведено графік значення максимальної просадки (збитку) роботи алгоритму за період з 5 листопада 2020 року по 16 травня 2025 року, загалом 168 угод, з якого можна побачити, що найбільший збиток-просадка дорівнюють 202 біпси.

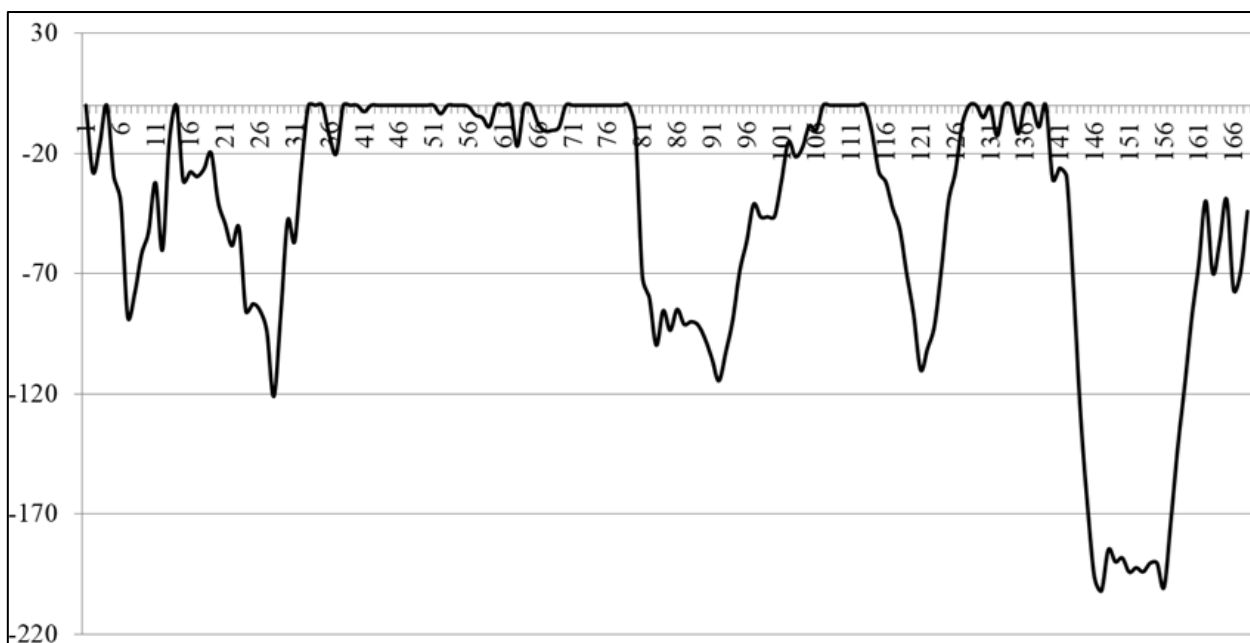


Рисунок 3. Ось абсцис – номер угоди, ось ординат – значення максимальної просадки (збитку) роботи алгоритму за період з 5 листопада 2020 року по 16 травня 2025 року, загалом 168 угод

Основний висновок, який можна зробити з отриманих результатів торгівлі, наступний – співвідношення прибуткових угод до збиткових становить 61% прибуткових проти 39% збиткових, при чому:

- Отримано загалом 65 збиткових угод.
- Середній розмір збиткової угоди 15 біпсів.
- Максимальний розмір збиткової угоди 60 біпсів.
- Сумарно 993 біпси збитків, по збиткових 65 угодах.
- Отримано загалом 103 прибуткові угоди.
- Середній розмір прибуткової угоди 16 біпсів.
- Максимальний розмір прибуткової угоди 50 біпсів.
- Сумарно 1626 біпсів прибутку, по 103 прибуткових угодах.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Howell, S. T., Niessner, M., Yermack, D. Initial Coin Offerings: Financing Growth with Cryptocurrency Token Sales // The Review of Financial Studies. – 33(9), 2019. — p. 3925-3974.
2. CoinMarketCap – <https://coinmarketcap.com/> (дата звернення: 20.11.2025).
3. Сайт біржи Binance – <https://www.binance.com/> (дата звернення: 20.11.2025).

МОДЕЛЮВАННЯ ТОРГОВИХ СИГНАЛІВ НА ОСНОВІ АНАЛІЗУ ФІНАНСОВИХ ЧАСОВИХ РЯДІВ ЗА ДОПОМОГОЮ МЕРЕЖ КОЛМОГОРОВА-АРНОЛЬДА

Деркач А.С.¹, Шаповал Н.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ derkach,oni@gmail.com, ² shapoval.nataliia@lil.kpi.ua

Прогнозування фінансових часових рядів, зокрема на ринках криптовалют, є складним завданням через високу волатильність, нестационарність даних та значний рівень шуму. Мета цього дослідження полягає у розробці адаптивної системи торгових сигналів на основі новітньої архітектури нейронних мереж Колмогорова-Арнольда для підвищення стійкості стратегій та забезпечення їх інтерпретованості. У ході роботи було проведено порівняльний аналіз ефективності KAN з класичними нейромережами MLP та базовими ринковими стратегіями. Її наукова новизна полягає у використанні навчальних сплайнових функцій активації для виявлення прихованих нелінійних залежностей та проведенні символічної дистиляції моделі, що дозволило отримати прозорі аналітичні формули прийняття рішень. Результати підтверджують перевагу методу у співвідношенні ризик/прибуток.

Ключові слова: мережі Колмогорова-Арнольда, фінансові часові ряди, алгоритмічна торгівля, символічна регресія, управління ризиками, волатильність.

1. ВСТУП

Сучасні фінансові ринки, і особливо сектор криптовалют, характеризуються високим рівнем невизначеності, нелінійною динамікою та наявністю важких хвостів у розподілі прибутковості. В умовах, коли традиційні методи технічного аналізу втрачають ефективність через алгоритмізацію торгівлі, виникає потреба у застосуванні адаптивних методів штучного інтелекту. Однак, попри високу точність сучасних методів глибокого навчання, їх широке впровадження у фінансовий сектор нашою є фундаментальна проблема.

Класичні архітектури, такі як багатошарові персептрони або рекурентні мережі, функціонують за принципом «чорної скриньки». Внутрішня логіка прийняття рішень у таких моделях прихована за мільйонами вагових коефіцієнтів, що унеможливує пояснення причини генерації того чи іншого торгового сигналу. Для інституційних інвесторів та регуляторів це створює неприйнятні ризики, оскільки неможливо розрізнити, чи базується прибуток на виявленні реальної ринкової закономірності, чи на випадковій підгонці під історичні дані [6].

Дана робота пропонує вирішення цієї дилеми шляхом використання новітньої архітектури — мереж Колмогорова-Арнольда. На відміну від MLP, де функції активації фіксовані, KAN використовують навчальні сплайни, розташовані на ребрах мережі. Це дозволяє отримати високу точність прогнозування, характерну для нейромереж, та

забезпечити інтерпретованість результатів, характерну для символічної регресії. Дослідження спрямоване на розкриття чорної скриньки та отримання прозорих формул торгових стратегій.

2. МЕТОДОЛОГІЯ ДОСЛІДЖЕННЯ ТА АРХІТЕКТУРА НЕЙРОННОЇ МЕРЕЖІ

Торгівля криптовалютами характеризується високим рівнем шуму та нестационарністю цінових рядів. Для забезпечення стабільної роботи нейронної мережі було застосовано метод трансформації сирих ринкових даних у стаціонарні технічні індикатори. Замість абсолютних цін, які постійно зростають і виходять за межі діапазону активації мережі, у роботі використано відносні величини [3]:

1. Логарифмічна дохідність: $R_t = \ln\left(\frac{P_t}{P_{t-1}}\right)$. Цей показник нормалізує цінові стрибки та робить розподіл даних більш симетричним.
2. Нормований індикатор RSI. Класичний індикатор Relative Strength Index було масштабовано у діапазон $[-1, 1]$. Це дозволяє моделі чітко розрізняти стани перекупленості та перепроданості.
3. Коефіцієнт волатильності (Volatility Ratio). Відношення короткострокової волатильності до довгострокової. Це дозволяє детестувати зміну ринкових режимів.
4. Відхилення від тренду. Відносна відстань поточної ціни від ковзної середньої, що слугує індикатором для стратегій повернення до середнього.

Для боротьби з аномальними викидами було застосовано метод Robust Scaling, який використовує медіану та інтерквартильний розмах замість середнього значення та дисперсії. Вхідні дані додатково обмежувалися діапазоном $[-1.5, 1.5]$ для стабільної роботи сплайнів.

Основним інструментом дослідження обрано архітектуру KAN [1]. На відміну від класичних багатопараметричних перцептронів, де функції активації фіксовані на вузлах, у KAN навчальні функції розташовані на ребрах графа.

Математичною основою моделі є теорема подання Колмогорова-Арнольда, згідно з якою будь-яка багатовимірна неперервна функція може бути представлена як суперпозиція одновимірних функцій. У програмній реалізації ці одновимірні функції апроксимуються за допомогою В-сплайнів порядку $k=3$.

Процес проходження сигналу через шар KAN описується формулою [2]:

$$\phi(x) = \omega_b \cdot \sigma(x) + \omega_s \cdot B(x)$$

де:

- $\sigma(x)$ – базова функція активації, що відповідає за загальний тренд;
- $B(x)$ – лінійна комбінація В-сплайнів, що відповідає за моделювання локальних нелінійностей;
- w_b, w_s – навчальні вагові коефіцієнти.

Така структура дозволяє мережі локально змінювати свою форму залежно від вхідних даних, що забезпечує вищу точність при меншій кількості параметрів порівняно з MLP.

Наукова новизна методології полягає у впровадженні механізму символічної дистиляції, який дозволяє перетворити навчену нейромережу у набір зрозумілих аналітичних формул. Процес екстракції знань, реалізований у програмному коді, складається з трьох етапів:

1. Ізоляція факторів. Для аналізу впливу конкретного індикатора на вихідний сигнал, ми фіксуємо всі інші вхідні параметри на рівні їх середніх значень, а досліджуваний параметр змінюємо у діапазоні від -1 до 1 .

2. Сканування реакції. Згенерований вектор значень подається на вхід навченої моделі KAN. На виході фіксується масив ймовірностей росту ціни $P(U_p)$. Це дозволяє отримати графік функції активації, яку вивчила мережа для конкретного фактора.
3. Поліноміальна апроксимація. Отримана крива апроксимується поліномом другого ступеня методом найменших квадратів:

$$f(x) \approx ax^2 + bx + c$$

Аналіз коефіцієнтів a та b дозволяє інтерпретувати логіку моделі:

- а. Якщо $|a| \gg |b|$, залежність є нелінійною, що вказує на стратегію торгівлі волатильністю.
- б. Якщо $|b| \gg |a|$, залежність є лінійною, що вказує на трендову стратегію.

На основі прогнозів моделі KAN побудовано адаптивну торгову стратегію з механізмом фільтрації сигналів. Алгоритм прийняття рішень виглядає наступним чином:

1. Генерація ймовірностей. Мережа видає вектор ймовірностей $[P_{down}, P_{up}]$ через функцію Softmax.
2. Фільтрація шуму. Замість прийняття рішення за принципом більшості ($P > 0.5$), введено поріг впевненості $\theta = 0.52$.
 - а. Якщо $P_{up} > \text{Сигнал на купівлю}$.
 - б. Якщо $P_{down} > \text{Сигнал на продаж}$.
 - с. В іншому випадку (діапазон невпевненості) позиція не змінюється.
3. Виконання угод. Симуляція торгівлі проводиться на історичних даних з урахуванням реінвестування прибутку.

Такий підхід дозволяє відсіяти слабкі сигнали, характерні для бокового руху ринку, та суттєво знизити максимальну просадку капіталу порівняно з базовими стратегіями.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕНЬ

3.1. Валідація математичної точності моделі

Першим етапом дослідження була перевірка здатності архітектури KAN апроксимувати складні нелінійні функції. Як еталонну модель було обрано формулу Блека-Шоулза, яка є стандартом для оцінки вартості опціонів [4].

Було згенеровано навчальну вибірку, що складається з вхідних параметрів (ціна активу, страйк, час, ставка, волатильність) та цільового значення (теоретична ціна).

Після навчання мережа досягла MSE Loss \$0.00007\$, побудована крива (рис. 1).

Низьке значення помилки підтверджує, що сплайнові функції активації KAN здатні з високою точністю емулювати складні математичні операції, закладені в теорію ціноутворення.

3.2. Апробація моделі на реальних ринкових даних

Наступним кроком стало тестування моделі в умовах реального ринку, де діють фактори, не враховані в ідеальній формулі. Для експерименту було завантажено ланцюжок опціонів на акції компанії Apple (AAPL), рис. 2.

Модель KAN навчалася відновлювати залежність премії опціону від його параметрів без використання формули Блека-Шоулза, спираючись виключно на ринкові ціни угод.

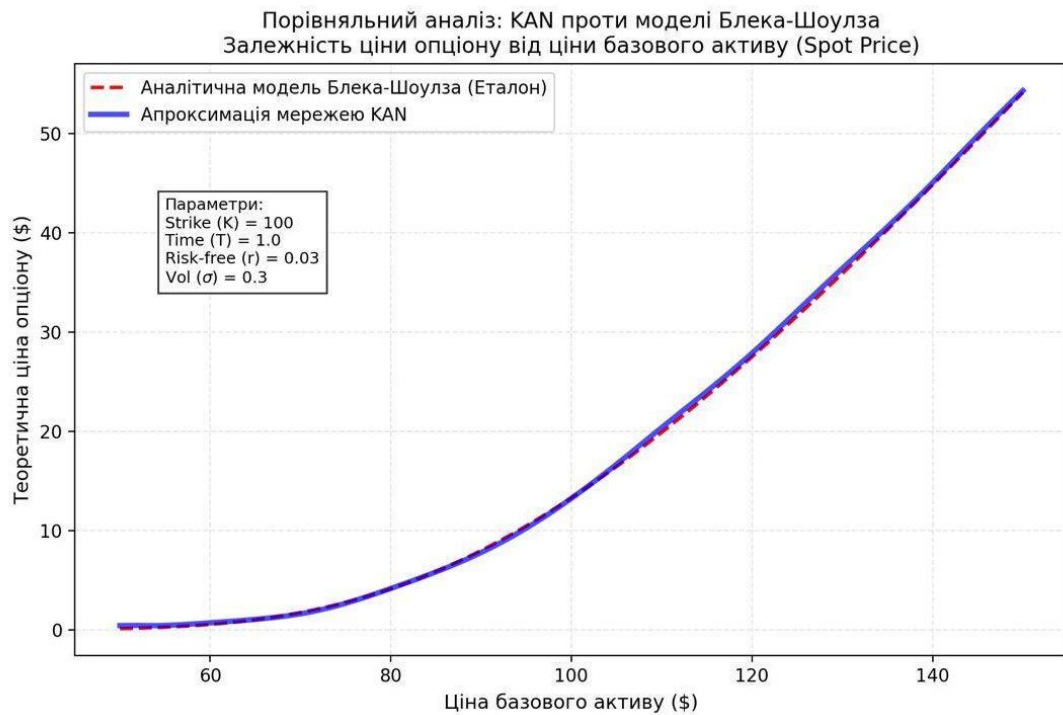


Рисунок 1. Результат апроксимації формули Блека-Шоулза на синтетичних даних

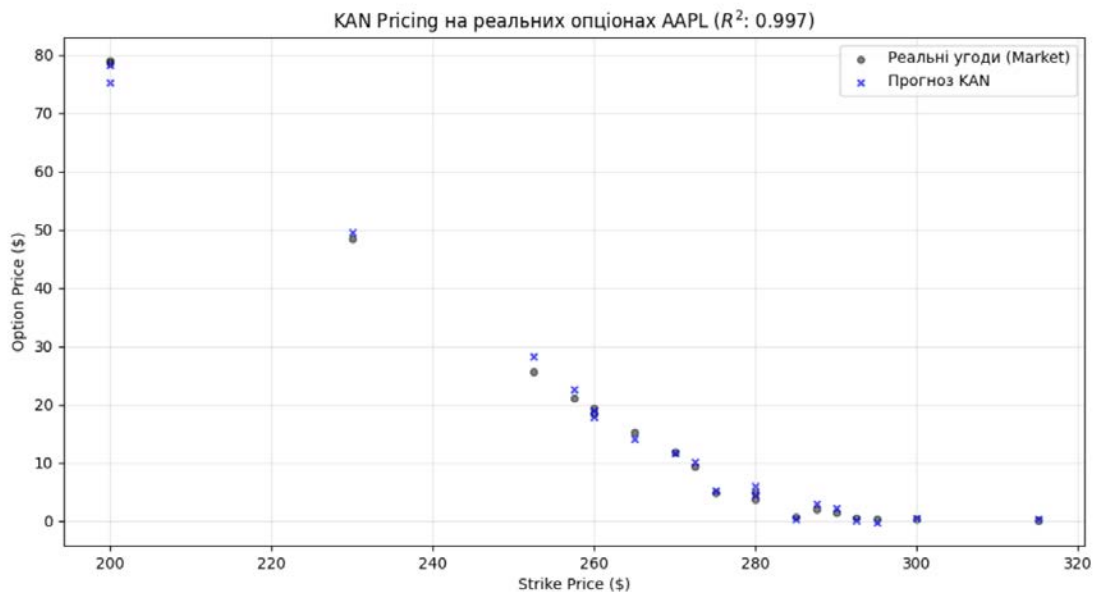


Рисунок 2. Діаграма розсіювання прогнозованих та реальних цін опціонів AAPL

Після навчання на реальних даних отримано такі метрики (рис. 3):

```
=====
РЕЗУЛЬТАТИ (REAL MARKET DATA)
=====
R2 Score (Точність KAN): 0.9972
Середня помилка (MSE): $1.24
```

Рисунок 3. Метрики отримані після навчання

Помилка у розмірі \$1.24 є допустимою для високоліквідних активів з високою вартістю контракту. Це доводить, що KAN здатна адаптуватися до емпіричних ринкових аномалій.

3.3. Порівняльний аналіз ефективності торгових стратегій

Заключним етапом експерименту стала перевірка ефективності розробленої архітектури UltraKAN у задачі алгоритмічної торгівлі. Для всебічного аналізу було протестовано чотири стратегії:

1. UltraKAN (Aggressive): Стратегія, що відкриває угоду при будь-якому, навіть мінімальному перехилі ймовірності ($P > 0.5$).
2. UltraKAN (Safe): Консервативна стратегія, що використовує поріг впевненості (Confidence Threshold $P > 0.52$). Угоди відкриваються лише тоді, коли мережа впевнена у сигналі, що дозволяє фільтрувати ринковий шум.
3. MLP (Standard): Класичний багатошаровий перцептрон, що слугує еталоном для перевірки переваг сплайнової архітектури.
4. Buy & Hold: Пасивна стратегія утримання активу.

Результати моделювання. Експерименти проводилися на трьох активах: Bitcoin (BTC) (рис. 4), Ethereum (ETH) та Meta Platforms (META). Зведені показники ефективності наведені у таблиці 1.

Таблиця 1. Порівняльна ефективність стратегій

Актив	Стратегія	Прибуток (Return)	Ризик (Max DD)	Sharpe Ratio [5]
BTC-USD	UltraKAN (Safe)	106.23%	-26.44%	1.33
	UltraKAN (Aggr)	99.25%	-25.07%	1.30
	MLP (Standard)	18.71%	-31.21%	0.44
	Buy & Hold	41.58%	-32.15%	0.66
ETH-USD	UltraKAN (Safe)	102.96%	-41.27%	1.07
	UltraKAN (Aggr)	38.23%	-49.70%	0.62
	MLP (Standard)	41.08%	-55.40%	0.66
	Buy & Hold	30.10%	-63.24%	0.55
META	UltraKAN (Safe)	11.00%	-30.46%	0.41
	UltraKAN (Aggr)	7.69%	-30.63%	0.33
	MLP (Standard)	-1.34%	-25.94%	0.11
	Buy & Hold	27.82%	-34.15%	0.70

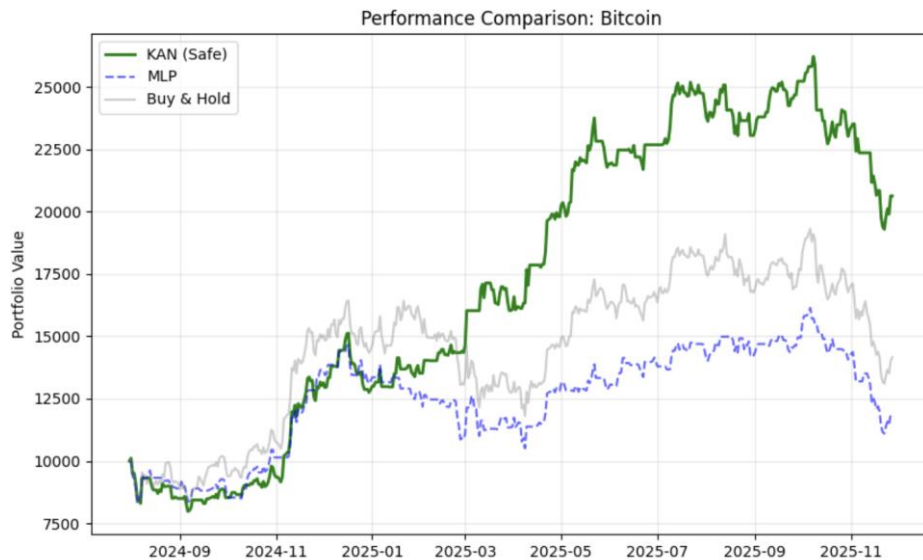


Рисунок 4. Динаміка капіталу стратегій на активі Bitcoin

Аналіз отриманих даних.

1. Ефект фільтрації шуму (Safe vs Aggressive):

Найбільш показовим є результат на активі Ethereum (ETH). Агресивна версія KAN заробила 38.23%, тоді як безпечна версія (Safe) досягла 102.96%. Це пояснюється тим, що ETH є високоволатильним активом з великою кількістю хибних рухів. Введення порогу впевненості дозволило моделі ігнорувати слабкі сигнали в періоди невизначеності, зберігаючи капітал для сильних трендових рухів.

2. Домінування над MLP:

У всіх трьох тестах архітектура KAN перевершила класичний персептрон. Особливо це помітно на акціях META, де MLP показала збиток (-1.34%), тоді як UltraKAN (Safe) вийшла в плюс (+11.00%). Це підтверджує, що навчальні сплайни KAN краще узагальнюють дані та менше схильні до перенавчання, ніж фіксовані нейрони MLP.

3. Перевага над ринком:

На ринку криптовалют (BTC, ETH) стратегія UltraKAN (Safe) показала кращі результати, ніж просте утримання активів, як за прибутком, так і за коефіцієнтом Шарпа.

3.4. Символьна інтерпретація навченої мережі

Однією з ключових переваг KAN є можливість розкрити чорну скриньку та зрозуміти логіку прийняття рішень. За допомогою процедури символьної дистиляції було відновлено функціональні залежності ймовірності росту ціни $P(Up)$ від вхідних індикаторів.

Результати символьної регресії представлені нижче:

1. Фактор RSI (RSI_Norm):

$$P(Up) \approx 0.11x^2 + 0.03x + 0.51$$

Інтерпретація: Отримана параболічна залежність (U-подібна форма) вказує на те, що модель ідентифікувала патерн Volatility Smile. Найбільша ймовірність успішного входу в позицію виникає при екстремальних значеннях RSI, що свідчить про використання стратегії пробою рівнів.

2. Фактор Волатильності (Vol_Ratio):

$$P(Up) \approx -0.09x^2 - 0.03x + 0.46$$

Інтерпретація: Обернена парабола. Від'ємний коефіцієнт при квадратичному члені свідчить про те, що ймовірність росту ціни максимальна при стабільній волатильності. При аномально високій волатильності модель знижує прогноз росту, ефективно уникаючи ризикових ситуацій.

3. Трендові фактори (Dist_SMA та Mom):

Для індикаторів відхилення від середньої та моментуму модель вивела залежності з меншим впливом нелінійності:

a. Dist_SMA: $0.04x^2 - 0.01x$

b. Mom: $0.00x^2 - 0.05x$

4. ВИСНОВОК

У дипломній роботі вирішено науково-прикладну задачу підвищення ефективності алгоритмічної торгівлі на високоризикових фінансових ринках шляхом розробки адаптивної системи на базі мереж Колмогорова-Арнольда. Результати порівняльного експерименту довели, що запропонована архітектура завдяки використанню навчальних сплайнових функцій активації здатна перевершити класичні багатошарові перцептрони та пасивні стратегії за показниками прибутковості та стійкості до ризику, досягаючи на криптовалютних активах коефіцієнта Шарпа понад 1.3.

Ключовим науковим результатом роботи стала демонстрація здатності мережі KAN до символічної дистиляції знань, що дозволило вирішити проблему інтерпретованості рішень чорної скриньки. Зокрема, модель самостійно ідентифікувала нелінійний ефект усмішки волатильності, який є відомим емпіричним виключенням із класичної моделі ціноутворення Блека-Шоулза. Цей факт свідчить про те, що архітектура KAN не просто апроксимує вхідні дані, а здатна виявляти приховані структурні аномалії та відхилення від теоретичних правил. Отримані результати відкривають перспективи подальшого використання мереж Колмогорова-Арнольда не лише для прогнозування трендів, а й як інструменту фінансового інжинірингу для автоматизованого пошуку ринкових неефективностей та уточнення існуючих економічних моделей на основі реальних даних.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Liu Z., Wang Y., Vaidya S. et al. KAN: Kolmogorov-Arnold Networks. *arXiv preprint arXiv:2404.19756*. 2024. DOI: <https://doi.org/10.48550/arXiv.2404.19756>.
2. Kolmogorov A. N. On the representation of continuous functions of many variables by superposition of continuous functions of one variable and addition. *Doklady Akademii Nauk SSSR*. 1957. Vol. 114. P. 953–956.
3. Lopez de Prado M. *Advances in Financial Machine Learning*. Hoboken : John Wiley & Sons, 2018. 400 p.
4. Black F., Scholes M. The Pricing of Options and Corporate Liabilities. *Journal of Political Economy*. 1973. Vol. 81, No. 3. P. 637–654.
5. Sharpe W. F. The Sharpe Ratio. *Journal of Portfolio Management*. 1994. Vol. 21, No. 1. P. 49–58.
6. Goodfellow I., Bengio Y., Courville A. *Deep Learning*. Cambridge : MIT Press, 2016. 800 p.

ЗАСТОСУВАННЯ МЕТОДІВ МАТЕМАТИЧНОГО МОДЕЛЮВАННЯ В СИСТЕМАХ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ З ПОВОЄННОГО ВІДНОВЛЕННЯ

Дуда В.О.¹, Терент'єв О.М.²

Інститут телекомунікацій і глобального інформаційного простору
Національної академії наук України, Київ, Україна

¹ dudavolodimir@gmail.com [0009-0002-4278-4635], ² o.terentiev@gmail.com [0000-0002-4288-1753]

Метою роботи є розроблення методології створення системи підтримки прийняття рішень, призначеної для вирішення задач управління повоєнним відновленням як на рівні держави, так і регіону та галузі національної економіки. Основу цієї системи становлять нові інформаційні технології збору, обробки, накопичення великих обсягів структурованої та неструктурованої інформації, математичні моделі, методи інтелектуального аналізу даних та штучний інтелект.

Ключові слова: система підтримки прийняття рішень, управлінські рішення, інформаційна технологія, оброблення даних, повоєнне відновлення.

1. ВСТУП

Прийняття ефективних управлінських рішень у період війни відбувається в умовах постійної зміни зовнішнього середовища, браку ресурсів, неповноти інформації про ситуацію та необхідності швидкого вибору такого варіанту дій, який ураховує наявні ризики й можливі, у тому числі довгострокові, наслідки. Системна невизначеність є характерною рисою не лише воєнного стану – вона не менш важлива і в період повоєнного відновлення.

2. ОПИС МЕТОДИКИ

Для прогнозування соціально-економічних процесів розроблено багато різних методик застосування інтелектуального аналізу даних, математичного моделювання у системах підтримки прийняття рішень. Ці підходи представлені в роботах багатьох вітчизняних та закордонних вчених [1–4], це:

- регресійні моделі різних типів.
- моделі із включенням трендової складової та урахуванням залишків (різниця між реальним та прогнозним значенням), що включаються до моделі у вигляді ковзного середнього, за умови, що між залишками та цільовою змінною є кореляція.
- моделі класу експоненційного згладжування – Хольта, Брауна, Вінтерса (з адитивною або мультиплікативною сезонною складовою), з урахування дампа тренда та інші модифікації.
- нейронні мережі.
- узагальнені лінійні моделі.
- ймовірнісні моделі.
- нечіткі методи.

Архітектура пропонованої системи призначена для вирішення задач прийняття рішень в умовах невизначеності, тому обов'язковим компонентом є керуюча частина із зворотнім зв'язком, показана на рисунку 1.

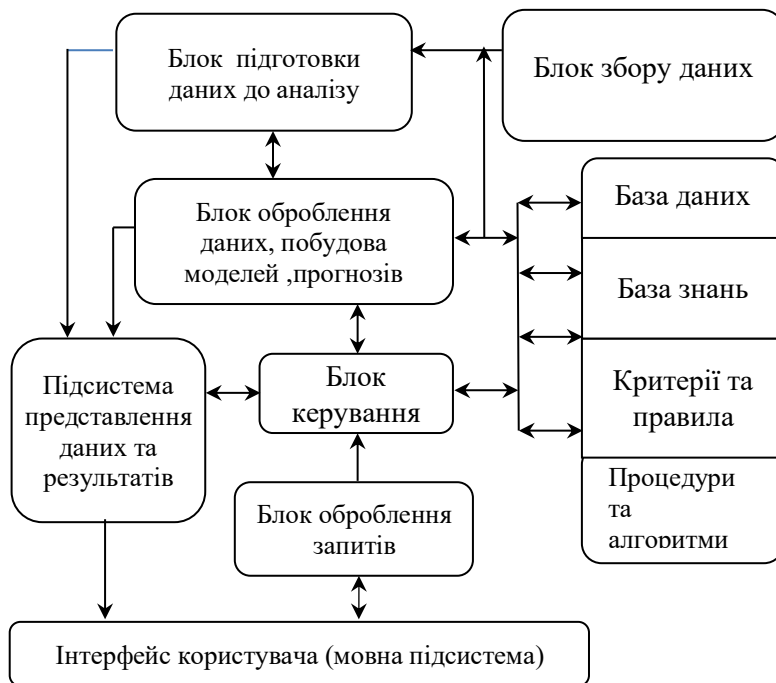


Рисунок 1. Архітектура системи підтримки прийняття рішень [4]

Як видно з наведеної схеми, підвищити якість прогнозування можна шляхом використання адаптивної моделювальної схеми та розширення функцій блоку попередньої обробки даних. Запропонований підхід було протестовано на кількох практичних задачах.

Зокрема, для аналізу перспектив повоєнного відновлення було застосовано запропоновану методику, реалізовану у кілька послідовних етапів. На першому етапі визначено просторові межі дослідження (територія України на макро-, мезо- та мікрорівнях), функціональні межі (галузі національної економіки та фінансовий сектор відповідно до класифікації видів економічної діяльності), а також часовий період дослідження (2015-2024 рр.). Для аналізу були використані відкриті дані: Порталу відкритих даних [5], Державної служби статистики України [6], Національного банку України [7], Міністерства фінансів України [8] та інших джерел.

На другому етапі проведено SWOT-, PEST- та SPACE-аналізи, а також кореляційний аналіз, доповнений матрицею впливу. Це дозволило поєднати кількісні показники – зокрема макроекономічні індикатори, кореляційні залежності та результати експертних оцінок. На основі матриці впливу сформовано мережевий граф, де вузли представляють галузі національної економіки та макропоказники, а ребра – зв'язки між ними із силою впливу понад 1,0. У результаті визначено ключові сфери, що формують стійкість системи – фінансову та транспортну, які потребують постійного моніторингу та управління [4].

На третьому етапі розроблено базовий, оптимістичний і песимістичний сценарії розвитку. Зокрема, базовий сценарій, наведений у таблиці 1, передбачає ведення бойових дій у вузькому коридорі, повільне відновлення інфраструктури та помірний темп реформ. За умови збереження поточного рівня міжнародної підтримки та стабільної глобальної

кон'юнктури прогнозується щорічне зростання ВВП на 3 %, інфляція на рівні 8-10 %, поступове зниження безробіття до 12 % і курс гривні в межах 40-45 грн/\$ [4].

Таблиця 1. Базовий сценарій [4]

Припущення	активні бойові дії припинені, інфраструктура відновлюється, хоча й повільно. реформи тривають у нинішньому темпі, але є пригальмовування із-за бюрократизації впровадження змін міжнародна допомога на високому рівні (на рівні траншів, отримуваних у період військових дій), потік прямих іноземних інвестицій помірний. На світовому ринку сприятлива цінова кон'юнктура на головні експортні товари – зерно та метал (ціни на зерно й метал залишаються близько до середньорічних 2022-2023 р. рівнів).
Макроекономічні показники (щорічно)	темپ приросту ВВП - 3 % індекс споживчих цін не перевищує: 10-15 % безробіття знижується з 15 % до 12 % дефіцит бюджету 15-20 % ВВП курс національної валюти стійкий – коливається в діапазоні 40-45 грн/дол.
Секторальні ефекти	промисловість: відновлюється, виробництва зростає на 60 % або на 20% від довоєнного рівня, зростання не менше 2 % щорічно. урожайність сільськогосподарських культур стало висока, експорт +1,5-2 % щорічно. ІТ-сектор зростає на 5-8 % на рік, сфера будівництва та реконструкції – відбудова 30 % зруйнованих об'єктів до 2027 р.

Побудовані сценарії дозволяють узагальнити впливові фактори, визначити їхню пріоритетність і структуру взаємозв'язків. На четвертому етапі здійснюється моделювання нелінійних залежностей між ключовими параметрами досліджуваного процесу. Для формування топології мережі Байєса застосовано метод Greedy Thick Thinning, а для обчислення ймовірнісних значень у мережі – алгоритм K2. Попередньо дані були підготовлені: інтервальні показники перетворено на бінарні – такі, що відображають два стани змінної (зростання або спад) у момент часу t порівняно з попереднім періодом $t-1$. Відповідний фрагмент побудованої мережі Байєса наведено на рисунку 2 [4].

Отримана мережа Байєса демонструє класифікаційну похибку на рівні 15%, тобто у 85% випадків вона коректно відтворює інвестиційний сценарій – зміну валового нагромадження капіталу залежно від станів інших вузлів мережі, що представляють макроекономічні показники [4].

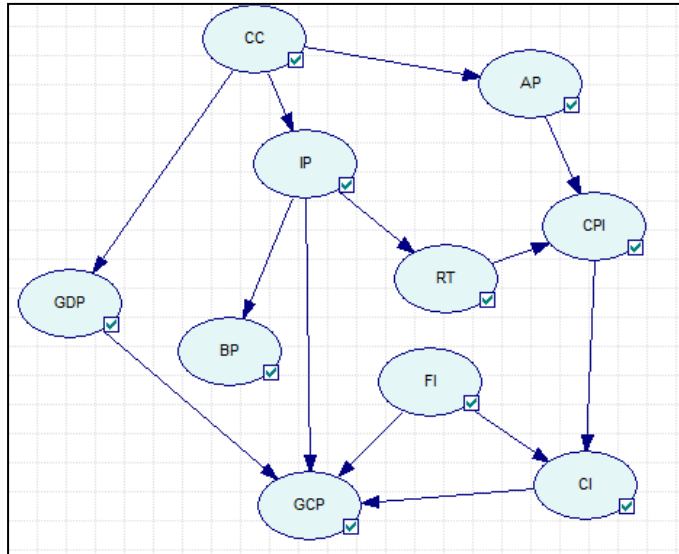


Рисунок 2. Фрагмент топології мережі Байєса, що моделює базовий сценарій повоєнного відновлення [4]

Загалом моделювання сценаріїв на основі мереж Байєса показало, що значне зростання цільової змінної GCP (валового нагромадження капіталу) забезпечується насамперед підвищенням значень таких показників, як GDP (валовий внутрішній продукт), RT (оборот роздрібною торгівлі), IP (обсяг реалізованої промислової продукції, товарів і послуг) та AP (виробництво сільськогосподарської продукції) [4].

На п'ятому етапі, спираючись на побудовані сценарії, здійснюється прогнозування різних цільових змінних на різних часових горизонтах. Наприклад, у межах «інвестиційного» сценарію для довгострокового прогнозу валового нагромадження основного капіталу було застосовано методи штучного інтелекту, зокрема нейронні мережі. Побудована модель характеризується високою якістю: коефіцієнт детермінації R^2 становить 0,93, а середня абсолютна процентна похибка (MAPE) на тренувальному наборі даних – 7,39% [4].

3. АНАЛІЗ РЕЗУЛЬТАТІВ

На рисунку 3 наведено результати довгострокового прогнозування ВВП та валового нагромадження основного капіталу за різними сценаріями на період до 2040 року. Для побудови моделі нейронної мережі було використано програмне забезпечення SAS Enterprise Miner з такими налаштуваннями: вхідний шар включає всі змінні, застосовані при побудові сценаріїв, зокрема лагові змінні першого порядку; один прихований шар містить п'ять вузлів (визначено аналітичним шляхом); для перетворення вхідних даних використовується функція softmax; кількість ітерацій навчання становить двадцять [4].

Основний результат проведеного дослідження полягає в тому, що відновлення економіки та позитивна динаміка до 2040 року можуть бути забезпечені за умови стабільного зростання інвестицій та збереження співвідношення між промисловістю та сільським господарством на довоєнному рівні.

Висока якість отриманих результатів прогнозного моделювання підтверджує перспективність застосування запропонованої методики для аналізу різних соціально-економічних процесів. Крім того, цей підхід має переваги в автоматизації ключових етапів попередньої обробки даних і прогнозного моделювання.

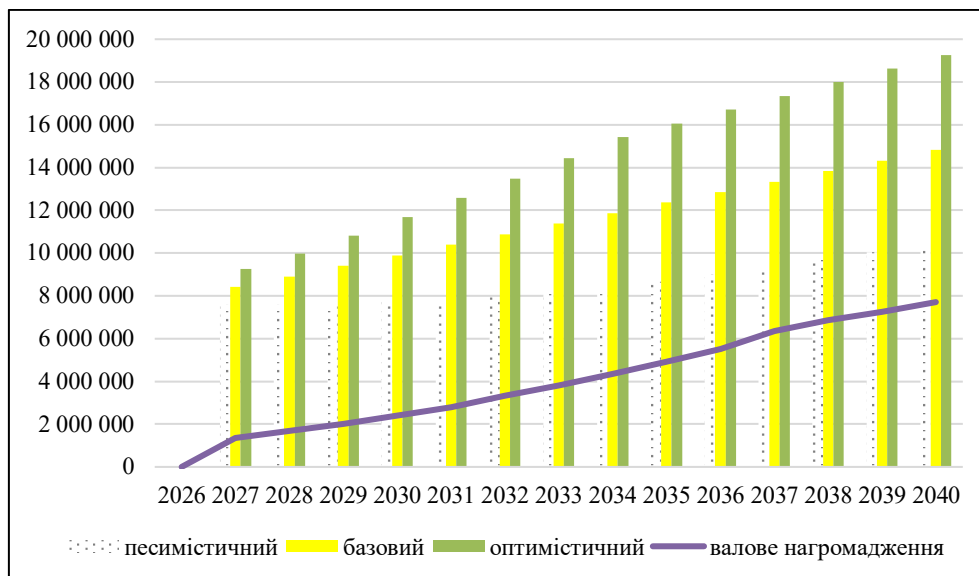


Рисунок 3. Результати довгострокового прогнозування макроекономічної динаміки за різними сценаріями інвестиційного забезпечення до 2040 року [4]

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Trofymchuk O. M., Bidiuk P. I., Prosyankina-Zharova T. I., Terentiev O. M. Decision Support Systems for Modeling, Forecasting and Risk Estimation (Monograph). – LAP LAMBERT Academic Publishing, 2019. – 176 p. – ISBN-13: 978-3-659-61214-5. – ISBN-10: 3659612146. – EAN: 9783659612145.
2. Trofymchuk O. M., Bidiuk P.I., Terentiev O.M., Klymenko V.I. The methodology for adaptive modeling and forecasting nonlinear and nonstationary processes // Міжнародний науково-технічний журнал Проблеми керування та інформатики. – 2024, № 1. – 63-79 с. – ISSN 2786-6491 – <https://jais.net.ua/index.php/files/article/view/216/302> (дата звернення: 20.11.2025).
3. Bidiuk P.I., Korshevnyuk L.O., Gozhyi A.P., Kalinina I.O., Prosyankina-Zharova T.I., Terentiev O.M. Modeling and forecasting financial and economic processes with decision support system // KPI Science News. – 2019. – No 5-6. – 7-17 с. – ISSN 2617-5509. – DOI: 10.20535/kpi-sn.2019.5-6.176835
4. Трофимчук О.М., Бідюк П.І., Терент'єв О.М., Просянкін-Жарова Т.І. Математичне моделювання, інтелектуальний аналіз даних та штучний інтелект для підтримки прийняття рішень з повоєнного відновлення // Екологічна безпека та природокористування. – 2025. – вип. 3 (55). – с. 33-49. – ISSN: 2411-4049. – <https://doi.org/10.32347/2411-4049.2025.3.33-49> (дата звернення: 20.11.2025).
5. Набори даних. <https://data.gov.ua/dataset>. <https://data.gov.ua>. (дата звернення: 20.11.2025)
6. Публікації. <https://www.ukrstat.gov.ua> Державної служби статистики України. Офіційний веб-сайт. (дата звернення: 20.11.2025)
7. Набори даних. Національний банк України: <https://bank.gov.ua/ua/statistic/> (дата звернення: 20.11.2025)
8. Набори даних. Міністерство фінансів України. Статистичний збірник. <https://mof.gov.ua/uk/statistichnij-zbirnik> (дата звернення: 20.11.2025)

СПОТОВА ТА Ф'ЮЧЕРСНА ТОРГІВЛЯ ЯК ЗАСІБ ХЕДЖУВАННЯ ДЛЯ АЛГОРИТМІЧНОЇ ТОРГІВЛІ НА КРИПТОБІРЖІ У СХІДНОЄВРОПЕЙСЬКОМУ РЕГІОНІ

Загарук А.Я.¹, Касьянов П.О.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ andrijko.zagaruk@gmail.com

Метою даної роботи є дослідження того, як алгоритмічні методи аналізу волатильності та марківського моделювання можуть бути використані для підтримки прийняття рішень у спотовій та ф'ючерсній торгівлі на криптовалютному ринку. Робота включає побудову системи, яка автоматично оцінює поточний ринковий стан, формує матрицю переходів волатильності та визначає оптимальні дії для трейдера за допомогою алгоритму Value Iteration. На відміну від традиційних підходів, що базуються на прогнозуванні цін, дана система використовує моделювання ринкових режимів і стохастичні процеси, що дозволяє отримати стратегію хеджування, стійку до високої волатильності криптовалют. У дослідженні застосовано історичні дані криптоактивів, отримані з біржі Kraken, та проведено класифікацію ринкових режимів на основі реалізованої волатильності. Результатом є оптимальна політика для спотових і ф'ючерсних позицій, яка демонструє переваги марківського підходу у прийнятті торгових рішень в умовах нестабільного криптовалютного ринку.

Ключові слова: криптовалюта, спотова торгівля, ф'ючерси, хеджування, волатильність

1. ВСТУП

Останніми роками криптовалютний ринок демонструє стрімке зростання складності та масштабу, що зумовлено розвитком фінансових технологій, появою нових біржових інструментів та збільшенням обсягу доступних ринкових даних. Криптовалютний ринок як високоволатильна та децентралізована система формується під впливом таких факторів, як глобальні макроекономічні тенденції, регуляторні рішення, активність бірж і загальна спекулятивність трейдерів. Особливу актуальність у цьому контексті набувають питання керування ризиками та побудови ефективних стратегій хеджування, оскільки швидкі й різкі коливання цін значно ускладнюють процес прийняття рішень.

Через нестабільність криптовалют традиційні інструменти аналізу ринку часто виявляються недостатньо ефективними. Саме тому все більшого значення набувають алгоритмічні та математичні методи, що ґрунтуються на моделюванні ринкових режимів, оцінці волатильності та ймовірнісних переходах між станами ринку. Сучасні підходи, включаючи аналіз реалізованої волатильності, кластеризацію ринкових умов, марківські процеси та методи динамічного програмування, дозволяють будувати торгові стратегії, здатні працювати в умовах постійної невизначеності крипторинку.

Метою даної роботи є дослідження застосування методів аналізу волатильності та марківського моделювання для формування оптимальної стратегії хеджування у спотовій та

ф'ючерській торгівлі на криптовалютному ринку Східноєвропейського регіону. У роботі проведено класифікацію ринкових режимів для основних криптовалют, побудовано матрицю переходів волатильності та визначено оптимальні торгові дії на основі алгоритму Value Iteration.

2. МОДЕЛІ ДЛЯ ПРОГНОЗУВАННЯ ЦІН АКЦІЙ

У межах даної роботи розглянуто моделі та методи, які дозволяють аналізувати ринкові умови криптовалют та формувати оптимальну стратегію хеджування для спотової та ф'ючерсної торгівлі. Оскільки криптовалютний ринок характеризується високою мінливістю та нерегулярними коливаннями, ключовим елементом дослідження стала оцінка реалізованої волатильності. Для цього було використано логарифмічні доходності, що визначаються як

$$r_t = \ln(P_t) - \ln(P_{t-1}),$$

де:

P_t – ціна активу у момент часу t ,

P_{t-1} – ціна активу у попередній момент часу,

R_t – логарифмічна доходність, яка дозволяє нормалізувати зміну ціни та зменшити вплив різких стрибків.

Ковзний рівень волатильності формувався на основі послідовності доходностей як середнє квадратичне відхилення за певний період, що дозволяє оцінити ступінь ринкової мінливості.

Після отримання ряду волатильності виникла необхідність у введенні дискретних режимів ринку. Для цього було застосовано класифікацію на три рівні волатильності: низьку, середню та високу. Поділ здійснювався за квантильним підходом, який враховує особливості конкретного активу та масштабу його історичних коливань. У деяких випадках альтернативно використовувався метод кластеризації KMeans, що дозволяв автоматично групувати значення волатильності відповідно до їх подібності. У результаті кожна точка історичного ряду була віднесена до одного з ринкових режимів.

Оскільки волатильність криптовалют має випадкову природу, подальший аналіз проводився на основі марківської моделі. Було сформовано матрицю переходів між станами ринку, що дозволило оцінити ймовірності переходів від одного режиму до іншого. Марківське припущення описується залежністю:

$$P(V_{t+1} = j \mid V_t = i),$$

де

V_t – режим волатильності у момент часу t ,

i, j – можливі стани ринку (низька, середня або висока волатильність),

$P(\cdot)$ – ймовірність переходу між станами.

Це припущення означає, що майбутній стан залежить лише від поточного, що суттєво спрощує моделювання ринкової динаміки та дозволяє формувати стохастичну поведінку ринку.

Для побудови логіки прийняття рішень було розширено простір станів шляхом поєднання ринкового режиму та типу позиції — спотової або ф'ючерсної. Для кожної з таких комбінацій визначалися можливі дії: купити, продати, утримувати або хеджувати позицію. Винагороди для цих дій встановлювалися вручну на основі принципів керування ризиками, з урахуванням того, що у високій волатильності доцільніше зменшувати ризики, тоді як у низькій волатильності ринок створює умови для акумулювання позицій.

Оптимальна стратегія визначалася методом Value Iteration, який дозволяє знайти політику, що максимізує сумарну винагороду з урахуванням ймовірностей переходу між станами та довгострокових наслідків кожної дії. Алгоритм ґрунтується на рівнянні оптимальності Беллмана:

$$V(s) = \max_a \left[R(s, a) + \gamma \sum_{s'} P(s'|s, a) \cdot V(s') \right],$$

де

$V(s)$ – значення стану s ,

$R(s, a)$ – негайна винагорода за виконання дії a у стані s ,

γ – коефіцієнт дисконтування, що визначає вагу майбутніх вигод,

$P(s'|s, a)$ – ймовірність переходу до стану s' після виконання дії a ,

Max_a – вибір найкращої дії, яка забезпечує максимальну очікувану вигоду.

У підсумку отримана модель здатна аналізувати волатильні ринкові режими, оцінювати ймовірності переходів між ними та пропонувати оптимальні торгові дії залежно від поточної ринкової ситуації.

3. МОДЕЛЮВАННЯ ТА РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

У процесі дослідження було проведено моделювання ринкової поведінки криптовалют на основі історичних даних та побудовано алгоритм, здатний визначати оптимальну стратегію хеджування для спотової та ф'ючерсної торгівлі. На першому етапі було здійснено збір та попередню обробку даних основних криптоактивів. Дані отримувалися у вигляді цін закриття, після чого обчислювалися логарифмічні доходності та реалізована волатильність. Паралельно проводилася візуалізація динаміки цін та волатильності (рис. 1), що дозволило оцінити характер змін і виявити періоди підвищеної ринкової напруги.

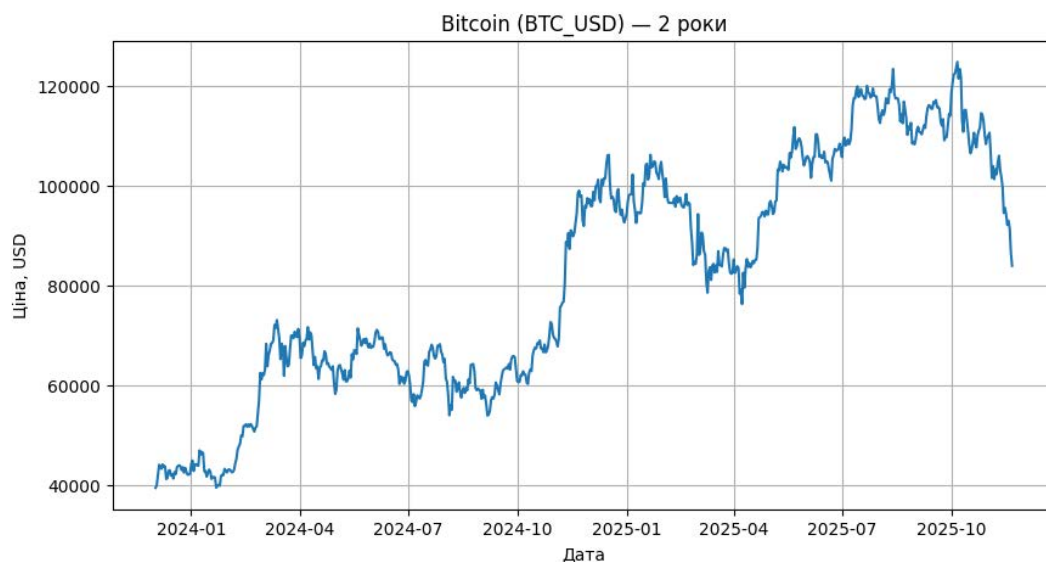


Рисунок 1. Динаміка змін цін криптовалюти Bitcoin

Після формування ряду волатильності було виконано дискретизацію ринку на три режими: низький, середній та високий. Цей процес дозволив перетворити неперервний часовий ряд у послідовність станів, придатних для марківського моделювання. Для кожного активу було побудовано власну матрицю переходів (рис. 2), між режимами, що дала змогу оцінити ймовірність зміни ринкових умов у майбутніх періодах. Результати показали, що

криптовалютні ринки демонструють значну інерційність: перебування у високій волатильності часто супроводжується збереженням цього стану протягом кількох днів, тоді як перехід із низької волатильності у високу зазвичай має низьку ймовірність і відбувається у вигляді різких ринкових сплесків.

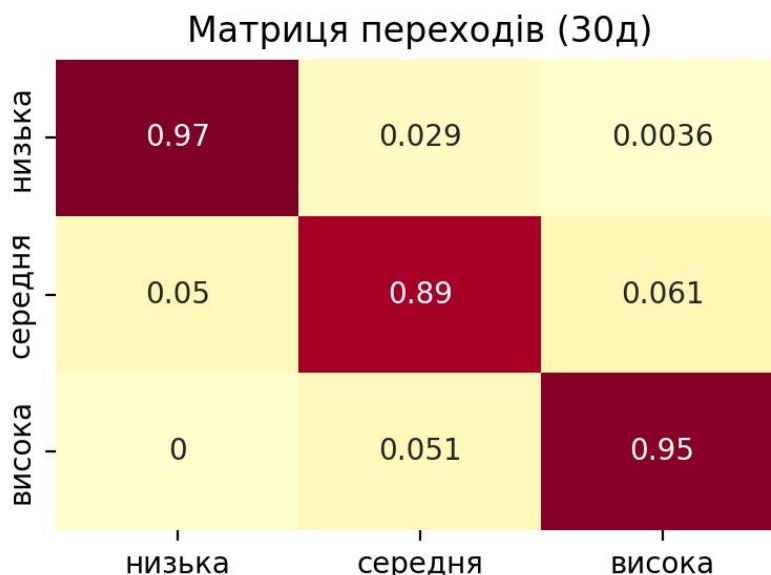


Рисунок 2. Матриця переходів

На основі марківської моделі було сформовано розширений простір станів, що поєднував ринковий режим та тип поточної позиції (спотова чи ф'ючерсна). Для кожної пари «стан–дія» визначалися очікувані винагороди, що моделювали прибутковість та ризиковість торгових рішень у різних ринкових умовах. Таким чином було побудовано середовище для роботи алгоритму Value Iteration.

Запуск алгоритму дозволив отримати оптимальні значення станів та відповідну політику дій. Після кількох десятків ітерацій значення станів збіглися, що свідчить про досягнення стабільної оптимальної стратегії. Результати показали закономірні властивості моделі: у фазах високої волатильності оптимальною дією найчастіше є хеджування або закриття позицій, тоді як у низькій волатильності система рекомендує акумулювати або утримувати позиції. Середня волатильність формує змішану політику, яка залежить від попереднього стану та типу інструмента.

Таблиця 1. Значення $V(s)$ для Bitcoin

	Спот	Ф'ючерс
Низька	28.20	21.25
Середня	21.39	25.12
Висока	27.36	28.51

На основі результатів алгоритму Value Iteration було отримано значення функції стану $V(s)$, які відображено на таблиці 1. Вони дозволяють оцінити очікувану довгострокову вигоду для кожної комбінації ринкового режиму та типу позиції. Для Bitcoin найвищі значення очікуваної винагороди спостерігаються у станах високої волатильності для ф'ючерсних позицій ($V = 28.51$) та у низькій волатильності для спотових позицій ($V = 28.20$). Це свідчить

про те, що за умов сильних коливань ринку ф'ючерси забезпечують потенційно кращий результат – здебільшого завдяки можливості відкривати короткі позиції або здійснювати більш гнучке хеджування. Натомість у низьковолатильних умовах ринок є більш передбачуваним, що робить спотові стратегії оптимальними.

Аналіз оптимальної політики підтверджує ці спостереження: у режимі низької волатильності модель схиляється до дій, спрямованих на утримання або нарощення спотових позицій, тоді як у режимах середньої та високої волатильності алгоритм частіше обирає хеджування або роботу з ф'ючерсами. Це узгоджується зі значеннями $V(s)$, де ф'ючерсні позиції отримують найвищу оцінку саме за умов підвищеної ринкової турбулентності. Таким чином, результати моделювання свідчать, що запропонована система прийняття рішень коректно відображає поведінкові особливості криптовалютного ринку та формує збалансовану стратегію хеджування, адаптовану до різних режимів волатильності.

4. ВИСНОВКИ

У результаті проведеного дослідження було розроблено та протестовано підхід до прийняття торгових рішень для спотової та ф'ючерсної торгівлі на криптовалютному ринку, заснований на аналізі волатильності та марківському моделюванні. Використання логарифмічних доходностей та реалізованої волатильності дозволило отримати інформативне представлення ринкових умов, а класифікація на ринкові режими забезпечила можливість формального опису динаміки криптовалютного середовища.

Побудовані матриці переходів підтвердили, що криптовалютні ринки мають високу інерційність, особливо в періоди підвищеної волатильності. Це свідчить про те, що ринкові стани не змінюються хаотично, а демонструють певні закономірності, які можна використовувати під час формування стратегій хеджування.

Застосування алгоритму Value Iteration дозволило визначити оптимальні дії для кожного поєднання ринкового режиму та типу позиції. Отримана політика виявила логічну структуру: у спокійних ринкових умовах доцільніше утримувати або нарощувати спотові позиції, тоді як у періоди підвищеної волатильності алгоритм зміщує акцент на ф'ючерси та хеджування ризиків. Таким чином, оптимальні рішення формуються не на основі прогнозу ціни, а на основі очікуваних умов ринку та ймовірностей їх зміни.

Підсумовуючи, запропонований підхід підтверджує ефективність марківських моделей і методів динамічного програмування у задачах підтримки прийняття рішень на криптовалютному ринку. Отримані результати демонструють, що впровадження таких алгоритмів може стати надійним інструментом хеджування та оптимізації торгових стратегій в умовах високої ринкової невизначеності. Модель є масштабованою та може бути розширена на інші активи, більш високочастотні дані або складніші структури станів, що відкриває перспективи для подальших досліджень та вдосконалення алгоритмічних торгових систем.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. J. Wooldridge. *Introductory Econometrics: A Modern Approach*. Boston, MA: Cengage Learning, 2020. 912 p.
2. R. Sutton, A. Barto. *Reinforcement Learning: An Introduction*. Cambridge, MA: MIT Press, 2018. 552 p.
3. R. Hyndman, G. Athanasopoulos. *Forecasting: Principles and Practice*. Melbourne: OTexts, 2021. 412 p.
4. S. Shreve. *Stochastic Calculus for Finance II: Continuous-Time Models*. New York: Springer, 2004. 550 p.

КЛАСТЕРИЗАЦІЯ РЕГІОНІВ УКРАЇНИ ЗА ДИНАМІКОЮ ПОВІТРЯНИХ ТРИВОГ МЕТОДОМ K-СЕРЕДНІХ ДЛЯ ЧАСОВИХ РЯДІВ (DTW/SOFT-DTW) У КОНТЕКСТІ РЕЛОКАЦІЙ БІЗНЕСУ

Кабанова М.І.¹, Кузнєцова Н.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ kabanova.marina@lil.kpi.ua [0009-0004-0983-0710],

² n.kuznietsova@kpi.ua

Запропоновано просторово-часовий підхід до аналізу безпекової ситуації та напрямів релокацій бізнесу в Україні. Регіони кластеризовано за формою динаміки повітряних тривог (08.2024–08.2025) методом k-середніх для часових рядів із відстанями DTW/Soft-DTW. Найкращу конфігурацію (Soft-DTW) використано для побудови профілів кластерів та їх зіставлення з показниками міжрегіональних переміщень бізнесу (інтенсивність потоків, частки притоків/відтоків, MER, MEI). Показано відмінності профілів. Отримані результати можуть бути використані для ідентифікації регіонів, привабливих для релокації бізнесу з урахуванням ризиків, а також для підтримки рішень міських і регіональних адміністрацій щодо посилення інвестиційної та безпекової привабливості.

Ключові слова: кластеризація часових рядів, DTW, Soft-DTW, Silhouette Score, повітряні тривоги, релокація бізнесу, MER, MEI, навчання без учителя, Python, економічна аналітика.

1. ВСТУП

Повномасштабне вторгнення росії спричинило нерівномірність економічної активності в регіонах України, створивши нові виклики для бізнесу, домогосподарств і державного управління. Частота, тривалість та сезонність повітряних тривог впливають на операційні процеси, освітній процес, логістику, а також на рішення бізнесу щодо релокації у безпечніші регіони. У таких умовах виникає потреба у системному кількісному аналізі динаміки повітряних тривог не лише як безпекового показника, а й як фактора регіональної економічної стійкості та інвестиційної привабливості. Аналіз безпекових профілів регіонів та міжрегіональних потоків релокації бізнесу дозволяє ідентифікувати зони підвищеного ризику, перспективні напрямки переміщення та потенційні місця зростання.

2. ПОСТАНОВКА ЗАДАЧІ

Метою дослідження є ідентифікація регіонів України з різними профілями безпекової ситуації за динамікою повітряних тривог і визначення на цій основі найбільш привабливих напрямків для релокації бізнесу. Потрібно здійснити аналіз та групування регіонів за показниками повітряних тривог за період серпень 2024 – серпень 2025 методом k-середніх для часових рядів із відстанями DTW та Soft-DTW та зіставити із агрегованими показниками міжрегіональних переміщень бізнесу.

3. МЕТОДИ

3.1. Обґрунтування вибору методу кластеризації часових рядів

Для часових даних стандартна евклідова відстань не завжди коректно описує подібність форм, адже порівнює точки на основі їх положення в часовій послідовності й не враховує зсуви у часі [1, 2, 3]. Тому у даному дослідженні використовуються методи, засновані на відстані між формами (distance-based methods) [1]. Алгоритм динамічного трансформування часу (Dynamic Time Warping, DTW) відомий завдяки надзвичайній ефективності для виявлення подібних форм з різними фазами [1, 2, 3]. Він мінімізує вплив зсуву та спотворення в часі, дозволяючи «розтягування» чи «стискання» часової осі, щоб оптимально вирівняти послідовності [4]. У свою чергу, Soft-DTW – це диференційована версія DTW, що використовує операцію «soft-min» замість жорсткого мінімуму, згладжуючи поверхню функції втрат, що дає стабільніші барицентри [5].

3.2. Опис та підготовка навчальних даних

Джерела: ALERTS.IN.UA (кількість та тривалість тривог, кількість зафіксованих вибухів) [6], ACLED (атаки по інфраструктурі) [7]. Економічні показники змін місцезнаходження (релокації) бізнесу надані командою Опендатабот [8]. Дослідження не включало Автономну Республіку Крим, місто Севастополь, та Луганську область, так як із причин тимчасової окупації цих регіонів дані частково відсутні, що може вплинути на якість результатів.

Період зібраних даних: серпень 2024 – серпень 2025; 24 області України, включаючи м. Київ. Ознаки для кластеризації: середня кількість тривог на день, середні години тривог на день, години на одну тривогу, частка нічних тривог. Економічні дані: міжрегіональні переміщення бізнесу (прибули/вибули та похідні метрики).

Щоб зосередитися на формі динаміки, незалежно від абсолютних масштабів регіонів, часові ряди (24×13×4) було стандартизовано за допомогою z-score окремо для кожної ознаки в межах регіону.

3.3. Кластеризація

Кластеризацію виконано за допомогою мультимірного алгоритму TimeSeriesKMeans із метриками DTW (metric="dtw") та Soft-DTW (metric="softdtw"), реалізованого в бібліотеці tslearn Python [9].

Формальне визначення відстані DTW задається формулою (1) [2]:

$$DTW_q(A, B) = \min_{\pi \in P} \left(\sum_{(i,j) \in \pi} d(a_i, b_j)^q \right)^{\frac{1}{q}}, \quad (1)$$

де P – множина допустимих шляхів вирівнювання; $d(a_i, b_j)$ – відстань між точками a_i з ряду A та b_j з ряду B , та сума береться по всіх точках (i, j) в оптимальному шляху вирівнювання з P .

Soft-DTW як згладжений аналог DTW визначається за формулою (2) [10]:

$$\text{softDTW}_\gamma(A, B) = \min_\gamma \sum_{(i,j) \in \pi} d(a_i, b_j)^2, \quad (2)$$

де $\min^\gamma(a_i, \dots, a_n) = -\gamma \log \sum_i e^{-a_i/\gamma}$, а γ – коефіцієнт згладжування ($\gamma > 0$).

3.4. Метрика оцінювання результатів роботи

Якість кластеризації оцінюється за допомогою коефіцієнта силуету (Silhouette Score, SS), який вимірює відстань між кластерами без потреби у «ground truth». Для кожної точки даних з усього набору обчислюється за наступною формулою (3) [11]:

$$s(i) = \frac{b(i) - a(i)}{\max\{b(i), a(i)\}} \in [-1, 1] \quad (3)$$

де $a(i)$ – середня відстань від точки i до власного кластера, $b(i)$ – мінімальна середня відстань до точок іншого кластера.

Значення, близьке до 1, свідчить про добре відокремлений кластер, тоді як значення, близьке до -1, – про погіршення кластеризації. Загальний бал SS обчислюється як середнє значення коефіцієнтів силуету по всіх точках даних.

Розрахунок коефіцієнту виконано на попередньо згенерованій матриці відстаней DTW/Soft-DTW, побудованій із тими самими параметрами, що й у моделі кластеризації, із використанням функції `silhouette_score (metric='precomputed')` бібліотеки `scikit-learn` у Python [12].

3.5. Підбір гіперпараметрів (Grid Search)

Початковий аналіз коефіцієнту силуету для кількості кластерів k у межах від 3 до 10 показав, що обидва методи потребують подальшої оптимізації, оскільки Soft-DTW досяг лише $SS=0.35$ (при $k=6$), а DTW – $SS=0.2$ (при $k=7$). Виконавши Grid Search, фінальні конфігурації гіперпараметрів для кожного алгоритму наступні:

- DTW із $SS=0.2175$: '*n_clusters*': 6, '*global_constraint*': '*sakoe chiba*', '*radius or slope*': 2, '*n_init*': 10, '*max_iter_barycenter*': 10.
- Soft-DTW із $SS=0.4873$: '*n_clusters*': 8, '*gamma*': 3.0, '*n_init*': 3, '*max_iter_barycenter*': 20.

3.6. Економічні метрики

Для зіставлення з безпековими профілями кластерів використано такі агреговані показники.

- Інтенсивність потоків на регіон (*flow_per_region*):

$$flow_per_region_{c,t} = \frac{gross_sum_{c,t}}{n_regions_c}$$

де c – кластер; t – місяць; *gross_sum* – сума прибулих та вибулих компаній у кластері; $n_regions_c$ – кількість регіонів кластеру.

- Частка притоків/відтоків кластера:

$$in_share_{c,t} = \frac{D_{c,t}}{D_{c,t} + O_{c,t}}, out_share_{c,t} = \frac{O_{c,t}}{D_{c,t} + O_{c,t}}$$

де $D_{c,t}$ – кількість прибулих компаній у кластері c за місяць t ; $O_{c,t}$ – кількість вибулих компаній у кластері c за місяць t .

- Коефіцієнт ефективності міграції (Migration Efficiency Ratio (MER))[13].

$$MER_{c,t} = 100 \times \frac{D_{c,t} - O_{c,t}}{D_{c,t} + O_{c,t}}$$

- Індекс ефективності міграції (Migration Effectiveness Index (MEI)).

MEI може приймати значення від 0 до 100, де низькі значення означають, що міжрегіональні потоки більш збалансовані [14].

$$MEI_{c,t} = 100 \times \frac{\sum_c |D_{c,t} - O_{c,t}|}{\sum_c (D_{c,t} + O_{c,t})}$$

Для узагальнення подано медіани та зважені середні (вага = *gross_sum* для часток/MEI/MER).

4. РЕЗУЛЬТАТИ КЛАСТЕРИЗАЦІЇ

Хоча обидві метрики демонструють домінування одного із кластерів за чисельністю регіонів (рис. 1), Soft-DTW сформував більшу кількість кластерів, виділивши більш тонкі закономірності та відмінності груп. Вищий показник Soft-DTW ($SS=0.4873$) свідчить про кращу відокремленість форм часових рядів, тому надалі для подальшого дослідження та економічного аналізу використовуватиметься саме ця кластеризація як основна. Центроїди кластерів відображають спільні сезонні хвилі й фазові зсуви, що не фіксуються евклідовими метриками (рис. 2).

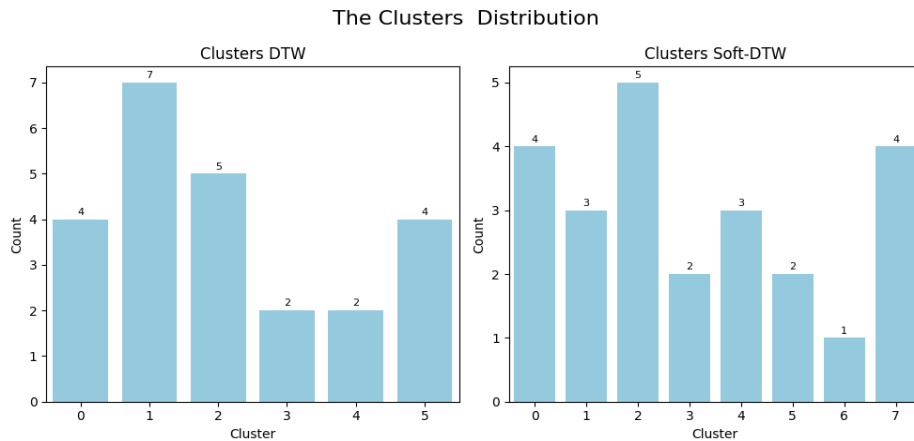


Рисунок 1. Розподіл кластерів (DTW та Soft-DTW)

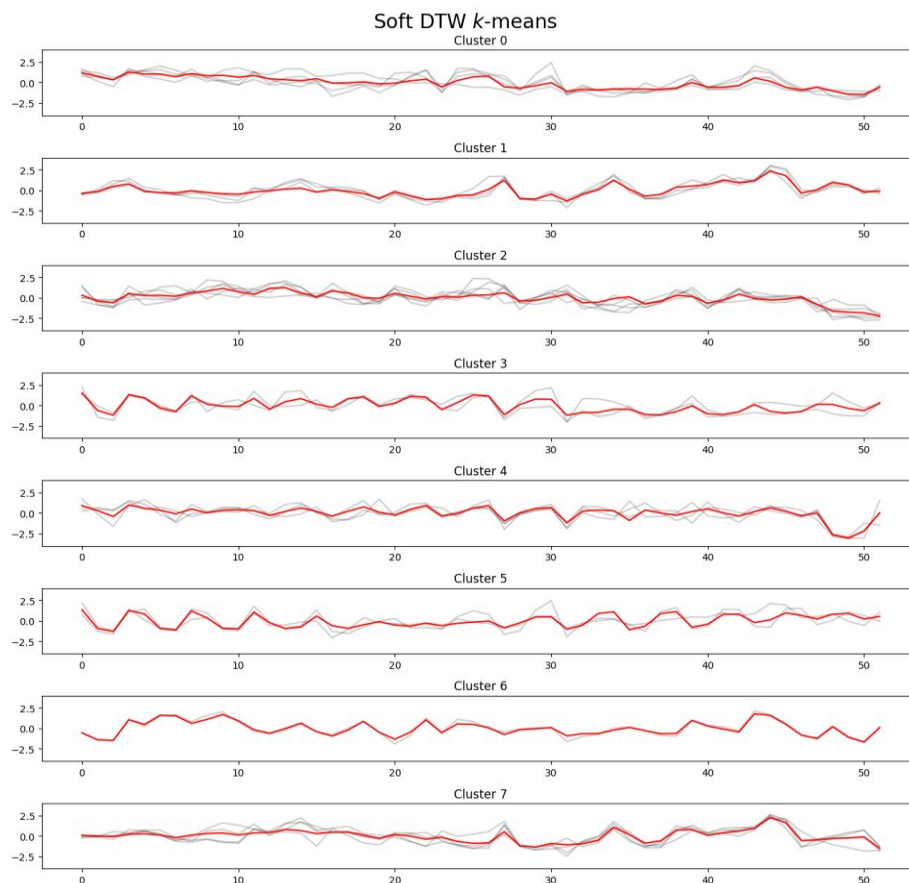


Рисунок 2. Центроїди кластерів Soft-DTW

Опишемо профілі підгруп областей, сформованих Soft-DTW ($k=8$) (табл. 1). Для формування чітких та економічно інтерпретованих профілів використовуємо агреговані ознаки: інтенсивність, тривалість та добову структуру тривог, наявність уражень критичної інфраструктури та прикордонний/прифронтовий статус (рис. 3).

Таблиця 1. Стислий профіль кластерів (на основі Soft-DTW, $k=8$)

Кластер	Області кластеру	Рівень ризику
K-0	Кіровоградська, Черкаська, Херсонська, Миколаївська	Середній: відсутні кордони з РФ/РБ, але помірно-висока денна частота тривог і помірні інфраструктурні удари (переважно житлова/енергетична); частка нічних тривог нижча за середню по країні.
K-1	Львівська, Волинська, Закарпатська	Низький: найменша тривалість та інтенсивність тривог; помірна частка нічних тривог; ураження рідкі, переважно житлова/енергетична; з Білоруссю межує лише Волинська.
K-2	Хмельницька, м. Київ, Житомирська, Київська, Вінницька	Середньо-високий: довші за середні епізоди та найвища частка нічних тривог; ураження переважно житлової інфраструктури; частина регіонів має кордон з Білоруссю.
K-3	Дніпропетровська, Чернігівська	Середньо-високий: підвищені частота і добові години тривог; тривалі одиничні події; нічні тривоги – нижче середнього; помітні удари по житловій/енергетичній; прикордонний/прифронтовий фактор.
K-4	Полтавська, Харківська, Сумська	Високий: комбінація частих і довгих тривог, значних уражень (житлова, енергетична, об'єкти освіти/медицини); прикордонний/прифронтовий фактор.
K-5	Донецька, Запорізька	Екстремально високий: найвищі частоти й добові години; тривалі одиничні події; занижена частка нічних тривог; дуже високі ураження (житлова/енергетична та об'єкти освіти/медицини); прифронтові регіони, частково прикордонні з РФ.
K-6	Одеська	Помірний: сукупне навантаження менше, ніж у східних кластерів, але тривоги наявні впродовж всієї доби; прибережний /портовий та інфраструктурний фактор ризику.
K-7	Чернівецька, Івано-Франківська, Рівненська, Тернопільська	Низький: низька інтенсивність та мінімальні ураження, але підвищена частка нічних тривог (вища за більшість кластерів); з Білоруссю межує лише Рівненська.

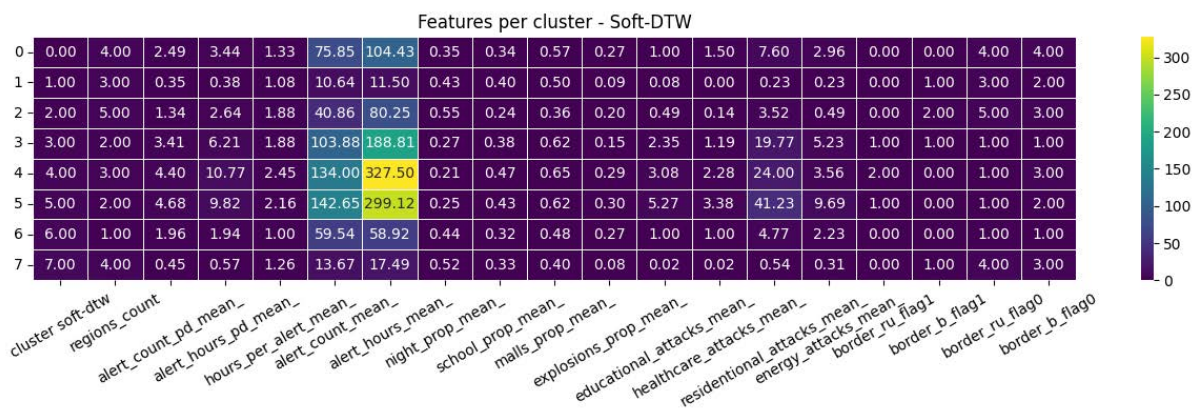


Рисунок 3. Теплова карта ознак Soft-DTW

5. ЕКОНОМІЧНІ ПРОФІЛІ КЛАСТЕРІВ

За період серпень 2024 – серпень 2025, найбільша інтенсивність релокацій бізнесу (flow_per_region) спостерігається у кластерах К-2 (175,2), К-6 (110,46) та К-3 (99,62). Найнижчі показники зафіксовано у кластерах К-7 (21,1) та К-0 (30,38) (табл. 2).

Частка притоків (in_share) є найбільшою у кластерів К-4 (0,61), далі К-7 (0,56) та К-1 разом із К-5 з (по 0,54). В той час, як у кластерів К-0 та К-2 (0,46-0,47) сильніше домінує відтік (out_share) за інші кластери.

За підрахунками, кластери К-1, К-4, К-5, К-7 мають позитивний зважений середній MER (від +8,78 до +22,28%), тобто чистий притік бізнесу. Натомість кластери К-0, К-2 та К-3 демонструють нетто-відтік, та К-6 близький до балансу (із середнім значенням MER -2,09% та позитивною медіаною +1,52%).

Зважений середній MEI є найвищим у кластерах К-4 (29,86) і К-5 (26,54), що вказує на більш виражену асиметрію потоків релокацій. Найнижчі показники ступеню незбалансованості потоків MEI у кластерах К-6 (9,33), К-3 (10,19) та К-2 (12,38).

Узагальнений розподіл показників in_per_region, out_per_region та flow_per_region за кластерами наведено на рис. 4.

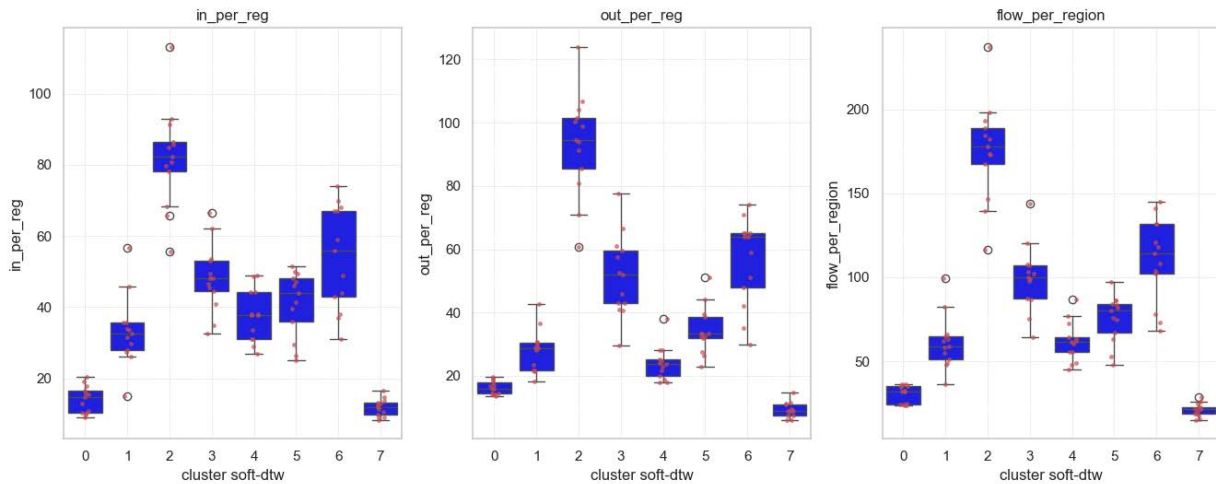


Рисунок 4. Коробкові графіки (boxplot) in_per_region, out_per_region, flow_per_region за кластерами

Таблиця 2. Медіани та зважені середні показників релокацій за кластерами

K	n_region	flow_per_region_avg	flow_per_region_median	in_share_wavg	in_share_median	out_share_wavg	out_share_median	MEI_wavg (0-100)	MEI_median (0-100)	MER_wavg (%)	MER_median (%)
0	4	30,38	32,00	0,46	0,44	0,54	0,56	24,30	24,00	-7,34	-12,06
1	3	61,05	58,67	0,54	0,55	0,46	0,45	19,53	18,88	8,78	9,68
2	5	175,20	177,60	0,47	0,46	0,53	0,54	12,38	11,50	-6,50	-7,16
3	2	99,62	100,0	0,48	0,46	0,52	0,54	10,19	9,03	-3,40	-7,14
4	3	61,56	61,67	0,61	0,61	0,39	0,39	29,86	29,87	22,28	22,58
5	2	75,50	80,00	0,54	0,55	0,46	0,45	26,54	27,17	8,81	10,67
6	1	110,46	114,0	0,49	0,51	0,51	0,49	9,33	10,26	-2,09	1,52
7	4	21,10	21,50	0,56	0,55	0,44	0,45	22,52	21,84	11,94	10,53

6. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Отримані безпекові профілі кластерів за Soft-DTW демонструють деяку узгодженість із економічною поведінкою бізнесу за показниками релокацій. Нічні та тривалі події створюють додатковий тиск на продуктивність і якість робочого дня наступної доби, тоді як часті денні тривоги призводять до прямих втрат операційних годин.

Отримані «тихі» західні кластери К-1 та К-7 із мінімальною інтенсивністю та тривалістю подій мають вищу частку притоку і стійкий плюс по MER. При цьому К-7 загалом менш інтенсивний за величиною потоків, з чого можна висунути припущення про вплив вищої частки нічних тривог. Кластер К-0 з частими, але короткими денними тривогами має нижчу інтенсивність релокацій і відносно вищий MEI (асиметрія потоків), порівняно з іншими кластерами. Тоді як, кластер К-3 із частими та тривалими денними тривогами демонструє вищу інтенсивність релокацій та від'ємний MER (нетто-відтік). В той час, як прикордонно/прифронтові кластери з вищою інтенсивністю тривог і довшими денними епізодами (К-4, К-5) мають одночасно високий MEI (висока асиметрія потоків) та додатний MER, тобто чистий притік бізнесу до окремих ризикованих регіонів. Натомість столично-центральный кластер К-2, що має довші тривоги із найбільшою часткою нічних тривог, поєднує дуже високий потік релокацій із від'ємним коефіцієнтом ефективності міграції MER – переважає відтік бізнесу. Помірний за рівнем ризику кластер К-6 й за показниками також виявився ближче до балансу, маючи невеликий від'ємний коефіцієнт ефективності міграції MER у середньому та позитивну медіану.

7. ВИСНОВКИ

Отже, в результаті виконаного дослідження було отримано вісім кластерів за формою динаміки тривог за допомогою Soft-DTW із параметром згладжування $\gamma = 3.0$, який контролює плавність функції втрат, забезпечуючи точніше обчислення центроїдів (барицентрів) за наявності шуму.

Зіставлення безпекових профілів кластерів із агрегованими показниками міжрегіональних релокацій бізнесу (flow, in/out-share, MER, MEI) показало: «тихі» кластери демонструють стійкий притік бізнесу; прикордонні/прифронтові кластери є інтенсивними та асиметричними за потоками, із зосередженням як ризиків, так і економічної активності; столично-центральный кластер поєднує високий потік релокацій із перевагою відтоку на тлі довших і переважно нічних тривог.

Отримані результати дозволяють ідентифікувати регіони з вищим потенціалом для відносно безпечної релокації бізнесу та можуть слугувати орієнтиром для місцевих і регіональних адміністрацій при розробці додаткових заходів з підвищення інвестиційної та безпекової привабливості регіонів, підтримання стійкості економічної активності.

Обмеження: коротке вікно спостереження; агрегація на рівні кластерів; потенційні похибки у зібраних даних та джерелах, вилучені регіони. Показаний зв'язок є асоціативним, тому для виявлення потенційних причинно-наслідкових залежностей необхідно провести панельний регресійний аналіз із фіксованими ефектами та додатковими індикаторами забезпеченості безпековою інфраструктурою. Зазначений аналіз заплановано у наших подальших дослідженнях.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Aghabozorgi S., Seyed Shirkhorshidi A., Wah T. Y. Time-series clustering – A decade review. *Information Systems*, 2015. Vol. 53. P. 16–38.

2. Tavenard R. *An introduction to Dynamic Time Warping*. Blog post, 08 June 2021. URL: <https://rtavenar.github.io/blog/dtw.html>
3. Müller M. Dynamic Time Warping. *Information Retrieval for Music and Motion*. Springer, 2007. 307 p.
4. Sakoe H., Chiba S. Dynamic Programming Algorithm Optimization for Spoken Word Recognition. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 1978, Vol. ASSP-26.
5. Cuturi M., Blondel M. Soft-DTW: a differentiable loss function for time-series. *Proceedings of the 34th International Conference on Machine Learning (ICML)*, 2017.
6. ALERTS.IN.UA. URL: <https://alerts.in.ua/en>
7. ACLED Data. *Ukraine Conflict Monitor*. URL: <https://acleddata.com/monitor/ukraine-conflict-monitor>
8. Опендатабот. URL: <https://opendatabot.ua/> (дата звернення: 24.09.2025).
9. Tavenard R. et al. Tsllearn: A machine learning toolkit for time series data. *Journal of Machine Learning Research*, 2020. Vol. 21, No. 118. P. 1–6. URL: <https://tslearn.readthedocs.io/en/stable/>
10. Tavenard R. *Differentiability of DTW and the case of soft-DTW*. Blog post, 22 June 2021. URL: <https://rtavenar.github.io/blog/softdtw.html>
11. Rousseeuw P. J. Silhouettes: a graphical aid to the interpretation and validation of cluster analysis. *Journal of Computational and Applied Mathematics*, 1987. Vol. 20, No 1. P. 53–65.
12. Pedregosa F. et al. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 2011. Vol. 12. P. 2825–2830. URL: <https://scikit-learn.org/stable/>
13. Galle O. R., Williams M. W. Metropolitan migration efficiency. *Demography*, 1972. Vol. 9, No 4. P. 655–664. DOI: 10.2307/2060672.
14. Amara E. F. L. *Migration measurement*. Lecture 05, Course SOCI 647, Texas A&M University, 22 September 2020. URL: <https://www.ernestoamaral.com/docs/soci647-20fall/Lecture05.pdf>

ІНФОРМАЦІЙНА СИСТЕМА ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ ПІДВИЩЕННЯ РІВНЯ ІНФОРМАЦІЙНОЇ БЕЗПЕКИ ОРГАНІЗАЦІЇ

Коломієць Д.С.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна
danylodomain@gmail.com [0009-0005-8811-7952]

У роботі розроблено та програмно реалізовано інформаційну систему підтримки прийняття рішень (ІСППР) для підвищення рівня інформаційної безпеки організації. Система виконує роль єдиної платформи, що інтегрує дані з різномірних джерел інформаційної безпеки (засоби виявлення загроз, сканери вразливостей, SIEM-системи тощо) та формує уніфіковане подання стану активів з урахуванням моделі конфіденційність-цілісність-доступність. На основі векторів ознак, побудованих з агрегованих показників загроз, реалізовано модель штучної нейронної мережі, яка оцінює інтегральний рівень ризику й забезпечує пріоритизацію інцидентів. Наведено архітектуру та реалізацію системи, описано етапи обробки даних і результати апробації, які демонструють, що інтегрований підхід дозволяє збільшити частку вчасно оброблених критичних інцидентів за незмінних ресурсів аналітиків, що інтерпретується як підвищення рівня інформаційної безпеки організації.

Ключові слова: інформаційна безпека, рівень інформаційної безпеки, інформаційна система підтримки прийняття рішень, інтеграція джерел даних, пріоритизація ризиків, штучні нейронні мережі.

1. ВСТУП

Сучасні організації функціонують в умовах постійно зростаючої кількості кіберзагроз та ускладнення інформаційної інфраструктури. Водночас зростає обсяг подій безпеки, що генеруються різними засобами захисту: системами виявлення та реагування на інциденти, сканерами вразливостей, системами управління журналами, мережевими сенсорами тощо. У результаті фахівці з інформаційної безпеки змушені працювати з великим масивом розрізнених даних, що ускладнює своєчасне прийняття обґрунтованих рішень.

На практиці широко застосовуються окремі продукти для захисту кінцевих точок, сканування вразливостей, централізованого збору журналів та аналізу подій безпеки. Прикладами таких рішень є засоби виявлення загроз класу EDR/XDR, платформи типу CrowdStrike Falcon, сканери вразливостей на кшталт Rapid7 Nexpose, а також різні SIEM-системи. Однак вони зазвичай функціонують у межах власних екосистем, мають різні формати представлення даних і не забезпечують єдиного інтегрованого показника ризику для конкретних активів [1–3].

Для вирішення зазначених проблем у роботі створено інформаційну систему підтримки прийняття рішень [4], яка об'єднує дані з різних джерел інформаційної безпеки, перетворює їх до уніфікованого вигляду, оцінює рівень загроз із позицій конфіденційності, цілісності та

доступності та забезпечує пріоритизацію ризиків для підтримки рішень фахівця. Метою роботи є розробка та програмна реалізація ІСППР, що виступає єдиною інтеграційною платформою для джерел даних інформаційної безпеки (зокрема, таких як засоби класу CrowdStrike Falcon, Rapid7 Nexpose, SIEM-системи), забезпечує автоматизовану пріоритизацію ризиків і сприяє підвищенню рівня інформаційної безпеки організації.

2. ПОСТАНОВКА ЗАДАЧІ

2.1. Об'єкт, предмет дослідження та вхідні/вихідні дані ІСППР

Об'єктом дослідження є процес управління інформаційною безпекою організації. Предметом дослідження є методи інтеграції та аналітичної обробки даних про події безпеки й вразливості, спрямовані на підтримку прийняття рішень щодо реагування на загрози.

Вхідними даними для реалізованої ІСППР є події та показники безпеки з різних технічних засобів, зокрема телеметрія від засобів захисту кінцевих точок та виявлення загроз (EDR/XDR-платформи, зокрема CrowdStrike Falcon), що містить інформацію про хости, виявлені інциденти, їхню критичність, час останньої активності та стан агента; результати сканування вразливостей (наприклад, Rapid7 Nexpose або інші сканери) з переліком вразливостей, CVE-ідентифікаторами, CVSS-оцінками, даними про наявність експлоїтів і час виявлення; події з інших систем моніторингу та збору журналів, зокрема SIEM, мережових сенсорів, систем контролю доступу та подібних рішень; а також довідкова інформація про активи, що описує тип системи, її роль у бізнес-процесах і критичність для організації.

Вихідними даними системи є інтегральний показник рівня загрози (ризик) для кожного активу, приведений до зручної шкали (наприклад, 0–1 або 0–100), оцінки впливу виявлених загроз на конфіденційність, цілісність та доступність (CIA), пріоритизований перелік інцидентів і вразливостей, які потребують першочергового реагування, а також ранжований список рекомендованих заходів, що може включати усунення критичних вразливостей, ізоляцію хостів, посилення контролю доступу, зміну конфігурацій тощо.

У рамках роботи розв'язано такі задачі: розроблено модель уніфікованого подання даних з різних джерел інформаційної безпеки; сформовано та реалізовано архітектуру ІСППР з модулями збору, попередньої обробки, аналітики та візуалізації; побудовано модель штучної нейронної мережі для оцінювання інтегрального рівня ризику; проведено апробацію системи та показано, що її використання дозволяє підвищити рівень інформаційної безпеки за рахунок кращої пріоритизації критичних інцидентів.

2.2. Рівень інформаційної безпеки та його оцінювання

Під рівнем інформаційної безпеки організації розумітимемо узагальнену характеристику ступеня захищеності її інформаційних ресурсів від порушення конфіденційності, цілісності та доступності. Чим нижчими є ймовірність реалізації загроз і очікувані наслідки інцидентів для зазначених властивостей, тим вищим вважається рівень інформаційної безпеки.

3. ОГЛЯД ІСНУЮЧИХ РІШЕНЬ

На ринку засобів інформаційної безпеки широко представлені системи виявлення вторгнень і атак, платформи захисту кінцевих точок і виявлення загроз (EDR/XDR), сканери вразливостей, а також SIEM- та SOAR-рішення для збору та кореляції подій. Більшість таких рішень орієнтована на свою предметну область: захист кінцевих точок, пошук вразливостей, аналіз мережевого трафіку тощо. Наприклад, EDR/XDR-платформи (зокрема CrowdStrike Falcon) забезпечують виявлення й блокування атак на робочих станціях і серверах, тоді як сканери вразливостей (на кшталт Rapid7 Nexpose) дозволяють виявляти помилки конфігурації та програмні вразливості, ранжуючи їх за CVSS-методологією.

Попри наявність потужних інструментів, на практиці виникають такі проблеми: по-перше, спостерігається фрагментованість інформації, оскільки різні продукти мають власні інтерфейси, формати логів та моделі даних; по-друге, відсутній єдиний, інтегральний показник ризику для активів, який би враховував дані з кількох джерел; по-третє, на рівні систем підтримки прийняття рішень для аналітиків досі обмежено застосовуються методи машинного навчання.

У рамках даної роботи ці проблеми вирішуються шляхом створення ІСППР, яка виступає надбудовою над існуючими засобами безпеки, консолідує їхні дані та надає інтегральні оцінки ризику й рекомендації щодо першочергових дій.

4. АРХІТЕКТУРА ТА РЕАЛІЗАЦІЯ СИСТЕМИ

Розроблена ІСППР виконує роль надбудови над існуючими засобами забезпечення інформаційної безпеки та виступає єдиною точкою доступу до даних про події й вразливості. Замість аналізу окремих звітів сканерів вразливостей, сповіщень EDR/XDR та записів SIEM аналітик отримує консолідоване подання стану активів, інтегральні оцінки ризику та сформований системою список першочергових заходів реагування. Узагальнену архітектуру ІСППР та приклад інтеграції декількох джерел даних (EDR/XDR, сканерів вразливостей, SIEM тощо) показано на рис. 1.

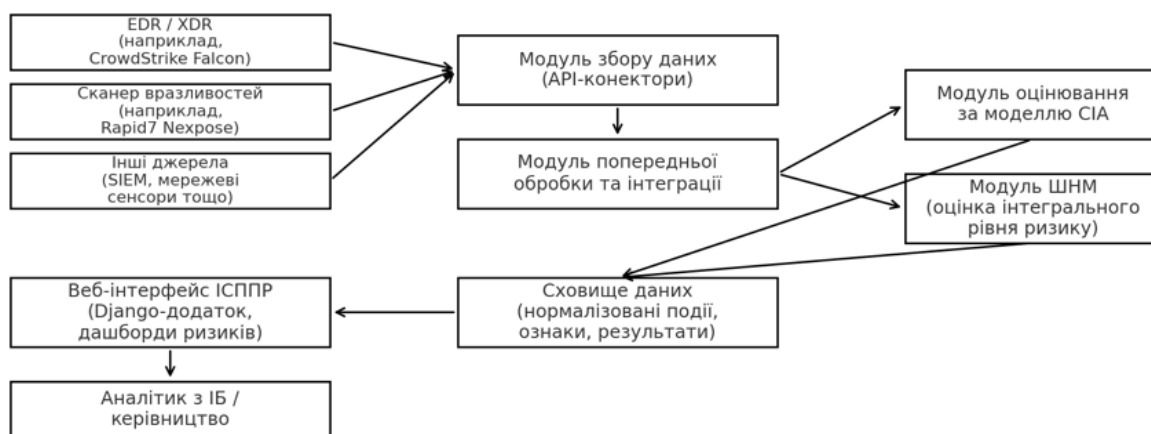


Рисунок 1. Узагальнена архітектура ІСППР та інтеграція джерел даних інформаційної безпеки

Архітектура системи має модульну структуру і включає кілька основних компонентів. Модуль збору даних забезпечує підключення до різних джерел інформаційної безпеки через API, конектори або агенти. Поточна реалізація підтримує інтеграцію з платформами типу CrowdStrike Falcon, сканерами вразливостей на кшталт Rapid7 Nexpose та іншими засобами, які надають машиночитаний інтерфейс [6; 7]. Архітектура модуля дозволяє розширювати перелік джерел без істотної модифікації решти системи.

Модуль попередньої обробки та інтеграції відповідає за нормалізацію й уніфікацію даних: перетворення форматів, очищення, обробку відсутніх значень, приведення категоріальних ознак (тип ОС, роль сервера, критичність активу) до числового вигляду. На цьому етапі події з різних джерел зіставляються з конкретними активами, формується вектор ознак для кожного активу.

Модуль оцінювання за моделлю CIA призначений для визначення орієнтовного впливу кожної вразливості чи інциденту на конфіденційність, цілісність та доступність. На основі отриманих оцінок обчислюються агреговані показники за CIA для кожного активу, які відображають накопичений вплив загроз.

Модуль машинного навчання реалізовано у вигляді окремого компонента на основі бібліотек для побудови штучних нейронних мереж. Він виконує навчання, валідацію, збереження та застосування моделі для нових даних, а виходом моделі є інтегральний рівень ризику для активу.

Підсистема зберігання даних забезпечує збереження нормалізованих подій, ознак і результатів аналізу в реляційній системі керування базами даних (наприклад, PostgreSQL), що спрощує інтеграцію з веб-інтерфейсом та гарантує цілісність інформації. Веб-інтерфейс реалізовано у вигляді Django-додатку, який надає користувачам доступ до результатів аналізу: списку активів з інтегральними показниками ризику, деталізації за CIA, відображення динаміки та перегляду рекомендацій. Інтерфейс підтримує фільтрацію активів за рівнем загрози, типом системи, джерелом подій тощо. На рис. 2 наведено фрагмент реалізованого інтерфейсу ІСППР із відображенням загальної кількості хостів, їхнього поточного статусу та мережевих адрес (частина адрес може бути прихована з міркувань конфіденційності).

Information Security DSS					Home About CrowdStrike
CrowdStrike – хости					
#	Hostname	Платформа	IP	Last seen	Онлайн-статус
1	vm116	Linux	192.168.61.116	2025-11-27T23:47:26Z	Online
2	SOE-1C	Windows	192.168.61.111	2025-11-27T23:26:33Z	Online
3	nexspose	Linux	192.168.61.115	2025-11-27T23:25:08Z	Online
4	WIN-DMEEK2GS945	Windows	192.168.10.26	2025-11-27T23:20:27Z	Online
5	M-MARTYNENKO	Windows	192.168.1.178	2025-11-27T23:13:35Z	Online
6	HP255-7	Windows	192.168.88.37	2025-11-27T23:06:19Z	Online
7	YELIZAVETA	Windows	192.168.0.105	2025-11-27T22:31:06Z	Offline

Рисунок 2. Фрагмент інтерфейсу ІСППР з відображенням загальної кількості хостів, їх статусів та адрес

Завдяки такій архітектурі система не прив'язана до конкретного вендора засобів безпеки: CrowdStrike, Nexpose чи інші рішення виступають лише окремими джерелами, які можуть бути замінені або доповнені.

5. МОДЕЛЬ ШТУЧНОЇ НЕЙРОННОЇ МЕРЕЖІ

Для оцінювання інтегрального рівня ризику в реалізованій системі використано багаторівневу перцептронну нейронну мережу (Multi-Layer Perceptron) [5]. Вхідним шаром моделі є вектор ознак, який формується для кожного активу та включає, зокрема:

- агреговані показники впливу загроз на конфіденційність, цілісність і доступність;
- кількість критичних, високих, середніх і низьких вразливостей;
- максимальні та середні CVSS-оцінки для вразливостей активу;
- наявність публічних експлойтів для виявлених вразливостей;
- частоту та нещодавність інцидентів, зафіксованих засобами виявлення загроз;
- критичність активу для бізнес-процесів організації.

Приклад структури моделі:

- вхідний шар розмірності n (кількість ознак);
- 2–3 приховані шари по 32–64 нейрони з активацією ReLU;
- вихідний шар з одним нейроном і сигмоїдною активацією (для нормованого ризику в діапазоні 0–1) або лінійною (для шкали 0–100).

Навчання моделі здійснюється на підготовленій вибірці, де цільовим значенням є експертна оцінка рівня ризику або інтегральний показник, сформований на основі правил. Для оцінювання якості моделі використовуються середня квадратична помилка, середня абсолютна помилка та коефіцієнт детермінації. Для запобігання перенавчанню застосовуються регуляризація та крос-валідація.

Після завершення етапу навчання модель використовується в режимі прогнозування для нових даних, що надходять від підключених джерел безпеки. Отриманий інтегральний рівень ризику надалі використовується для пріоритизації інцидентів і планування заходів реагування.

6. РЕЗУЛЬТАТИ ТА ОБГОВОРЕННЯ

Для апробації системи було сформовано тестовий набір даних, який відображає типову ситуацію для умовної організації: набір активів із різним рівнем критичності, вразливостями різної тяжкості та інцидентами, виявленими засобами моніторингу. Для кожного активу було побудовано вектор ознак і отримано інтегральний рівень ризику за допомогою нейронної мережі.

Окрему увагу приділено сценарію обмежених ресурсів, коли аналітик інформаційної безпеки може опрацювати лише перші N інцидентів у сформованій черзі. Було порівняно три підходи до сортування інцидентів за пріоритетністю:

- використання лише результатів сканування вразливостей (наприклад, на основі CVSS-балів сканера типу Nexpose);
- використання лише телеметрії засобів виявлення загроз (наприклад, сповіщення засобу класу CrowdStrike Falcon);
- інтегрований підхід, коли черга формувалася на основі інтегрального рівня ризику, розрахованого ІСППР з використанням нейронної мережі.

Для кожного значення N обчислювалася частка дійсно критичних інцидентів (за експертною оцінкою), які потрапили до перших N записів. Залежність цієї частки від N для трьох підходів наведено на рис. 3.

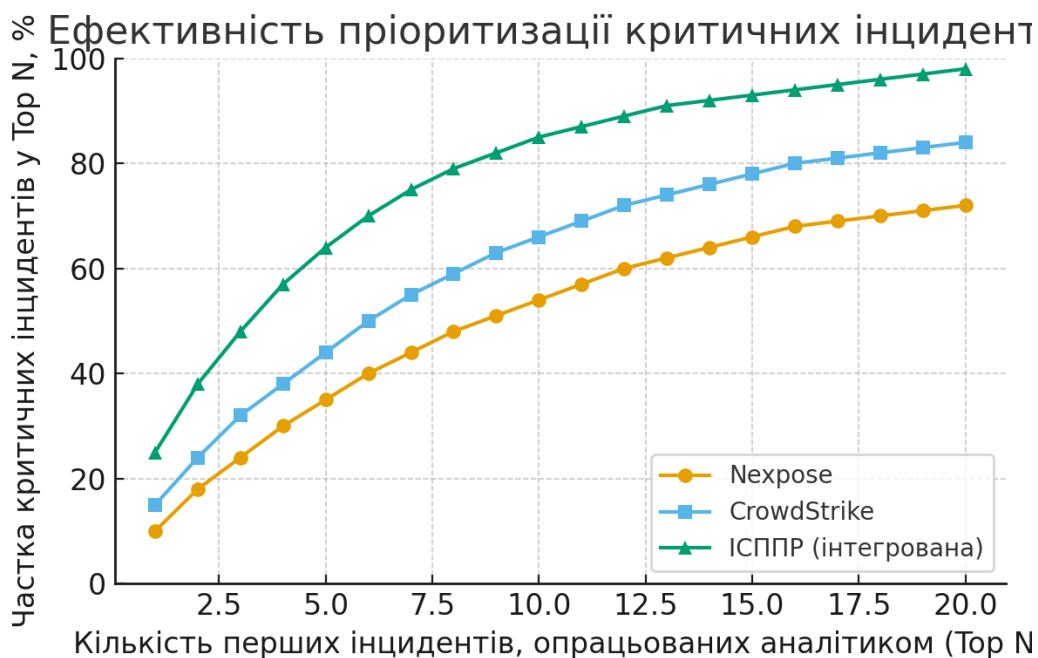


Рисунок 3. Ефективність пріоритизації критичних інцидентів для окремих джерел даних і інтегрованої ІСППР

Як показують отримані результати, інтегрована ІСППР забезпечує вищу частку вчасно відібраних критичних інцидентів при тій самій кількості інцидентів, які здатен опрацювати аналітик. Наприклад, для перших десяти інцидентів підхід, що базується лише на даних сканера вразливостей, дозволяє охопити орієнтовно 50–60 % критичних подій, аналіз тільки телеметрії EDR/XDR — близько 60–70 %, тоді як інтегрована ІСППР — понад 80 %. Це означає, що за однакових людських ресурсів більша частина найнебезпечніших інцидентів потрапляє до поля зору аналітика.

З позицій рівня інформаційної безпеки такі результати можна інтерпретувати як **зменшення нереалізованого ризику**: знижується ймовірність того, що критичні інциденти залишаться необробленими внаслідок їх низького пріоритету в черзі. На практиці це призводить до зростання частки активів, для яких інтегральний рівень ризику знижується до прийнятних значень, а отже, до підвищення рівня інформаційної безпеки організації в цілому.

Крім того, завдяки уніфікованому поданню даних і агрегованим показникам за СІА, аналітик отримує можливість не лише бачити інтегральний ризик, а й розуміти структуру загроз: які саме вразливості, інциденти та джерела даних роблять найбільший внесок у ризик для конкретного активу.

7. ВИСНОВКИ

У роботі розроблено та реалізовано інформаційну систему підтримки прийняття рішень для підвищення рівня інформаційної безпеки організації. Основні результати полягають у такому:

- створено ІСППР, яка виконує роль єдиної інтеграційної платформи для різномірних джерел даних інформаційної безпеки;
- розроблено модель уніфікованого подання подій та вразливостей із урахуванням моделі конфіденційність–цілісність–доступність;
- сформовано та реалізовано архітектуру системи, що включає модулі збору даних, попередньої обробки, СІА-оцінювання, машинного навчання та веб-візуалізації;
- реалізовано модель штучної нейронної мережі для оцінювання інтегрального рівня ризику для активів;
- проведено апробацію системи, яка показала, що інтегрована пріоритизація інцидентів на основі даних з кількох джерел дозволяє підвищити частку вчасно опрацьованих критичних інцидентів за незмінних ресурсів аналітиків.

Практичне значення роботи полягає в тому, що реалізована ІСППР дозволяє не лише зменшити інформаційне навантаження на фахівців з інформаційної безпеки, а й формалізувати поняття рівня інформаційної безпеки через інтегральний показник ризику та пріоритизацію інцидентів. Підвищення частки виявлених і своєчасно оброблених критичних подій без збільшення ресурсів захисту свідчить про реальне зростання рівня захищеності інформаційних активів.

Подальші дослідження доцільно спрямувати на розширення переліку підключених джерел даних, використання більш складних моделей глибинного навчання, а також на інтеграцію ІСППР з процесами.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. ISO/IEC 27001:2022 Information security, cybersecurity and privacy protection — Information security management systems — Requirements. Geneva : International Organization for Standardization, 2022. 36 p.

2. ISO/IEC 27005:2018 Information technology — Security techniques — Information security risk management. Geneva : International Organization for Standardization, 2018. 75 p.

3. NIST. Guide for Conducting Risk Assessments : NIST Special Publication 800-30 Rev. 1. Gaithersburg, MD : National Institute of Standards and Technology, 2012. 95 p.
4. Turban E., Sharda R., Delen D. Decision Support and Business Intelligence Systems. 9th ed. Upper Saddle River : Pearson, 2014. 720 p.
5. Goodfellow I., Bengio Y., Courville A. Deep Learning. Cambridge, MA : MIT Press, 2016. 775 p.
6. CrowdStrike. CrowdStrike Falcon Platform : Product Overview. URL: <https://www.crowdstrike.com/resources/> (дата звернення: 22.11.2025).
7. Rapid7. Vulnerability Management (InsightVM / Nexpose) Documentation. URL: <https://docs.rapid7.com/> (дата звернення: 22.11.2025).

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПРОГНОЗУВАННЯ ФІНАНСОВИХ ЧАСОВИХ РЯДІВ ТА РОБАСТНОГО АНАЛІЗУ ПРОГНОЗНОЇ НЕВИЗНАЧЕНОСТІ

Макітрук М.Т.¹, Селін Ю.М.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ vbeamv@gmail.com [0009-0006-4195-8006],

² selinyurij1963@gmail.com [0000-0002-7562-8586]

У роботі представлено інтелектуальну систему прогнозування фінансових часових рядів, що поєднує LSTM-моделі, статистичну оцінку невизначеності та робастні методи інтеграції прогнозів. Запропоновано підхід, який об'єднує залишкові розподіли, модель Black–Litterman, Risk Parity та робастну оптимізацію для формування збалансованих та стабільних рішень. Представлено поєднання глибинного прогнозування з оцінкою достовірності прогнозів а також можливість застосування системи для аналізу ризиків.

Ключові слова: фінансові часові ряди, LSTM, прогнозування, невизначеність, Black–Litterman, робастна оптимізація

1. ВСТУП

Фінансові часові ряди характеризуються високою волатильністю, нелінійністю та залежністю від широкого спектра економічних і поведінкових факторів, що ускладнює їх точне прогнозування за допомогою класичних статистичних моделей. У сучасних умовах стрімкого розвитку штучного інтелекту особливого значення набувають глибинні нейронні мережі та методи інтелектуального опрацювання даних, здатні виявляти приховані структури, враховувати довгострокові залежності та адаптуватися до складної динаміки ринку.

У цій роботі досліджено комплексний підхід до прогнозування фінансових часових рядів та аналізу достовірності прогнозу. Основою прогнозної моделі є рекурентна нейронна мережа типу LSTM [1], розроблена для роботи з послідовними даними та виявлення нелінійних патернів. Для кількісної оцінки невизначеності застосовано підхід на основі залишкових розподілів, що дозволяє формувати інтервали довіри та коригувати прогноз залежно від похибок моделі.

Другим елементом дослідження є інтелектуальне опрацювання прогнозних оцінок із використанням гібридних методів. На цьому етапі інтегруються історичні статистичні характеристики, результати прогнозів та робастні оптимізаційні підходи, зокрема модель Black–Litterman [2, 3], метод Risk Parity та Distributionally-Robust Optimization на основі кулі Васерштайна. Це забезпечує комплексний погляд на якість прогнозів та дозволяє формувати інтерпретовані оцінки, що враховують рівень ризику, нестабільність та довіру до прогнозу.

Запропонована система поєднує методи машинного навчання, статистики невизначеності та обчислювального моделювання, що робить її універсальним інструментом для дослідження поведінки фінансових часових рядів і побудови інтелектуальних рішень на основі прогнозних даних.

2. МЕТОДИ ТА АРХІТЕКТУРА ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ

2.1. Підготовка даних для прогнозування

Інтелектуальна система побудована як поєднання глибинного прогнозування, статистичного аналізу невизначеності та робастного опрацювання прогнозів. На першому етапі виконується очищення даних, нормалізація та перехід до лог-дохідностей, що дає змогу стабілізувати дисперсію та виділити пропорційні зміни ціни. Формально лог-дохідність визначається як

$$r_t = \ln \left(\frac{P_t}{P_{t-1}} \right),$$

а сформований із ковзних вікон тензор

$$X \in \mathbb{R}^{(T-n) \times n}$$

слугує вхідними даними для LSTM-моделі.

2.2. Прогнозування часових рядів за допомогою LSTM

Прогнозування ґрунтується на архітектурі LSTM, яка моделює довгострокові часові залежності через систему гейтів. Її динаміка описується рівняннями

$$\begin{aligned} f_t &= \sigma(W_f h_{t-1} + U_f x_t + b_f), \\ i_t &= \sigma(W_i h_{t-1} + U_i x_t + b_i), \\ o_t &= \sigma(W_o h_{t-1} + U_o x_t + b_o), \\ \tilde{c}_t &= \tanh(W_c h_{t-1} + U_c x_t + b_c), \end{aligned}$$

що формують новий стан пам'яті

$$c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t,$$

а вихід LSTM

$$h_t = o_t \odot \tanh(c_t)$$

використовується для формування прогнозу

$$\hat{y}_{t+1} = g(h_t).$$

Таке подання забезпечує здатність моделі виявляти як локальні, так і глобальні структури ринку.

2.3. Оцінка прогнозної невизначеності

Для оцінювання надійності прогнозу аналізуються залишки моделі

$$\varepsilon_t = y_t - \hat{y}_t.$$

Апроксимація залишків нормальним розподілом

$$\varepsilon_t \sim N(\mu_\varepsilon, \sigma_\varepsilon^2)$$

дозволяє формувати інтервали достовірності, де оцінки

$$\mu_\varepsilon = \frac{1}{T} \sum_{t=1}^T \varepsilon_t, \quad \sigma_\varepsilon^2 = \frac{1}{T-1} \sum_{t=1}^T (\varepsilon_t - \mu_\varepsilon)^2$$

використовуються для побудови 95% довірчого інтервалу прогнозу:

$$\hat{y}_{t+1}^{(low, up)} = \hat{y}_{t+1} + \mu_\varepsilon \pm 1.96 \sigma_\varepsilon.$$

Таким чином, система оцінює не лише значення прогнозу, але і ступінь його надійності.

2.4. Інтелектуальне опрацювання прогнозних оцінок

Середня прогнозна дохідність

Середня прогнозна дохідність для кожного активу обчислюється як

$$\mu_{mean,i} = \frac{1}{H} \sum_{k=1}^H \frac{F_{i,k} - P_{i,last}}{P_{i,last}}$$

Модель Black–Litterman

Для узгодження прогнозних оцінок моделі з історичними статистичними характеристиками застосовується модель Black–Litterman [2, 3], де постерійні дохідності визначаються рівнянням

$$\mu_{BL} = [(\tau\Sigma)^{-1} + P^T\Omega^{-1}P]^{-1}[(\tau\Sigma)^{-1}\mu_{prior} + P^T\Omega^{-1}q]$$

Цей підхід дозволяє “згладити” прогноз, зменшуючи вплив шуму.

Risk Parity

Risk Parity забезпечує рівномірний розподіл ризику між компонентами портфеля [4]. Кожен внесок ризику визначається як

$$RC_i = w_i(\Sigma w)_i.$$

Цільова умова

$$RC_1 = RC_2 \dots = RC_n$$

призводить до апроксимації ваг

$$w_i \approx \frac{1/\sigma_i^2}{\sum_j 1/\sigma_j^2}.$$

Distributionally-Robust Optimization

У DRO розглядається найгірший сценарій у межах кулі Васерштайна [5]. Оцінка дохідності коригується згідно:

$$\mu_{DRO,i} = \mu_{mean,i} - \rho\sigma_i.$$

2.5. Інтерактивне графічне середовище

Усі алгоритмічні блоки інтегровані в графічне середовище, яке забезпечує завантаження даних, побудову прогнозів, формування інтервалів довіри, порівняння моделей (Mean, BL, RP, DRO) та візуальний аналіз графіків і таблиць.

3. РЕЗУЛЬТАТИ РОБОТИ ІНТЕЛЕКТУАЛЬНОЇ СИСТЕМИ

Розроблена система забезпечує повний цикл опрацювання фінансових часових рядів — від завантаження та підготовки даних до побудови прогнозів і формування інтегрованих аналітичних оцінок. Інтерактивне середовище (рис. 1) реалізоване у вигляді кількох логічних модулів: імпорт даних, нормалізація цінкових рядів, прогнозування на основі LSTM та аналіз прогнозової невизначеності. Окремий блок забезпечує інтелектуальне опрацювання результатів, включаючи моделі Mean, Black–Litterman, Risk Parity та Wasserstein-DRO. Усі компоненти пов’язані єдиним інтерфейсом, який дає змогу керувати параметрами, переглядати графіки, автоматичні рекомендації та результати оптимізацій.

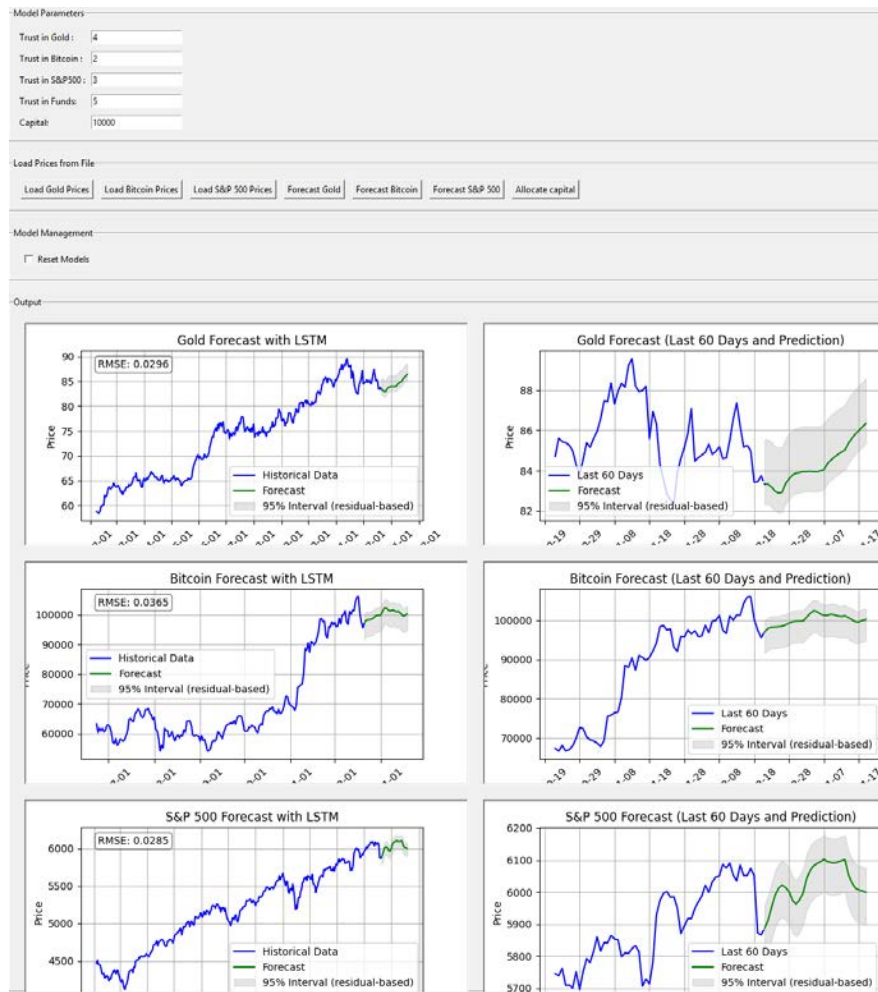


Рисунок 1. Головне вікно інтелектуальної системи.

Для кожного активу система генерує точковий прогноз на 30 днів та 95% інтервал довіри, побудований на основі нормального наближення залишків.

Кожен прогноз супроводжується автоматичною аналітикою, яка включає оцінку тренду, рекомендацію дії, значення RMSE, волатильність та дисперсію залишків (рис. 2).

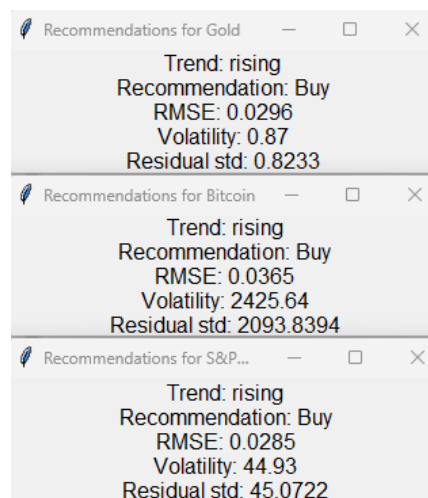


Рисунок 2. Автоматичні рекомендації системи для золота, Bitcoin і S&P500.

Для інтегральної оцінки прогнозів реалізовано чотири моделі: середня оцінка (Mean), Black–Litterman, Risk Parity, Distributionally-Robust Optimization (Wasserstein DRO).

Результат обчислення ваг портфеля та очікуваних доходностей (рис. 3, табл. 1). Ці дані відображають вплив різних припущень щодо ринку на структуру aloкації капіталу та стійкість моделі до невизначеності.

Таблиця 1. Порівняння ваг та прогнозних прибутків моделей

Актив	М_вага	BL_вага	RP_вага	М_приб.	BL_приб.	DRO_приб.
Gold	28.57%	28.57%	2.47%	0.65%	0.12%	- 0.89%
Bitcoin	14.29%	14.29%	0.35%	4.68%	0.67%	- 2.03%
S&P500	21.43%	21.43%	3.66%	2.75%	0.42%	- 0.42%
Funds	35.71%	35.71%	93.51%	0.83%	0.19%	0.09%

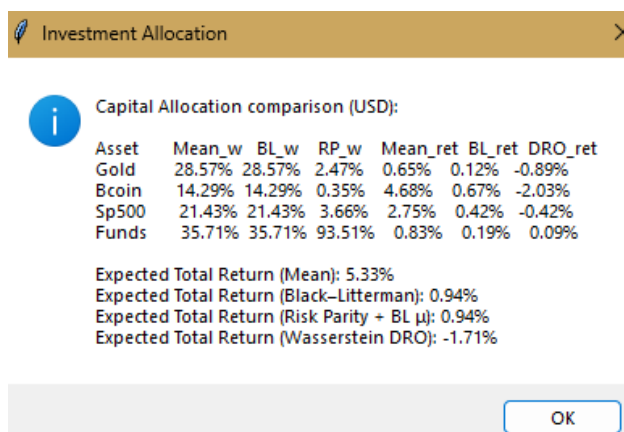


Рисунок 3. Порівняння ваг портфеля та прогнозних доходностей різних моделей

4. ВИСНОВКИ

У роботі представлено комплексну систему інтелектуального прогнозування фінансових часових рядів, що поєднує методи глибинного навчання, статистичної оцінки невизначеності та робастного моделювання. Результати дослідження підтверджують, що така інтеграція дозволяє суттєво підвищити точність прогнозів та підсилити інтерпретованість моделі порівняно з класичними підходами.

Запропонована LSTM-архітектура продемонструвала стабільну ефективність на різних класах фінансових активів, забезпечуючи низькі значення RMSE та здатність відтворювати нелінійну динаміку ринкових процесів. Включення резидуального статистичного аналізу та побудови 95% довірчих інтервалів створює механізм кількісної оцінки ризику прогнозу, що особливо важливо для нестабільних часових рядів із високою волатильністю.

Проведений синтез декількох інтелектуальних стратегій обробки прогнозованих значень: класичної очікуваної оцінки, байєсівської корекції Black–Litterman, risk-parity нормалізації та distributionally-robust моделювання на основі метрики Васерштайна. Усі методи інтегровані у єдине візуальне середовище, що дозволяє порівнювати їх поведінку, чутливість до невизначеності та узгодженість прогнозів між собою. Проведений аналіз показав, що кожен із підходів має власну зону ефективності, що підкреслює необхідність багатомодельного представлення прогнозу у задачах високої невизначеності.

Показано поєднання автоматизованої обробки даних, прогнозування та інтелектуальної оцінки результатів у зручному програмному середовищі. Це забезпечує прозорість обчислень, відтворюваність експериментів та можливість подальшого розширення інструментарію. Одержані результати свідчать про надійність обраної методології та її відповідність сучасним вимогам до систем прогнозування складних часових рядів.

Таким чином, розроблена система підтвердила свою ефективність як інструмент глибинного аналізу ринку та наукового дослідження динаміки фінансових процесів. Робота створює підґрунтя для подальшого розвитку, зокрема впровадження більш складних моделей, адаптивних оцінок невизначеності та розширення робастних методів у контексті складних реальних сценаріїв.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Idzorek T. A Step-By-Step Guide to the Black–Litterman Model: Incorporating User-Specified Confidence Levels. SSRN Electronic Journal. 2019. 48 p.
2. Roncalli T. Introducing Expected Returns into Risk Parity Portfolios. MPRA Paper. 2013. № 49821. 42 p.
3. Wu Z., Sun K. Distributionally Robust Optimization with Wasserstein Metric for Multi-Period Portfolio Selection Under Uncertainty. ResearchGate Preprint. 2022. 28 p.
4. Ernstsson H. The Black–Litterman Asset Allocation Model: A Study of Bayesian Portfolio Construction. KTH Royal Institute of Technology. Stockholm, 2021. 54 p.
5. Hochreiter S., Schmidhuber J. Long Short-Term Memory. Neural Computation. 1997. Vol. 9, No. 8. P. 1735–1780.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ АНАЛІЗУ ЕКОНОМІЧНИХ РИЗИКІВ

Мельник М.С.¹, Левенчук Л.Б.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ melnyk.mykola@lil.kpi.ua

У роботі досліджено розробку інтелектуальної системи підтримки прийняття рішень (ІСППР) для аналізу економічних ризиків та прогнозування банкрутства компаній. Метою є порівняння ефективності сучасних методів машинного навчання для цієї задачі. Дослідження включає застосування низки моделей, зокрема DecisionTree, RandomForest, SVM, GradientBoosting та XGBoost. Методи дослідження зосереджені на подоланні проблеми незбалансованості вхідних даних за допомогою техніки SMOTE та відборі найбільш значущих ознак з використанням RFE. Основні отримані результати демонструють, що хоча SVM показує найвищий F1 Score (~0.96), XGBoost є беззаперечним лідером за метрикою ROC AUC (~0.85), що вказує на його кращу розрізнявальну здатність. Наукова новизна полягає у комплексному порівнянні моделей в умовах дисбалансу класів та акценті на інтерпретованості результатів за допомогою методів, як-от SHAP. Практична значимість роботи полягає у проектуванні архітектури веб-орієнтованої СППР, яка надає ризик-менеджерам прозорий та ефективний інструмент для прийняття обґрунтованих рішень.

Ключові слова: інтелектуальна система підтримки прийняття рішень (ІСППР), аналіз економічних ризиків, прогнозування банкрутства, машинне навчання, XGBoost, SVM, незбалансовані дані, SMOTE.

1. ВСТУП

Сучасне економічне середовище є надзвичайно складною і динамічною системою, що піддається впливу безлічі факторів – від макроекономічних показників до поведінкових аспектів ринку [1]. Ефективний аналіз економічних ризиків у таких умовах вимагає використання аналітичних інструментів, здатних не лише швидко обробляти великі обсяги інформації, але й виявляти приховані закономірності та прогнозувати ймовірність настання кризових подій. У цьому контексті, інтелектуальні системи підтримки прийняття рішень (ІСППР) стають незамінним інструментом.

Такі системи є критично важливими для осіб, що приймають швидкі стратегічні рішення: фінансових аналітиків, ризик-менеджерів та керівників, які відповідають за стабільність компанії. ІСППР, що використовують потужні моделі машинного навчання, надають таким фахівцям змогу отримувати поглиблений аналіз та прогнози, що ґрунтуються на об'єктивних даних.

Застосування сучасних ансамблевих методів, таких як XGBoost та LightGBM, дозволяє досягти високої якості прогнозів, враховуючи навіть таку складну проблему, як незбалансованість вхідних даних [2].

2. ОПИС НАБОРУ ДАНИХ

В даній роботі було використано набір даних, який відображає історичні фінансові дані американських компаній. У вибірці містяться 18 фінансових показників (ознаки X1–X18, такі як: поточні активи X1, чистий прибуток X6, ринкова вартість X8, X12 та загальні активи X10). Цільова характеристика – status_label – відображає стан компанії («alive» або «failed»). Моделі машинного навчання мають певні вимоги до формату вхідних даних і, зокрема, більшість моделей потребує виключно чисельні значення. Оскільки цільова характеристика status_label в нашому наборі представляє категоріальні значення («alive»/«failed»), постала необхідність в її конвертуванні в чисельний формат (де «alive» = 0 та «failed» = 1). На наступному етапі було виконано масштабування всієї сукупності фінансових показників X1-X18 для покращення подальших результатів побудованих моделей.

Для аналізу вибірки була побудована кореляційна матриця (рис. 1). Аналіз матриці виявив наступні залежності.

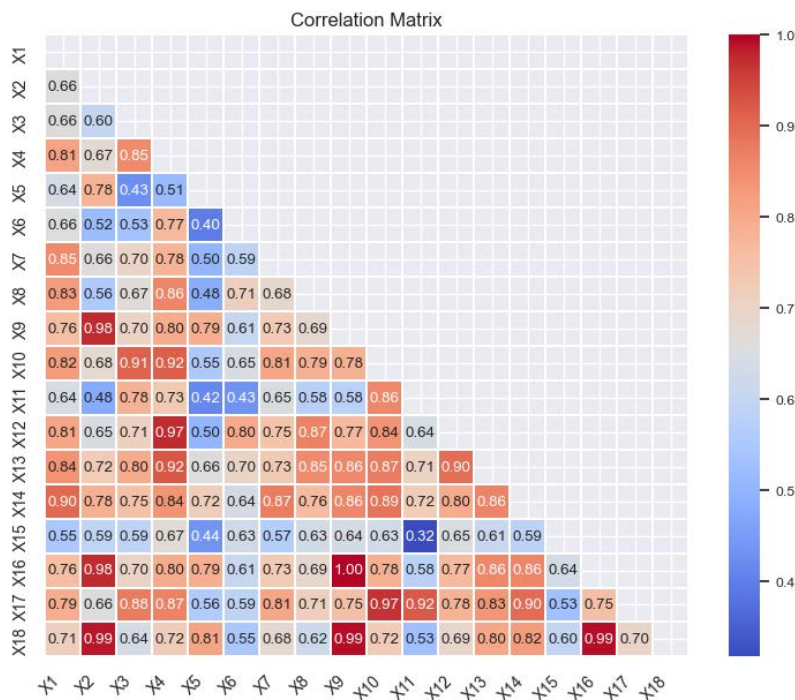


Рисунок 1. Кореляційна матриця

Сильні позитивні кореляції. Було виявлено зв'язки, що вказують на те, що змінні мають тенденцію зростати разом. Найсильніші з них спостерігаються між X3 та X16 (0.83), а також між X15 та X18 (0.78).

Помірні кореляції. Спостерігається значуща, але не абсолютна залежність, наприклад, між X7 та X14 (0.75).

Слабкі кореляції. Деякі пари, як-от X13 та X14 (0.01), демонструють практично повну відсутність лінійного зв'язку.

Негативні кореляції. Матриця також показала зворотні зв'язки, де збільшення однієї змінної відповідає зменшенню іншої, наприклад, між X1 та X15 (–0.67).

Ці результати вказують на потенційні проблеми мультиколінеарності, що може негативно вплинути на побудову прогностичних моделей. Для зменшення надмірності даних може знадобитися відбір ознак або зменшення розмірності.

3. ОПИС ВИКОРИСТАНИХ МОДЕЛЕЙ МАШИННОГО НАВЧАННЯ

У роботі було досліджено та порівняно низку моделей машинного навчання для вирішення задачі прогнозування банкрутства. Оскільки дані для такого типу завдань зазвичай є сильно незбалансованими (кількість компаній-банкрутів значно менша за кількість фінансово здорових), ключовим етапом підготовки даних стало застосування методу SMOTE (Synthetic Minority Over-sampling Technique). Цей підхід дозволяє збалансувати класи шляхом генерації синтетичних прикладів для класу меншості, що є критично важливим для коректного навчання моделей та уникнення їх упередженості в бік мажоритарного класу [3].

Основні методи, що розглянуті в роботі, включають:

Дерево рішень (DecisionTree) є одним з базових методів керованого машинного навчання, що застосовується для задач класифікації та регресії. Модель являє собою ієрархічну деревоподібну структуру, де кожен внутрішній вузол відповідає перевірці певної ознаки (фінансового показника), кожна гілка – результату цієї перевірки, а кожен листковий вузол – кінцевому рішенню або класу (наприклад, «alive» або «failed»). Побудова дерева відбувається шляхом рекурсивного поділу вибірки на підмножини на основі критеріїв, що максимізують чистоту вузлів. Незважаючи на високу інтерпретованість, метод схильний до перенавчання [4].

Випадковий ліс (RandomForest) – це ансамблевий метод машинного навчання, що розширює концепцію дерев рішень для підвищення їхньої стабільності та точності. Алгоритм будує велику кількість (ансамбль) незалежних дерев рішень під час тренування, причому кожне дерево навчається на випадковій підмножині даних (бутстреп-вибірка) та з використанням випадкової підмножини ознак для кожного поділу. Кінцевий прогноз для задачі класифікації визначається шляхом голосування більшості дерев в ансамблі, що дозволяє нівелювати схильність окремих дерев до перенавчання та покращити узагальнюючу здатність моделі [4].

Метод опорних векторів (Support Vector Machine, SVM) є потужним алгоритмом керованого навчання, що використовується для лінійної та нелінійної класифікації. Основна ідея методу полягає у побудові оптимальної розділяючої гіперплощини у багатовимірному просторі ознак, яка б максимізувала зазор (margin) між найближчими точками різних класів (опорними векторами). Для нелінійно розділюваних даних SVM використовує «ядровий трюк» (kernel trick) для проектування даних у простір вищої розмірності, де лінійний поділ стає можливим, що робить його ефективним для складних задач класифікації [4].

Гرادієнтний бустинг (GradientBoosting) – це ансамблевий метод машинного навчання, що будує моделі у послідовний, поетапний спосіб. На відміну від випадкового лісу, де дерева будуються паралельно та незалежно, градієнтний бустинг формує нову модель (зазвичай дерево рішень) на кожній ітерації таким чином, щоб вона виправляла помилки (залишки) попередньої моделі. Навчання відбувається шляхом оптимізації функції втрат за допомогою методу градієнтного спуску, що дозволяє ансамблю поступово «посилювати» (boost) слабкі класифікатори, створюючи високоточну прогностичну модель [4].

Екстремальний градієнтний бустинг (eXtreme Gradient Boosting) є не стільки окремим алгоритмом, скільки високоефективною та масштабованою реалізацією методу градієнтного бустингу [2]. Його архітектура оптимізована для швидкодії та продуктивності, зокрема за рахунок паралелізації обчислень при побудові дерев та ефективної обробки пропущених значень. Ключовою відмінністю є включення L1 та L2 регуляризації (штрафів за

складності моделі) безпосередньо у функцію мети, що дозволяє ефективно контролювати перенавчання та часто забезпечує вищу точність порівняно зі стандартними реалізаціями градієнтного бустингу.

Також у роботі було застосовано метод рекурсивного усунення ознак (RFE – Recursive Feature Elimination). Цей підхід використовується для відбору найбільш значущих фінансових показників, дозволяючи зменшити розмірність даних і складність моделей, що в свою чергу може покращити їхню здатність до узагальнення та знизити ризик перенавчання.

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Результати порівняння продуктивності моделей, представлені на графіках (рис. 2–4), демонструють суттєві відмінності в їхній ефективності залежно від обраної метрики. Було використано наступні метрики: accuracy, precision, recall, F1, AUC-ROC і крива precision-recall [5].

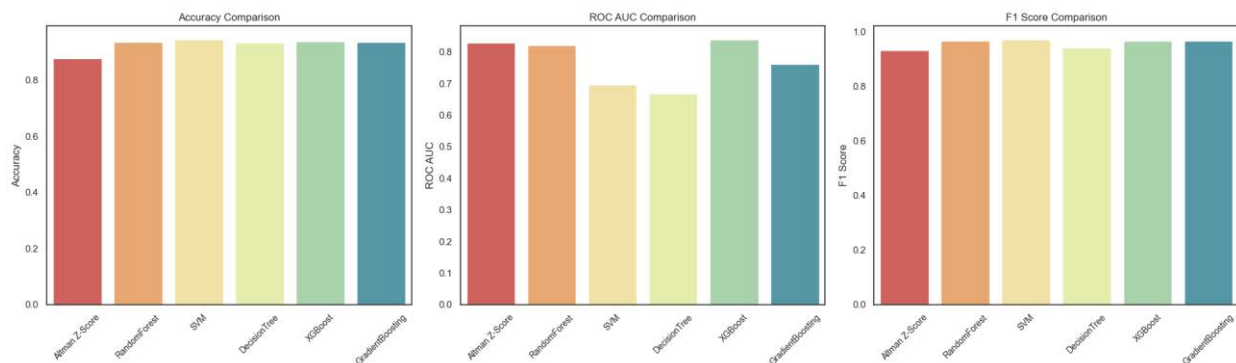


Рисунок 2. Метрики адекватності моделей на тестовій вибірці

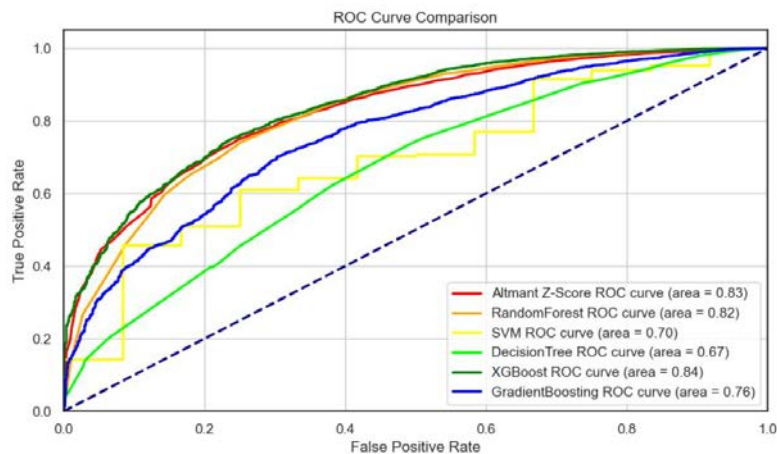


Рисунок 3. ROC-криві побудованих моделей

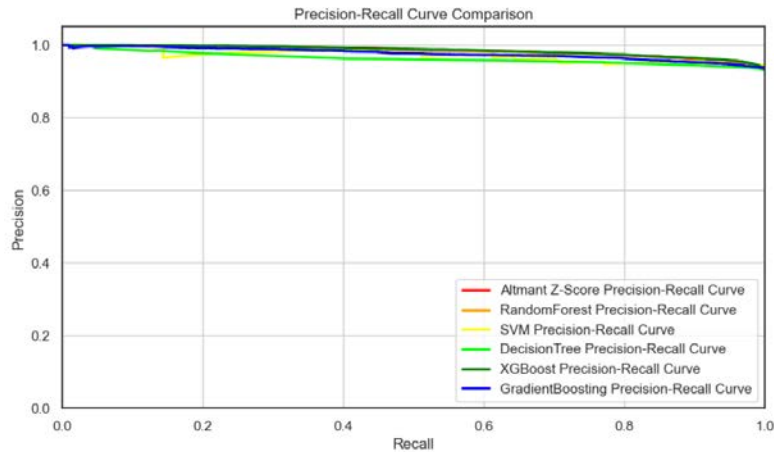


Рисунок 4. Precision-Recall криві побудованих моделей

За метрикою accuracy (рис. 2), всі моделі машинного навчання (RandomForest, SVM, DecisionTree, XGBoost, GradientBoosting) показали дуже високі та схожі результати, що коливаються в діапазоні ~ 0.93 – 0.94 . Це значно перевищує базовий показник Altman Z-Score (~ 0.88). Модель SVM демонструє незначну перевагу (~ 0.94).

При аналізі F1 Score (рис. 2), яка є більш показовою для даної задачі, SVM виразно лідирує з показником приблизно 0.96. RandomForest, XGBoost та GradientBoosting також мають дуже високі результати (~ 0.95). Це свідчить про те, що SVM найкраще збалансувала виявлення компаній-банкрутів, не створюючи при цьому надмірної кількості помилкових тривог.

Однак, аналіз ROC AUC (рис. 2) показує кардинально іншу картину. XGBoost є беззаперечним лідером з показником AUC ~ 0.85 , що вказує на його потужну здатність до розрізнення класів. RandomForest (~ 0.82) та, що цікаво, традиційний Altman Z-Score (~ 0.83) також демонструють високу ефективність. Натомість, SVM (~ 0.69) та DecisionTree (~ 0.67) показують значно слабші результати, що лише трохи краще за випадкове вгадування.

5. АРХІТЕКТУРА СППР

Архітектуру СППР наведено на рисунку 5.

СППР складається з наступних модулів:

- **Веб-застосунок.** Це загальна оболонка системи, з якою взаємодіє користувач.
- **Вікно: Введення даних.** Це вікно має одну дію: «Вибір файлу вхідних даних», воно відповідає за завантаження датасету.
- **Вікно: Візуалізація даних.** Користувач може обрати для перегляду графік. Доступні графіки: співвідношення збанкрутілих компаній до живих, матриця кореляцій та викиди.
- **Вікно: Підготовка вхідних даних.** У цьому вікні користувач може обрати модель для якої підготувати дані, а саме обрати ключові ознаки та зберегти оброблені дані у окремому файлі.
- **Вікно: Навчання моделей.** Це основне вікно системи, де після обробки даних стає доступна можливість навчання моделей. Як результат виводяться ключові метрики та їх графічне відображення.
- **Вікно: Візуалізація кінцевих результатів.** Вікно для співставлення метрик навчених моделей у вигляді графіків на вибір користувача. Вигляд доступних графіків наведений на рисунках 2-4.
- **Вікно: Прогнозування результатів.** Надає користувачеві можливість обрання однієї з навчених моделей та введення вхідних даних для прогнозування ризику банкрутства.

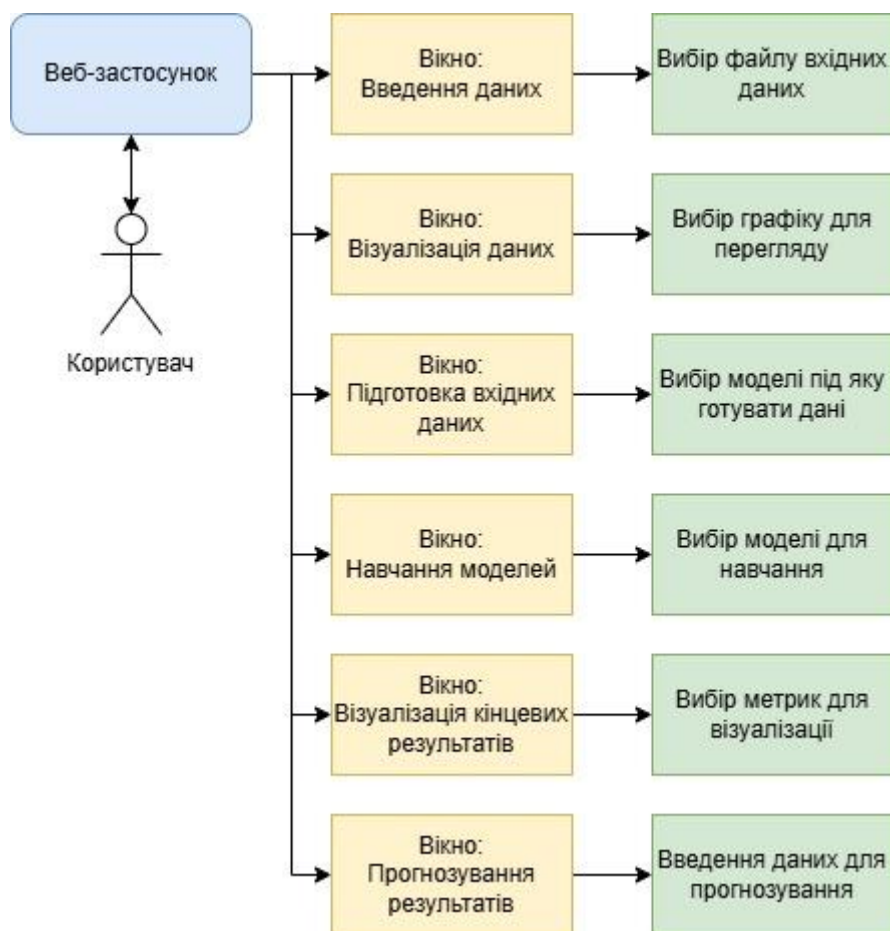


Рисунок 5. Архітектура СППР

6. ВИСНОВКИ

У даній роботі розглянуто можливості застосування моделей машинного навчання для вирішення задачі аналізу економічних ризиків та прогнозування банкрутства компаній. Було проведено порівняльний аналіз низки алгоритмів, включаючи DecisionTree, RandomForest, SVM, GradientBoosting та XGBoost, із застосуванням методу SMOTE для вирішення проблеми незбалансованості вхідних даних. Проведений аналіз і отримані результати підтверджують доцільність використання цих моделей: зокрема, XGBoost продемонстрував найвищу розрізнявальну здатність (ROC AUC ~ 0.85), а SVM – найкращий баланс точності та повноти (F1 Score ~ 0.96). Крім того, було спроектовано архітектуру системи підтримки прийняття рішень (СППР), що базується на даних моделях, у вигляді веб-застосунку з продуманою модульною структурою, яка забезпечує повний цикл роботи з даними та робить СППР ефективним інструментом для ризик-менеджерів.

У рамках подальших досліджень доцільно розглянути інші методи класифікації, зокрема нейронні мережі, які можуть покращити точність прогнозів за рахунок виявлення складніших нелінійних залежностей у фінансових показниках. Перспективним напрямком є також поглиблене впровадження методів інтерпретації, як-от SHAP, безпосередньо в архітектуру СППР, що дозволить надавати користувачам не лише прогноз, але й його миттєве пояснення. Такий підхід додасть новий рівень прозорості в оцінювання ризиків та допоможе фінансовим установам підвищити довіру до автоматизованих систем аналізу.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Python for Finance and Economics - Medium, 2025. URL: <https://medium.com/@quanticascience/python-for-finance-and-economics-baec27219da1> (дата звернення: 15.11.2025)
2. Chen, T., & Guestrin, C. (2016). XGBoost: A Scalable Tree Boosting System. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 785-794.
3. Lemaître, G., Nogueira, F., & Aridas, C. K. (2017). Imbalanced-learn: A Python Toolbox to Tackle the Curse of Imbalanced Datasets in Machine Learning. Journal of Machine Learning Research, 18(17), 1-5.
4. Hastie, T., Tibshirani, R., Friedman, J. The Elements of Statistical Learning: Data Mining, Inference, and Prediction. (2009). Сполучені Штати Америки: Springer.
5. Bishop, C. M. (2006). Pattern Recognition and Machine Learning. Швейцарія: Springer New York.

ПОВЕДІНКОВІ МОДЕЛІ ОЦІНЮВАННЯ УСПІШНОСТІ СТУДЕНТІВ У СМАРТ-КЛАСАХ

Ткач В.С.¹, Кузнєцова Н.В.²

Національний технічний університет України
«Київський політехнічний інститут ім. Ігоря Сікорського», Київ, Україна

¹ tkach.viktoria@lil.kpi.ua, ² n.kuznietsova@kpi.ua

У роботі досліджено поведінкові дані студентів у смарт-класах та представлено підхід до прогнозування їхньої залученості й ризику відтоку (відвідуваності) з класу. На основі набору даних OULAD побудовано моделі машинного навчання та моделі виживання, які дають змогу оцінювати ключові фактори успішності, відстежувати часову динаміку активності та своєчасно виявляти студентів із підвищеним ризиком припинення навчання.

Ключові слова: смарт-клас, поведінкові дані, машинне навчання, моделі виживання, прогнозування відтоку.

1. ВСТУП

Смарт-клас – це сучасне навчальне середовище, у якому поєднано цифрові технології та педагогічні підходи. Завдяки збору даних щодо активності студентів, їх результати та взаємодію з матеріалами та навчальними ресурсами стає можливим створення індивідуальних, орієнтованих на конкретних студентів, планів навчання та підвищення якості засвоєння знань. У період активної цифровізації такі класи особливо важливі, адже вони забезпечують гнучкість і допомагають швидко реагувати на потреби кожного студента.

Важливою частиною дослідження у даній сфері є аналіз поведінкових даних. Показники активності, участі, взаємодії та оцінювання можуть вказувати на рівень залученості, який в свою чергу впливає на ризик відтоку. Уміння прогнозувати ці процеси робить смарт-клас корисним інструментом для прийняття своєчасних рішень. Актуальність теми зумовлена тим, що сьогодні навчальні заклади мають потребу забезпечувати вищу якість освітнього процесу в умовах зростання вимог до результатів навчання і залучати для цього сучасні цифрові технології, які роблять навчальний процес більш інтерактивним. Для цього необхідно своєчасно виявляти студентів, які можуть втратити зацікавленість або зіткнутися з труднощами у навчанні. Саме тому поведінкових аналіз і прогнозування відтоку стають важливими інструментами підтримки студентів та прийняття своєчасних педагогічних рішень.

2. ПОСТАНОВКА ЗАДАЧІ ДОСЛІДЖЕННЯ

Метою даного дослідження є аналіз поведінкових даних студентів та прогнозування їхньої залученості й ризику відтоку у смарт-класі з використанням методів машинного навчання та моделей виживання. У роботі застосовуються різні підходи класифікації й аналізу часу до настання певної події-триггеру, проводиться їх порівняння та оцінювання точності для визначення найбільш ефективних методів прогнозування навчальної активності студентів. Такий підхід дозволяє визначити чинники, що впливають на успішність, та своєчасно виявляти студентів із підвищеним ризиком припинення навчання. Об'єкт дослідження – процес поведінкової та навчальної активності студентів у смарт-класі.

Предмет дослідження – моделі машинного та моделі виживання, що використовуються для прогнозування залученості та ризику відтоку студентів.

3. МЕТОДИ МАШИННОГО НАВЧАННЯ ТА МОДЕЛІ ВИЖИВАННЯ

3.1. Методи машинного навчання

У дослідженнях, пов'язаних із класифікацією та прогнозуванням, широко застосовуються різноманітні методи машинного навчання, які дозволяють моделювати взаємозв'язки між ознаками та визначити належність об'єктів до певних категорій. У даній роботі реалізовано п'ять підходів: логістичну регресію, дерево рішень, градієнтний бустинг, наївний байєсівський класифікатор та метод опорних векторів із RBF-ядром.

Логістична регресія – це статистичний метод бінарної класифікації, який моделює ймовірність належності об'єкта до певного класу за допомогою логістичної функції. Вона дозволяє оцінювати вплив кожної ознаки на результат [1].

Дерево рішень – метод, який формує ієрархічну структуру з послідовності умов «якщо – то». Модель на основі дерева рішень розбиває простір ознак на підмножини, будуючи логічні правила, які забезпечують зрозумілу та наочну класифікацію [2].

Градієнтний бустинг – ансамблевий метод, у якому набір простіших моделей, наприклад, дерев рішень, додається послідовно, кожна з яких коригує помилки попередньої. Завдяки поетапному покращенню метод здатен відтворювати складні закономірності [3].

Наївний байєсівський класифікатор ґрунтується на теоремі Байєса та припускає незалежність ознак між собою. Попри простоту і деяку наївність даний метод часто демонструє стабільну роботу та використовується для швидкої і надійної класифікації [4].

Метод опорних векторів з RBF-ядром – це алгоритм, який шукає оптимальну межу між класами, максимізуючи відстань між ними. Використовуючи радіально-базисної функції дозволяє моделювати нелінійні залежності та ефективно розділяти дані на класи у багатовимірному просторі [5].

3.2. Моделі виживання

Моделі виживання застосовуються для аналізу часу до настання події та дозволять працювати з цензурованими спостереженнями. У дослідженні використано три підходи: модель Каплана-Майєра, модель пропорційних ризиків Кокса та модель прискореного часу до події Вейбула.

Модель Каплана-Майєра є непараметричною моделлю, яка оцінює функцію виживання й показує частку об'єктів, що залишаються активними у різні моменти часу. Вона будує ступінчасту криву, що дозволяє порівнювати різні групи за динамікою виживання [6].

Модель пропорційних ризиків Кокса описує, як змінні впливають на ризик настання події. Основна формула має вигляд:

$$h(t|x) = h_0(t)\exp(\beta x),$$

де $h(t|x)$ – ризик настання події, $h_0(t)$ – базовий ризик, βx – лінійна комбінація коефіцієнтів і коваріант. Експонента коефіцієнтів $\exp(\beta)$ інтерпретується як hazard ratio, тобто у скільки разів змінюється ризик при зміні ознаки [7].

Модель прискореного часу до події Вейбулла напряду моделює тривалість до події. Вона ґрунтується на рівнянні:

$$\ln(T) = \beta x + \sigma W,$$

де T – час до події, β – коефіцієнти, σ – параметр масштабу, а W – випадкова складова залежно від вибраного розподілу. У цій моделі $\exp(\beta)$ інтерпретується як hazard ratio, який показує, як змінюється тривалість часу до події за зміни певних характеристик [8].

Представлені підходи охоплюють як лінійні так і не лінійні способи моделювання даних, що забезпечує комплексне дослідження різних типів залежностей. Сукупне використання цих методів створює теоретичне підґрунтя для подальшого аналізу та інтерпретації поведінкових процесів у навчальному середовищі.

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для аналізу поведінкових даних студентів було використано відкритий навчальний набір даних OULAD [9], який містить інформацію щодо активності здобувачів освіти, їхні демографічні характеристики, результати оцінювання. Попередня обробка даних включала очищення даних, узгодження таблиць та перетворення записів у формат який придатний для подальшого моделювання. Окремо було сформовано підмножини даних для класифікаційних методів і моделей виживання.

На наступному етапі побудовано та протестовано моделі машинного навчання й моделі виживання відповідно до описаних раніше підходів. Для кожного методу розраховано ключові метрики точності класифікації, які відображають здатність моделі коректно розмежовувати студентів на групи «відрахований» та «не відрахований». Узагальнене порівняння результатів наведено у таблиці 1, де подано значення метрик macro-F1, accuracy та ROC-AUC для всіх моделей.

Таблиця 1. Порівняння результатів моделей класифікації

Модель	Macro-F1	Accuracy	ROC-AUC
Логістична регресія	0.9362	0.9361	0.9804
Гradientний бустинг	0.9393	0.9392	0.9852
Дерево рішень	0.9187	0.9185	0.9185
Метод опорних векторів (RBF)	0.9408	0.9407	0.9811
Наївний байєсівський класифікатор	0.8673	0.8664	0.9443

Аналіз результатів показав, що найкращі значення основних показників точності продемонструвала модель методу опорних векторів із RBF-ядром. Gradientний бустинг і логістична регресія також показали високу якість класифікації, тоді як дерево рішень і наївний байєсівський класифікатор показали нижчі результати.

На основі результатів прогнозування, отриманих методами машинного навчання, було визначено ключові характеристики студентів, найбільш схильних до відрахування. Аналіз показав, що проблемні студенти мають суттєво нижчу активність у віртуальному навчальному середовищі: рідше заходять у навчальну систему, виконують менше завдань і раніше припиняють участь у курсі. Також виявлено відмінність у соціально-демографічних ознаках, такі студенти частіше мають нижчий рівень освіти та проживають у менш забезпечених регіонах. Сукупність цих факторів формує узагальнений профіль студента з підвищеним ризиком відтоку.

Аналіз часу до відрахування студентів був проведений на основі трьох моделей виживання. Загальна крива Каплана-Майєра (рис. 1) показує, що найбільше відрахувань студентів відбувається саме на початку навчання, а саме у перші тижні функція виживання помітно зменшується, після цього зниження стає більш плавним, але до кінця періоду спостерігається, що ймовірність для студента залишитися на курсі поступово зменшується

приблизно до 0.4. Це свідчить про те, що критичний період ризику припадає на першу чверть навчання.

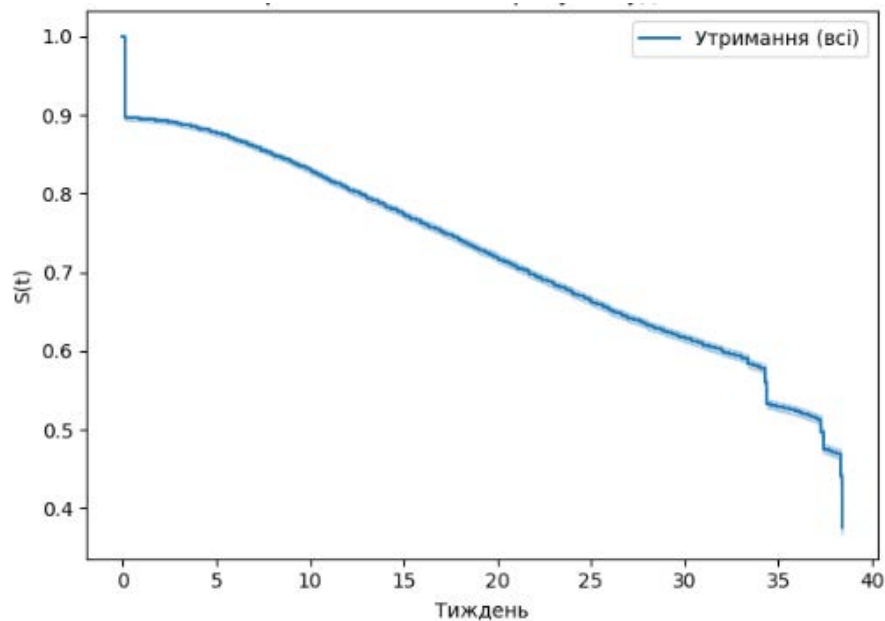


Рисунок 1. Крива Каплана-Майєра

Стратифіковані криві Каплана-Майєра за рівнем активності (рис. 2) демонструють чітку різницю у часовій динаміці між групами. Студенти з низькою кількістю кліків вибувають значно раніше і ця різниця формується вже з перших тижнів навчання. Натомість студенти з високою активністю залишаються на курсі довше. Це підкреслює те, що інтенсивність взаємодії з платформою прямо пов'язана з тривалістю навчання.

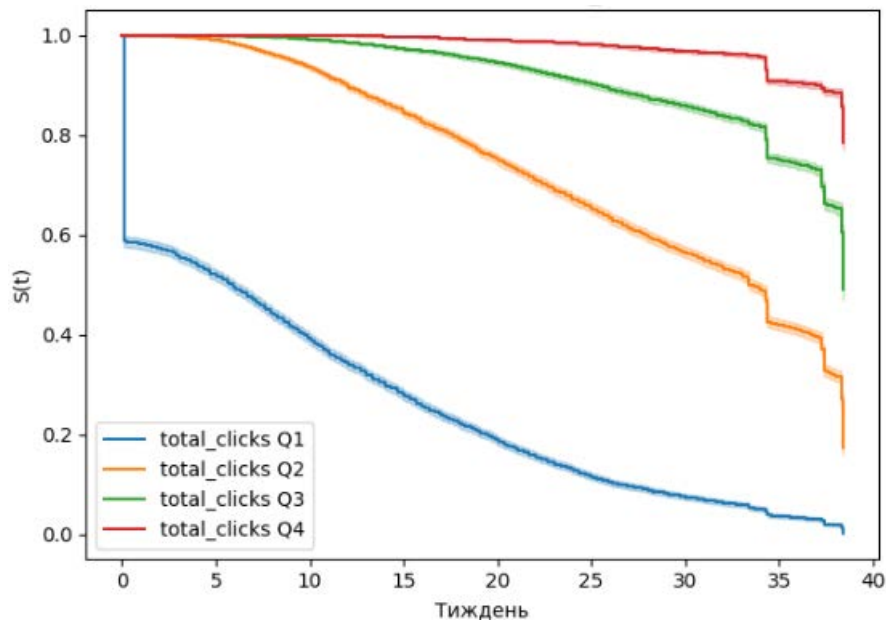


Рисунок 2. Стратифіковані криві Каплана-Майєра

На рисунку 3 наведено результати аналізу з використанням моделі пропорційних ризиків Кокса, яка дозволяє оцінити вплив окремих ознак на швидкість настання події. З графіка видно, що поведінкові показники: активні дні, кількість виконаних завдань і середній

бал зменшують ризик відрахування, тоді як окремі демографічні та регіональні характеристики пов'язані з більш раннім формуванням ризику. Це свідчить про те, що регулярна активність уповільнює настання події, тоді як соціальна вразливість її прискорює.

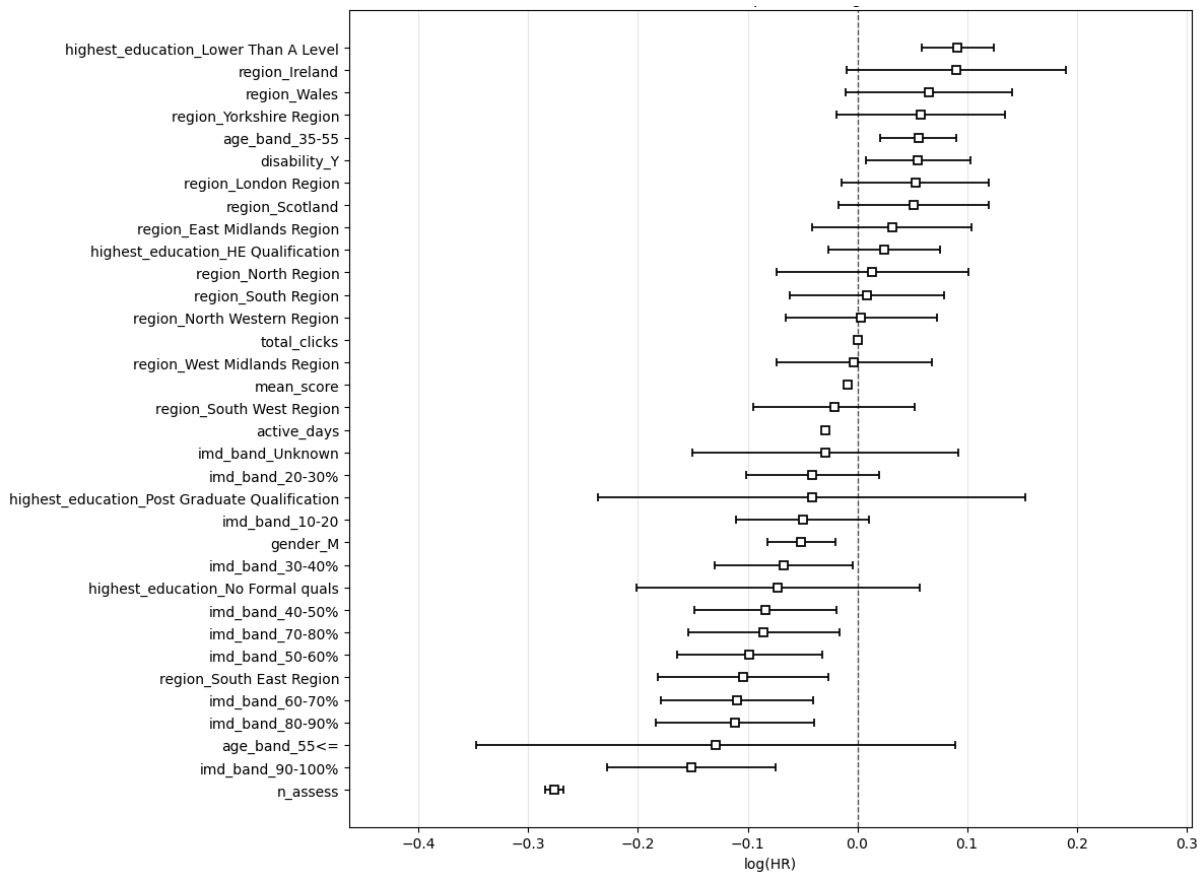


Рисунок 3. Коефіцієнти моделі Кокса

Прогнозні криві виживання отримані на основі моделі Вейбулла (рис. 4) показують часові відмінності між студентами різної активності. Студенти з низькою активністю мають суттєво короткий час до відрахування і вибувають приблизно до 25 тижня, тоді як активні студенти демонструють високу ймовірність залишитися на курсі.

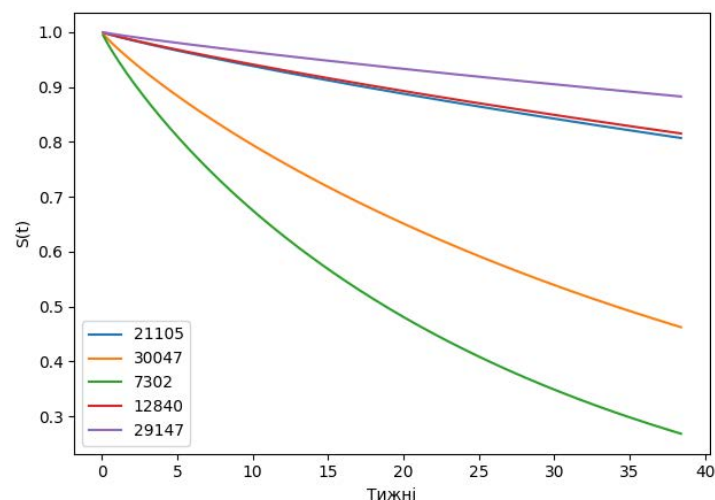


Рисунок 4. Прогнозні криві виживання моделі Вейбулла

5. ВИСНОВКИ

У ході дослідження було проаналізовано поведінкові дані студентів у смарт-класі та побудовано моделі машинного навчання і моделі виживання для прогнозування залученості студентів та ризику їх відтоку. Методи класифікації показали високу точність у визначенні студентів із підвищеною ймовірністю відрахування, а найкращі результати продемонстрував метод опорних векторів із RBF-ядром. Отримані моделі дали змогу виявити ключові чинники ризику, серед яких визначальне значення мають показники активності, а саме частота входу в навчальну систему, кількість виконаних завдань, активність та рівень успішності.

Аналіз часу до події підтвердив, що критичний період ризику припадає на початок навчання. Крива Каплана-Майєра показала швидке зниження виживання у перші тижні, а стратифікація за активністю продемонструвала суттєві часові відмінності між групами студентів за кількістю кліків. Модель Кокса дала змогу оцінити вплив окремих на швидкість настання відрахування, тоді як модель Вейбулла показала, що студенти з низькою активністю вибувають приблизно до 25 тижня.

Отримані результати свідчать, що поведінкові характеристики студентів є ключовими показниками їхньої стабільності у навчанні. Застосування моделей машинного навчання та аналізу виживання у смарт-класах дозволяє своєчасно виявляти студентів із підвищеним ризиком відтоку, визначити критичні часові точки та підтримувати здобувачів освіти на ранніх етапах в критичні моменти.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. LaValley M. P. Statistical Primer for Cardiovascular Research. Logistic Regression. *Circulation*. 2008. Vol. 117. P. 2395-2399.
2. Song Y., Lu Y. Decision tree methods: applications for classification and prediction. *Shanghai Archives of Psychiatry*. 2015. Vol. 27, No. 2. P. 130-135.
3. Friedman J.H. Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*. 2021. Vol.29, No. 5. P. 1189-1232.
4. Peretz O., Koren M., Koren O. Naïve Bayes classifier – An ensemble procedure for recall and precision enrichment. *Engineering Applications of Artificial Intelligence*. 2024. Vol. 136, 108972. 12 p.
5. Ding X., Liu J., Yang F., Cao J., Random radial basis function kernel-based support vector machine. *Journal of the Franklin Institute*. 2021. Vol. 358, No. 18. P. 10121-10140.
6. D'Arrigo G., Leonardis D., Abd Elhafeez S., Fusaro M., Tripepi G., Roumeliotis S. Method to Analyse Time-to-Event Data: The Kaplan-Maier Survival Curve. *Oxidative Medicine and Cellular Longevity*. 2021. Vol. 2021, 2290120, 7 p.
7. Fox J., Weisberg S., Cox Proportional-Hazards Regression for Survival Data in R. *An R Companion to Applied Regression*. 2023. Vol. 3, 20 p.
8. Liu E., Liu R., Lim K., Using the Weibull Accelerated Failure Time Regression Model to Predict Time to Health Events. *Applied Biostatistics for the Health Science and Epidemiology*. 2023. Vol. 13, No. 14. 13041
9. OULAD: Open University Learning Analytics Dataset. URL: <https://www.kaggle.com/code/vjcalling/oulad-open-university-learning-analytics-dataset>

РОЗРОБКА ТА ДОСЛІДЖЕННЯ ПРОГНОЗНИХ МОДЕЛЕЙ У СИСТЕМАХ АЛГОРИТМІЧНОЇ ТОРГІВЛІ

Ярінко Б.Б.

Національний технічний університет України
«Київський політехнічний інститут ім. Ігоря Сікорського», Київ, Україна

yarinko.b@gmail.com

У роботі запропоновано підхід до аналізу та вдосконалення прогнозних моделей у алгоритмічній торгівлі. Використано методи машинного навчання та розроблено інструментарій автоматизованої діагностики. Отримані результати демонструють підвищення точності моделей і підтверджують практичну ефективність систематизованого аналізу.

Ключові слова: алгоритмічна торгівля, прогнозні моделі, машинне навчання, діагностика моделей, фінансові ринки.

1. ВСТУП

Застосування методів машинного навчання в системах алгоритмічної торгівлі є однією з найактуальніших та водночас найскладніших задач сучасної фінансової інженерії. На відміну від класичних областей, де дані є відносно стабільними, фінансові ринки характеризуються унікальними властивостями: низьким співвідношенням сигнал/шум, нестаціонарністю та гетероскедастичністю. Як зазначається у сучасних дослідженнях [1, 2], ці характеристики суттєво ускладнюють виявлення стійких закономірностей. Унаслідок цього стандартні підходи до побудови прогнозних моделей, навіть із використанням потужних алгоритмів, таких як градієнтний бустинг чи рекурентні нейронні мережі, часто демонструють низьку та нестабільну ефективність.

Основною проблемою, з якою стикаються дослідники та практики, є значна складність та трудомісткість процесу діагностики та вдосконалення розроблених моделей. Аналіз причин низької продуктивності зазвичай потребує глибоких експертних знань і ручного, інтуїтивного пошуку помилок, що є неефективним, погано масштабованим і схильним до упередженості. Існуючі інструментальні засоби та MLOps-платформи, такі як MLflow або QuantConnect, ефективно автоматизують процеси навчання, версіонування та розгортання моделей, однак вони не надають засобів для глибокої, автоматизованої діагностики їхніх внутрішніх недоліків. Виникає потреба у створенні спеціалізованого інструменту, який би міг виконувати роль «інтелектуального асистента» для дослідника, автоматизуючи процес аудиту моделей.

Це зумовлює необхідність розробки систематизованого підходу до діагностики прогнозних моделей, що ґрунтується на формалізації експертних знань у вигляді структурованої методології. Метою даної роботи є дослідження прогнозних моделей через призму такого підходу: спочатку емпірично виявити типові патерни помилок, а потім розробити та апробувати програмний інструментарій, що автоматизує їх виявлення та пропонує шляхи вдосконалення. У статті представлено результати експериментального дослідження набору прогнозних моделей та продемонстровано ефективність запропонованої методології на практичному прикладі їх ітеративного покращення.

2. ПОСТАНОВКА ЗАДАЧІ ДОСЛІДЖЕННЯ

Загальною задачею даної роботи є дослідження ефективності та виявлення систематичних недоліків у роботі поширених прогнозних моделей при їх застосуванні до фінансових часових рядів, а також розробка та апробація систематизованого підходу для їх цілеспрямованого вдосконалення.

Для досягнення поставленої мети необхідно виконати такі конкретні завдання:

1. Провести експериментальне дослідження набору базових прогнозних моделей, що належать до різних класів (лінійні, ансамблеві на основі дерев, рекурентні нейронні мережі).
2. На основі результатів експериментального аналізу, систематизувати виявлені недоліки у вигляді формалізованих «патернів проблем» та розробити концепцію Базу Знань, що структурує методи їх діагностики й усунення.
3. Спроекувати та реалізувати програмний інструментарій, який автоматизує процес діагностики моделей на основі розробленої Базу Знань
4. Продемонструвати практичну роботу розробленого інструментарію, застосувавши його для автоматизованого аналізу набору базових моделей та формування комплексного плану їх вдосконалення на основі згенерованих рекомендацій.
5. Провести фінальну верифікацію запропонованого підходу, порівнявши продуктивність базових та вдосконалених за допомогою інструментарію моделей на відкладеній тестовій вибірці для оцінки його практичної ефективності.

3. ОСНОВНІ МЕТОДИ

З метою проведення всебічного порівняльного дослідження було обрано чотири алгоритми, що представляють різні парадигми машинного навчання: Логістична регресія як базовий лінійний метод; Випадковий ліс як представник ансамблевих методів; Екстремальний градієнтний бустинг як високоефективна реалізація градієнтного бустингу; та LSTM як модель рекурентної нейронної мережі. Вибір цієї архітектури зумовлений її здатністю ефективно враховувати часові залежності, що, за даними літератури [3], часто забезпечує їй перевагу над традиційними алгоритмами у задачах фінансового прогнозування.

Ключовим елементом дослідження є застосування розробленої методології систематизованого аналізу, яка була реалізована у вигляді спеціалізованого програмного модуля. Ця методологія базується на використанні Базу Знань, що реалізована у вигляді документо-орієнтованої структури. База Знань містить формалізований опис типових проблем прогнозних моделей, систематизованих за класами, що стосуються даних, архітектури моделі та її оцінки. Кожне правило в базі має уніфіковану структуру, що включає об'єкт-детектор, який описує логіку виявлення проблеми, та об'єкт-коректор, що містить пріоритезовані рекомендації щодо її вирішення. Процес аналізу моделей забезпечується аналітичним ядром, яке послідовно викликає програмні реалізації детекторів для перевірки наявності «симптомів» кожної проблеми та генерує фінальний діагностичний звіт.

Методологія експерименту розроблена таким чином, щоб забезпечити об'єктивність та відтворюваність результатів. Дослідження проводилося на історичних денних даних ETH/USD. Весь набір даних було хронологічно розділено на тренувальну, валідаційну та тестову вибірки у співвідношенні 70%, 15% та 15% відповідно. Валідаційна вибірка використовувалася для проміжної оцінки моделей та їх діагностики, тоді як тестова вибірка залишалася недоторканою до фінального етапу. Для оцінки якості було обрано метрику F1-Score. Її використання є рекомендованим для задач із дисбалансом класів, оскільки, на відміну від Accuracy, вона забезпечує об'єктивнішу оцінку балансу між точністю та повнотою [5].

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

В рамках експериментального дослідження було проведено навчання та оцінку чотирьох базових моделей, здійснено їх діагностику за допомогою розробленого програмного інструментарію та проведено фінальну верифікацію покращених версій моделей.

На першому етапі було проведено аналіз продуктивності базових моделей на відкладеній тестовій вибірці. Результати, представлені в табл. 1, демонструють значну диференціацію в ефективності між різними класами алгоритмів. Моделі, що працюють з даними як з таблицею (Logistic Regression, Random Forest, XGBoost), показали точність, близьку до 50-52%, причому XGBoost продемонстрував найкращий F1-Score (0.479) серед цієї групи. Разом з тим, модель LSTM показала найвищу продуктивність за ключовою метрикою F1-Score (0.574). Отриманий результат узгоджується із загальною тенденцією, описаною в літературі [3], щодо вищої ефективності рекурентних архітектур при роботі зі складними послідовними даними. Це свідчить про її вищу здатність ідентифікувати позитивний клас (торговий сигнал), хоча й за рахунок більшої загальної кількості помилок. Тим не менш, той факт, що жодна з моделей не досягла стабільно високих показників, підтверджує вихідну гіпотезу про обмежену ефективність стандартного підходу до моделювання та необхідність глибокого аналізу помилок

Таблиця 1. Продуктивність базових моделей на тестовій вибірці

Модель	Accuracy	F1-Score
Logistic Regression	0.507	0.412
Random Forest	0.523	0.454
XGBoost	0.505	0.479
LSTM	0.483	0.574

Для глибокого та систематизованого аналізу причин низької продуктивності було застосовано розроблений в рамках даної роботи програмний інструментарій для автоматизованої діагностики моделей.

Розроблений програмний інструментарій є модульною аналітичною платформою, призначеною для автоматизації процесу діагностики та підтримки вдосконалення прогностичних моделей. Концептуальна архітектура системи, представлена на рис. 1, ілюструє взаємодію її ключових компонентів.

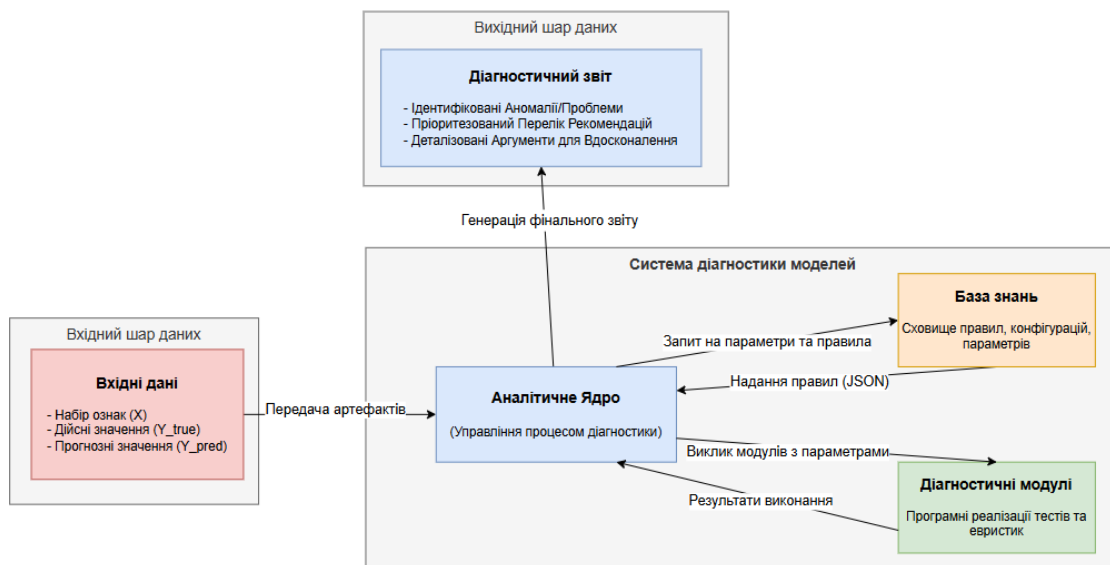


Рисунок 1. Концептуальна архітектура системи діагностики моделей

В основі архітектури лежать три основні компоненти:

1. Аналітичне ядро: Центральний компонент, що керує процесом аналізу. Він приймає на вхід артефакти роботи моделі (вхідні дані, прогнози, метрики), послідовно звертається до бази знань для отримання правил діагностики та викликає відповідні програмні модулі (детектори) для їх перевірки.

2. База знань: Серце системи, що реалізоване у вигляді гнучкої документо-орієнтованої структури. Вона містить формалізований опис типових проблем, з якими стикаються моделі в алгоритмічній торгівлі, та є легко розширюваною.

3. Модулі детекторів та коректорів: Це програмна реалізація логіки, закладеної в базі знань. Детектори є набором функцій, кожна з яких виконує специфічний тест для виявлення «симптомів» певної проблеми. Коректори містять логіку для генерації рекомендацій щодо вирішення виявлених проблем.

Така модульна архітектура дозволяє легко додавати нові правила діагностики, не змінюючи основну логіку системи, що забезпечує її гнучкість та масштабованість.

На другому етапі було проведено автоматизовану діагностику кожної з чотирьох моделей за допомогою цього інструментарію. Для візуалізації та підтвердження закономірностей, виявлених програмним інструментарієм, було проведено сегментований аналіз продуктивності моделей, результати якого представлено на рис. 2.

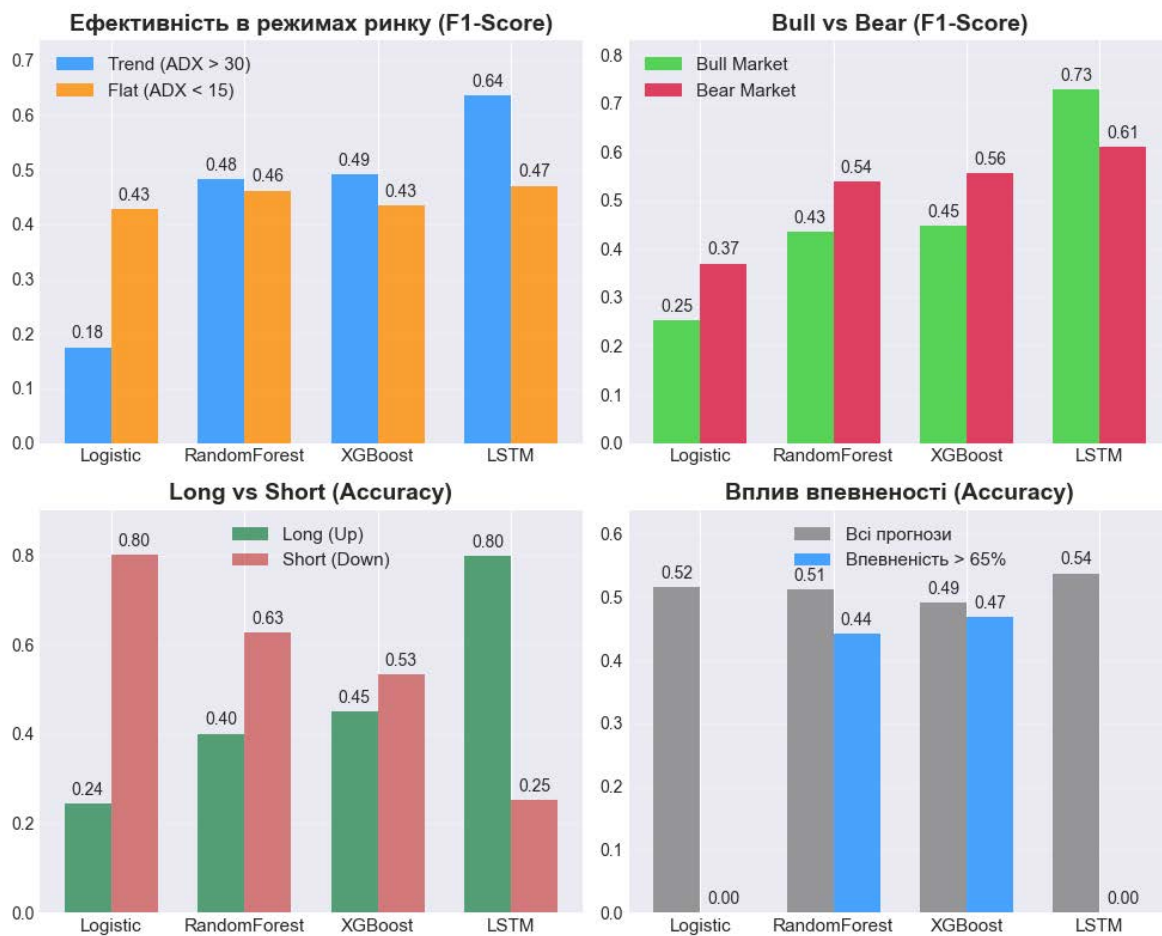


Рисунок 2. Аналіз продуктивності моделей у різних ринкових режимах та умовах

Аналіз звітів системи та візуалізацій на рисунку дозволив виявити кілька ключових, спільних для більшості моделей, патернів неефективності. По-перше, спостерігається чітка залежність від динаміки ринку, де лінійна модель краще працює у боковому русі з показником F1-Score 0.43 проти 0.18 у тренді, тоді як LSTM досягає найвищої ефективності 0.64 саме на трендових ділянках. По-друге, зафіксовано асиметрію щодо фаз ринку, де нейромережа єдина демонструє домінування на ринку, що зростає, зі значенням метрики 0.73. По-третє, аналіз напрямку угод виявив полярну спеціалізацію, при якій класичні моделі прогнозують сигнали на продаж з точністю до 0.80, а LSTM показує аналогічну високу точність 0.80 виключно у сигналах на купівлю. По-четверте, фільтрація прогнозів за високим рівнем впевненості виявилася неефективною, призвівши до зниження точності випадкового лісу з 0.51 до 0.44 та виявивши повну відсутність впевнених сигналів у інших алгоритмів. Окрім цього, автоматизований аналіз виявив низку інших спільних недоліків: надлишковість ознак, перенавчання для ансамблевих моделей, автокореляцію помилок та високі методологічні ризики. Ці спостереження доводять, що поверхневий аналіз метрик є недостатнім. Відсутність стандартизованих протоколів оцінювання, на яку вказують дослідники [3], вимагає застосування глибокої діагностики для виявлення справжніх причин неефективності моделей.

На основі узагальненого аналізу діагностичних звітів було сформульовано єдиний, комплексний план вдосконалення, спрямований на вирішення найбільш критичних та поширених проблем. Цей план включав три основні напрямки: інжиніринг ознак (створення комбінованих індикаторів для зменшення мультиколінеарності та додавання ознаки

ринкового режиму на основі ADX), оптимізацію гіперпараметрів (проведення пошуку по сітці для боротьби з перенавчанням деревних моделей) та вдосконалення торгової стратегії (запровадження фільтра сигналів за рівнем «впевненості» прогнозу). Використання індикатора ADX є загальноприйнятою практикою для кількісної ідентифікації сили тренду та розмежування ринкових фаз [6]. Всі чотири моделі були перенавчені з урахуванням цих комплексних корекцій.

Фінальна верифікація покращених моделей на тестовій вибірці (табл. 2) показала надзвичайно показові та позитивні результати, що підтверджують ефективність запропонованої методології. Застосування комплексного плану вдосконалення, сформованого на основі діагностичних звітів, призвело до значного покращення ключової метрики F1-Score для всіх без винятку моделей.

Таблиця 2. Порівняння продуктивності моделей до та після застосування рекомендацій

Модель	F1-Score (Baseline)	F1-Score (Final)	Зміна F1
Logistic Regression	0.412	0.676	+64.1%
Random Forest	0.454	0.550	+21.1%
XGBoost	0.479	0.527	+10.0%
LSTM	0.574	0.672	+17.1%

Найбільш драматичне зростання продемонструвала Логістична регресія, F1-Score якої зріс на 64.1%, досягнувши найвищого значення серед усіх моделей (0.676). Це свідчить про те, що для лінійних моделей правильний інжиніринг ознак та усунення базових проблем даних є критично важливим і здатний кардинально змінити їхню прогностичну силу. Модель LSTM, яка була лідером на базовому етапі, також показала суттєве зростання на 17.1%, що підтверджує ефективність додавання нових контекстних ознак навіть для рекурентних архітектур. Ансамблеві моделі Random Forest та XGBoost також продемонстрували впевнене покращення, що вказує на успішність боротьби з перенавчанням та оптимізації їхньої структури.

Ці результати однозначно демонструють, що систематичний, керований даними підхід до аналізу та вдосконалення, реалізований в розробленому програмному інструментарії, є значно ефективнішим, ніж стандартне застосування моделей. Запропонована методологія дозволила не просто покращити метрики, а й підвищити якість та надійність прогнозів для широкого спектру алгоритмів, що підтверджує її практичну цінність.

5. ВИСНОВОК

У ході дослідження було розроблено та апробовано систематизований підхід до аналізу та вдосконалення прогностичних моделей для задач алгоритмічної торгівлі. Проведені експерименти підтвердили, що стандартне застосування поширених алгоритмів машинного навчання демонструє обмежену ефективність, проте глибокий аналіз їхніх помилок виявляє наявність спільних, систематичних недоліків, таких як чутливість до ринкових режимів, інформаційна надлишковість ознак та процедурні ризики.

Наукова новизна роботи полягає у розробці концепції та прототипу програмного інструментарію, що базується на формалізованій Базі Знань для автоматизованої діагностики моделей. Такий підхід дозволяє перейти від ручного, інтуїтивного аналізу до стандартизованого та відтворюваного процесу аудиту прогностичних систем, що підвищує об'єктивність та швидкість дослідницького циклу.

Практична значущість дослідження підтверджена результатами фінальної верифікації, які продемонстрували, що застосування рекомендацій, згенерованих розробленим інструментарієм, призвело до статистично значущого покращення прогнозної здатності для всіх протестованих моделей, включаючи лінійні, ансамблеві та рекурентні нейронні мережі. Це доводить, що запропонована методологія є дієвим інструментом для створення більш робастних та ефективних прогнозних моделей.

Напрямами подальших досліджень є розширення Баз Знань новими, більш специфічними для різних класів моделей правилами, а також інтеграція розробленого інструментарію з MLOps-платформами для повної автоматизації життєвого циклу розробки торгових стратегій.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Wang Z., Ventre C. A Financial Time Series Denoiser Based on Diffusion Model. arXiv preprint arXiv:2409.02138. 2024. URL: <https://arxiv.labs.arxiv.org/html/2409.02138> (дата звернення: 10.11.2025)
2. Zhang L., Hua L. Major Issues in High-Frequency Financial Data Analysis: A Survey of Solutions. Mathematics. 2025. Vol. 13, No. 3. P. 347. URL: <https://www.mdpi.com/2227-7390/13/3/347> (дата звернення: 15.11.2025)
3. Giantsidi S., Tarantola C. Deep Learning for Financial Forecasting: A Review of Recent Trends. SSRN Electronic Journal. 2025. URL: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5263710 (дата звернення: 15.11.2025)
4. Wang T. Automated Model Testing and Monitoring: The Bedrock of Startup MLOps. Medium. 2024. URL: <https://medium.com/@tommyadeliyi/automated-model-testing-and-monitoring-the-bedrock-of-startup-mlops-1061c0b78e2e> (дата звернення: 15.11.2025)
5. F1 Score in Machine Learning Explained. Encord. 2023. URL: <https://encord.com/blog/f1-score-in-machine-learning/> (дата звернення: 20.11.2025)
6. ADX: The Trend Strength Indicator. Investopedia. URL: <https://www.investopedia.com/articles/trading/07/adx-trend-indicator.asp> (дата звернення: 20.11.2025)

СИСТЕМА ОЦІНЮВАННЯ І ПРОГНОЗУВАННЯ ВРОЖАЙНОСТІ ЗЕРНОВИХ КУЛЬТУР МЕТОДАМИ МАШИННОГО НАВЧАННЯ НА ОСНОВІ СУПУТНИКОВИХ І КЛІМАТИЧНИХ ДАНИХ

Ярко А.Ю.¹, Кузнєцова Н.В.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ a.yarko03@gmail.com

Дослідження присвячене розробленню інтелектуальної системи оцінювання і прогнозування врожайності культур на основі супутникових та кліматичних даних. Використано моделі машинного навчання, оптимізацію гіперпараметрів і валідацію LOYO. Отримано високу точність прогнозів та виявлено ключові фактори впливу, що дозволяють прогнозувати врожайність різних культур. Новизна роботи полягає в поєднанні гетерогенних даних і стійких алгоритмів, практична значимість – доцільність застосування такої системи під час страхування та управління аграрними ризиками.

Ключові слова: машинне навчання, врожайність, ризики, супутникові дані, прогнозування, аграрний сектор.

1. ВСТУП

Сучасні аграрні системи стикаються зі зростаючою невизначеністю, пов'язаною з кліматичними змінами, деградацією ґрунтів та економічною волатильністю ринків. Це актуалізує необхідність застосування інтелектуальних технологій для прогнозування врожайності та оцінювання ризику збитків. Традиційні статистичні методи часто не здатні адекватно моделювати складні нелінійні взаємозв'язки між екологічними та кліматичними змінними. Саме тому методи машинного навчання стають ефективним інструментом для аналізу великих гетерогенних даних.

У роботі представлено інформаційну технологію прогнозування врожайності пшениці в Україні на основі супутникових (MODIS), кліматичних (ERA5-Land) та картографічних даних (Copernicus EU), а також проведено порівняльний аналіз ефективності шести моделей машинного навчання [1, 2].

Метою дослідження є розробити інтелектуальну систему прогнозування врожайності та оцінювання аграрних ризиків, що об'єднує модулі збору даних, аналітичної обробки, моделювання, оцінки якості та формування рекомендацій для страхування і управління аграрними ризиками.

2. ВХІДНІ ДАНІ ТА ПОПЕРЕДНЯ ОБРОБКА

У дослідженні використано багатоджерельні дані з різною просторовою та часовою роздільною здатністю:

- індекси рослинності NDVI/EVI, біофізичні параметри (LAI, FPAR), температура поверхні (LST) з MODIS [1];
- кліматичні змінні ERA5-Land (t2m, tp, swvl1, swvl2, ssr);

- растрова карта культур CoprMap EU для ідентифікації площ вирощування;
- адміністративні межі областей України.

Розроблено трирівневу схему обробки:

1. Фільтрація растрових даних і створення TIF-файлів;
2. Генерація вибірових точок із розрахунком коефіцієнтів crop_ratio та overlap_ratio;
3. Обчислення середньозважених місячних агрегатів для кожної області.

Кінцевий датасет є панельною таблицею з понад 50 ознаками для кожного регіону та місяця вегетаційного сезону, як показано на рис. 1.

region_name	year	month	t2m	stl1	stl2	ssr	tp	swvl1	swvl2	LST_Night	EVI	LAI	LST_Day	NDVI	FPAR
Poltavska	2010.00	3.00	276.60	277.36	277.25	10635102.67	0.00	0.35	0.35	266.31	0.07	0.00	276.11	0.19	0.02
Poltavska	2010.00	4.00	283.70	284.10	283.51	14755818.60	0.00	0.30	0.30	277.56	0.17	0.02	296.07	0.35	0.08
Poltavska	2010.00	5.00	291.03	291.54	290.35	17502033.46	0.00	0.23	0.23	284.90	0.32	0.05	300.62	0.51	0.18
Poltavska	2010.00	6.00	296.48	297.52	296.28	19755509.42	0.00	0.22	0.21	272.47	0.39	0.07	301.78	0.57	0.24
Poltavska	2010.00	7.00	298.67	299.42	298.38	19227009.77	0.00	0.26	0.26	292.41	0.36	0.09	304.37	0.53	0.25
Poltavska	2010.00	8.00	300.29	302.20	301.42	19000049.07	0.00	0.14	0.16	288.34	0.19	0.04	303.26	0.34	0.17
Poltavska	2010.00	9.00	291.92	293.76	294.19	13108775.58	0.00	0.18	0.16	279.12	0.15	0.02	227.97	0.32	0.09
Poltavska	2010.00	10.00	281.78	283.14	284.32	6930946.01	0.00	0.31	0.29	272.87	0.14	0.01	269.92	0.33	0.08

Рисунок 1. Вигляд датасету

3. МЕТОДИ МОДЕЛЮВАННЯ

Для прогнозування річної врожайності пшениці (ц/га) обрано шість регресійних моделей [3, 4]:

- Метод опорних векторів;
- Випадковий ліс;
- Методи градієнтного та екстримального градієнтного бустингу;
- Дерево рішень;
- Метод k-найближчого сусіда.

Оптимізація гіперпараметрів здійснювалася методом багесівської оптимізації, а оцінювання моделей – за допомогою крос-валідації Leave-One-Year-Out (LOYO), що повністю усуває витік часових даних. Принцип роботи LOYO показано на рис. 2.



Рисунок 2. Метод крос-валідації LOYO

Для кількісної оцінки використано метрики: RMSE, MAE, MAPE, RRMSE, MSE.

4. РЕЗУЛЬТАТИ ЕКСПЕРИМЕНТІВ

Всі моделі продемонстрували конкурентоспроможну точність, проте їх якість суттєво відрізняється за стабільністю та здатністю відтворювати пікові значення врожайності. В таблиці 1 показано порівняння результатів прогнозування річної врожайності для всіх моделей.

Таблиця 1. Порівняльна таблиця якості прогнозів річної врожайності різними моделями

Модель	RRMSE	RMSE	MAE	MSE	MAPE
XGBoost	10.70167	4.530092	3.636955	20.52173	8.969499
KNN	11.19026	4.736914	3.794581	22.43836	9.763631
DT	13.44328	5.690634	4.567624	32.38332	10.94073
GB	10.97346	4.64514	3.690874	21.57733	9.124724
RF	10.81002	4.575956	3.630003	20.93938	8.908937
SVM	10.28598	4.354125	3.389671	18.9584	8.122624

Найвищу точність прогнозування врожайності продемонструвала модель опорних векторів (SVM), яка забезпечила найменші значення помилок RMSE, MAE та RRMSE, що свідчить про її здатність ефективно відтворювати як рівень, так і динаміку врожайності. Модель випадкового лісу (RF) виявилася найбільш стабільною в динаміці (для прогнозування за різними роками) та показала високу узагальнювальну здатність у регіональному розрізі, що робить її придатною для практичного застосування в умовах змінного клімату. Методи градієнтного бустингу (GB) та екстремального градієнтного бустингу (XGBoost) також продемонстрували високу точність і добре моделюють нелінійні залежності між ознаками, хоча їх результати виявилися дещо менш стабільними у періоди (роки) з аномальними погодними умовами. Методи k-найближчого сусіда (KNN) та дерев рішень (DT) істотно поступаються за точністю ансамблевим методам і методу на основі опорних векторів (SVM) через недостатню здатність узагальнювати дані та підвищену чутливість до локальних особливостей вибірки. Загалом, моделі опорних векторів і випадкового лісу забезпечують оптимальний баланс між точністю, стійкістю та інтерпретованістю, а тому є найбільш доцільними для задачі прогнозування врожайності на основі супутникових та кліматичних даних.

У межах проведених експериментів усі ансамблеві моделі (Random Forest, Gradient Boosting, XGBoost) продемонстрували дуже схожу структуру важливості ознак, що є підтвердженням узгодженості результатів. Найвищу прогностичну силу мали три групи параметрів:

- вологість ґрунту (swvl1, swvl2);
- індекси рослинності (NDVI, EVI, FPAR);
- опади (tp).

Для прикладу на рисунку 3 показано топ-20 важливих ознак для моделі випадкового лісу.

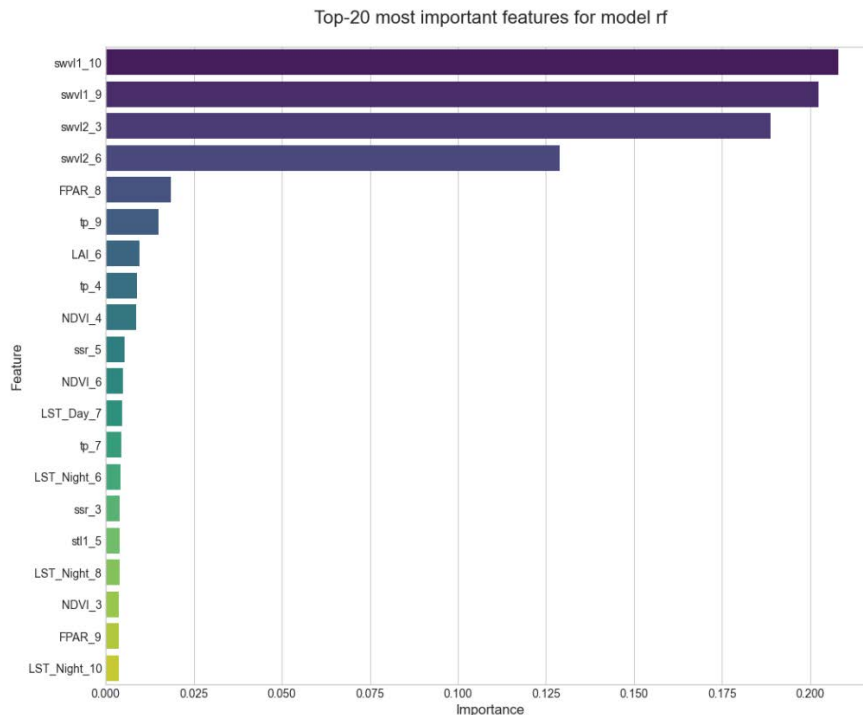


Рисунок 3. Важливість ознак

Показники об'ємної вологості ґрунту (особливо swvl1_10, swvl1_9, swvl2_3) стабільно входили у топ-10 ознак у всіх ансамблевих моделях, що підтверджує критичну залежність культур зернової групи від водного забезпечення на різних етапах розвитку [5].

- swvl1 (перший шар ґрунту, 0–7 см) відображає короткострокову доступність води для проростання і ранніх фаз вегетації.
- swvl2 (7–28 см) відіграє ключову роль у фазах активного наростання листкової маси, накопичення азоту та формування колосу.

Динаміка важливості ознак показує, що період вересень–жовтень попереднього року (swvl1_9, swvl1_10) має несподівано сильний вплив на урожай наступного сезону. Це пояснюється вологісним режимом ґрунту після збирання та перед зимівлею, що формує стартові умови для озимих культур. Це узгоджується з агрономічними дослідженнями, де саме достатній рівень води восени визначає кількість продуктивних стебел навесні [6].

Другу за значущістю групу становлять супутникові індекси NDVI/EVI та біофізичні параметри LAI і FPAR:

- NDVI добре корелює зі щільністю рослинного покриву.
- EVI краще реагує на зміни у густих посівах та при наявності аерозолів.
- FPAR є індикатором активності фотосинтезу – фактично описує “зелену енергію” рослин.

У моделях RF та XGBoost FPAR_5, NDVI_6 та LAI_4 входили у топ-20 ознак, особливо для періоду активного формування біомаси навесні.

Це означає, що система не лише враховує кліматичні параметри, але й оцінює фактичний фізіологічний стан рослин, що підвищує точність прогнозу.

Проведений аналіз похибок прогнозування врожайності (RMSE, RRMSE, MAPE) на рівні областей України показав значні просторові відмінності як показано на рис. 4.

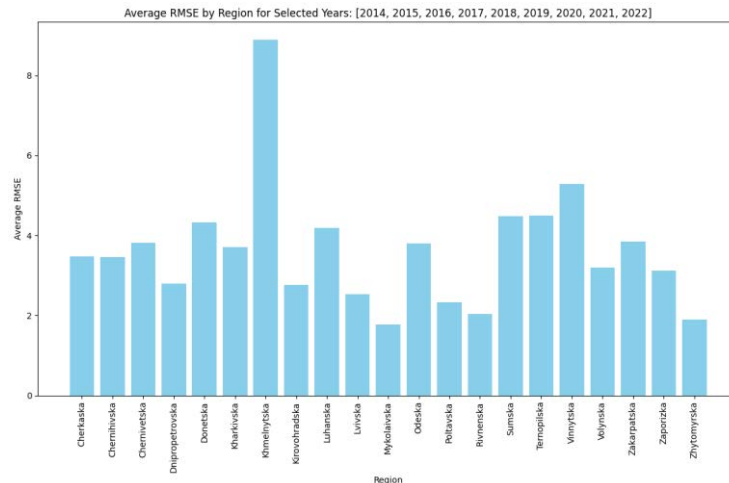


Рисунок 4. Порівняння похибок для різних областей

Різниця у точності прогнозів між областями зумовлена неоднорідністю природно-кліматичних умов, різною структурою ґрунтів, відмінностями в агротехнологіях та варіативністю якості статистичних даних. Регіони зі стабільними кліматичними умовами й однорідними ґрунтами (Львівська, Закарпатська, Житомирська області) забезпечують моделі більш передбачувані закономірності, що знижує похибку. Натомість території з високою внутрішньою гетерогенністю, складною рельєфною структурою, різко змінними погодними умовами або неповними даними створюють більшу варіативність у врожайності, через що моделі гірше узагальнюють залежності та демонструють вищу помилку прогнозування.

На рисунку 5 відображено динаміку зміни середньоквадратичної помилки (RMSE) для шести алгоритмів машинного навчання протягом 2011–2023 років. Графік демонструє, що всі моделі реагують на міжрічні коливання врожайності подібним чином, проте амплітуда похибок суттєво різниться. Дерево рішень (DT) стабільно показує найвищі значення середньоквадратичної помилки, що свідчить про слабку здатність моделі до узагальнення. Алгоритми градієнтного бустингу, випадкового лісу та екстремального градієнтного бустингу демонструють більш плавні зміни та нижчий рівень похибок, підтверджуючи їх перевагу в моделюванні нелінійних залежностей.

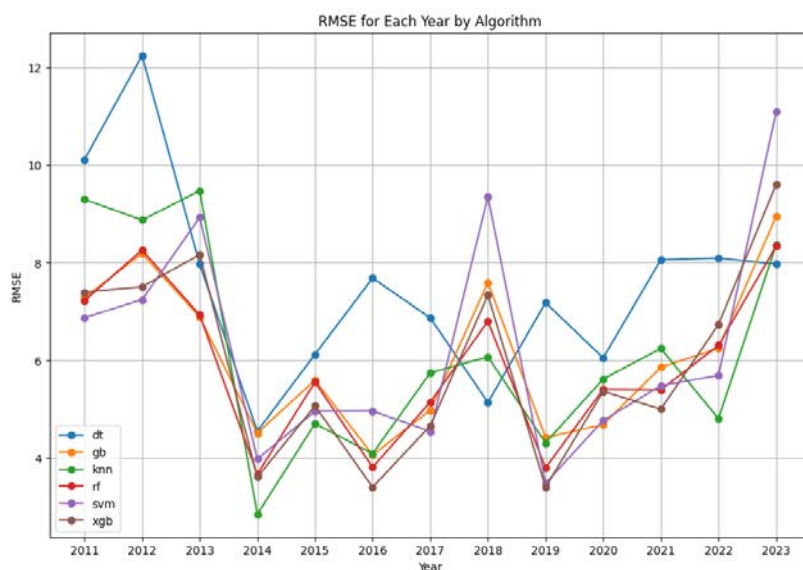


Рисунок 5. Динаміка змін похибок за роками

Модель на основі методу опорних векторів у більшості років показує найменший RMSE, проте водночас вона найбільш чутлива до “аномальних” років – наприклад, різке зростання похибки спостерігається у 2018 та 2023 роках. Це вказує на те, що метод опорних векторів добре відтворює типові закономірності, але важче адаптується до різких погодних змін. Натомість випадковий ліс та екстремальний градієнтний бустинг демонструють кращу стійкість – їхні RMSE менше реагують на екстремальні роки, що робить ці моделі більш надійними для практичного прогнозування врожайності протягом тривалих періодів.

5. ВИСНОВКИ

Результати дослідження показали, що застосування методів машинного навчання дає змогу з високою точністю прогнозувати врожайність пшениці на основі інтегрованих супутникових та кліматичних даних. Найкращу точність демонструє модель на основі методу опорних векторів, тоді як випадковий ліс забезпечує найвищу стабільність результатів у міжрічному та регіональному розрізі. Ансамблеві моделі градієнтного бустингу та екстремального градієнтного бустингу також показали добру продуктивність, підтверджуючи здатність ефективно відтворювати нелінійні залежності між ознаками. Загалом, інтеграція різних джерел даних і використання сучасних алгоритмів дозволили створити надійну систему оцінювання показників урожайності в умовах високої кліматичної мінливості.

Аналіз важливості ознак засвідчив ключову роль вологості ґрунту, індексів рослинності та опадів у формуванні точності моделей, а оцінка регіональних відмінностей підтвердила суттєву просторову неоднорідність агрокліматичних умов в Україні. Найменші похибки спостерігаються в областях із більш передбачуваними кліматичними умовами, тоді як регіони з високою внутрішньою гетерогенністю та різкими погодними коливаннями характеризуються значно вищою варіативністю результатів. Це підкреслює важливість подальшої адаптації моделей до локальних особливостей та розвитку регіонально орієнтованих систем прогнозування.

Прогнозована врожайність використана для кількісної оцінки страхових ризиків шляхом розрахунку ймовірності дефіциту врожаю, очікуваного збитку, страхової премії, VaR та Expected Shortfall. Це забезпечує пряме практичне застосування моделі у страхуванні сільськогосподарських культур та підвищує точність оцінки ризику порівняно з традиційними статистичними підходами.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. C. Justice, E. Vermote, J. Townshend, et al., “An overview of MODIS Land data processing and product status,” *Remote Sensing of Environment*, vol. 83, no. 1–2, pp. 3–15, 2002.
2. H. Hersbach, B. Bell, P. Berrisford, et al., “The ERA5 global reanalysis,” *Quarterly Journal of the Royal Meteorological Society*, vol. 146, no. 730, pp. 1999–2049, 2020.
3. A. Bégué, D. Arvor, B. Bellon, et al., “Remote sensing and crop monitoring: A review,” *Remote Sensing*, vol. 10, no. 7, pp. 1–32, 2018. DOI: <https://doi.org/10.3390/rs10071032>
4. L. Breiman, “Random Forests,” *Machine Learning*, vol. 45, pp. 5–32, 2001.
5. J. H. Friedman, “Greedy function approximation: A gradient boosting machine,” *Annals of Statistics*, vol. 29, no. 5, pp. 1189–1232, 2001.
6. S.-J. Lee, E. Sohn, M. Kim, K.-H. Park, K. Park, and Y. Lee, “Real-time retrieval of daily soil moisture using IMERG and GK2A satellite images with NWP and topographic data: A machine learning approach for South Korea,” *Remote Sensing*, vol. 15, no. 18, pp. 1–20, 2023.

ДЕЦЕНТРАЛІЗОВАНА СИСТЕМА КОНТРОЛЮ ЛАНЦЮГІВ ПОСТАВОК ГУМАНІТАРНОЇ ДОПОМОГИ НА ОСНОВІ ТЕХНОЛОГІЇ БЛОКЧЕЙН

Андрушко А.І.¹, Гіоргізова-Гай В.Ш.²

¹ andrushko.andry@lil.kpi.ua, ² high.victoria@lil.kpi.ua

У статті досліджується розробка децентралізованої системи контролю ланцюгів поставок гуманітарної допомоги на основі технології блокчейн. Метою дослідження є вирішення проблем корупції, відсутності прозорості та затримок у доставці гуманітарних ресурсів. Використано методи аналізу існуючих блокчейн-рішень та проектування консорціумного блокчейну з механізмом консенсусу Proof-of-Authority. Розроблено п'ятирівневу архітектуру системи зі смарт-контрактами для автоматизації процесів та багаторівневою системою безпеки. Створено прототип на .NET 9.0, який забезпечує прозорість, зниження корупції та автоматизацію бюрократичних процесів.

Ключові слова: блокчейн, гуманітарна допомога, ланцюги поставок, смарт-контракти, Proof-of-Authority, децентралізована система.

1. ВСТУП

Традиційні централізовані системи управління ланцюгами поставок гуманітарної допомоги стикаються з критичними викликами. За оцінками, до 30–50% гуманітарних ресурсів втрачається через корупцію [1] та брак прозорості. Відсутність надійних механізмів відстеження призводить до затримок у доставці критично необхідних ресурсів.

Основні проблеми включають: можливість підробки документів, обмежену видимість для донорів, складність швидкого виявлення крадіжок, затримки в розслідуванні інцидентів (тижні або місяці).

Існуючі блокчейн-рішення орієнтовані на комерційний сектор і вимагають значних ресурсів, що робить їх непридатними для локальних організацій з обмеженими бюджетами.

Аналіз останніх досліджень і публікацій. Аналіз вже реалізованих проєктів в гуманітарних ланцюгах поставок показав ефективність застосування в них блокчейну [2]. Проєкт Building Blocks від WFP продемонстрував скорочення адміністративних витрат на 98% [3]. Ініціатива Oxfam показала зменшення втрат на 30-50% [4]. Проєкт Start Network скоротив час реакції на кризи на 30% [5].

У рамках дослідження децентралізованої системи для контролю в ланцюгу поставок гуманітарної допомоги було проведено консультацію з представниками Міжнародного Комітету Червоного Хреста (ICRC) в Києві. ICRC виявила ряд специфічних вимог: баланс між прозорістю та конфіденційністю, робота в офлайн-режимі, інтеграція з існуючими системами, простий інтерфейс.

Публічні блокчейни не забезпечують конфіденційності [6]. Приватні створюють ризики централізації. Комерційні рішення вимагають значних інвестицій [7].

2. ДЕЦЕНТРАЛІЗОВАНА СИСТЕМА КОНТРОЛЮ ЛАНЦЮГІВ ПОСТАВОК ГУМАНІТАРНОЇ ДОПОМОГИ НА ОСНОВІ ТЕХНОЛОГІЙ БЛОКЧЕЙН

Пропонується побудова системи на основі консорціумного блокчейну з механізмом консенсусу Proof-of-Authority [8] для підвищення прозорості, надійності та ефективності контролю в ланцюгах поставок гуманітарної допомоги.

Архітектура системи. П'ятирівнева архітектура: презентація (Blazor PWA, MAUI), API-шар (REST, gRPC, JWT), бізнес-логіка (сервіси управління), блокчейн-рівень (Transaction Manager, Smart Contract Runtime, PoA Engine, P2P), дані (розподілений реєстр, PostgreSQL, Redis).

Ролі учасників системи. Система підтримує чотири типи ролей з розмежованими повноваженнями: Донор – надає фінансування через Donation Escrow Contract, кошти утримуються в ескроу до підтвердження доставки; Замовник – постачальник товарів, який може надавати продукцію авансом на основі репутаційної системи або після верифікації донорського фінансування; Координатор – організація, яка керує закупівлями через Procurement Contract після підтвердження фінансування та організовує логістику; Отримувач – бенефіціар допомоги, який проходить верифікацію через КЕП (Кваліфікований Електронний Підпис) у Delivery Verification Contract для підтвердження особи та отримання вантажу. Всі дії учасників записуються в блокчейн для аудиту та забезпечення прозорості.

Механізм консенсусу. Proof-of-Authority з рандомним призначенням валідаторів. Пропускна здатність 500-1000 tps, латентність 3-5 секунд, Byzantine Fault Tolerance до 1/3 валідаторів.

Смарт-контракти. Shipment Tracking Contract для автоматичного відстеження життєвого циклу поставок на основі GPS-даних; Donation Escrow Contract для управління донорськими коштами із автоматичним утриманням до завершення поставки; Procurement Contract для автоматизації закупівель координатором після надходження донорських коштів; Delivery Verification Contract для верифікації отримувача через Кваліфікований Електронний Підпис (КЕП) та підтвердження доставки; Penalty Enforcement Contract для автоматичних штрафів при порушеннях умов [9, 10].

Багаторівнева безпека. На рівні мережі застосовується TLS 1.3 для REST API та mTLS для міжвузлової gRPC-комунікації. На рівні автентифікації використовується комбінація JWT-токенів, багатофакторної автентифікації та ECDSA-підписів для блокчейн-транзакцій. Приватні ключі зберігаються зашифрованими через AES-256 та розшифровуються лише локально. На рівні даних реалізовано RBAC/ABAC для гранулярного контролю доступу, чутливі дані зберігаються зашифрованими, а в блокчейні фіксуються лише SHA-256 хеші. На рівні блокчейну криптографічні хеші зв'язують блоки в незмінний ланцюг, а Consensus Engine забезпечує Byzantine Fault Tolerance до 1/3 валідаторів.

Технології та інструментальні засоби проектування: .NET 9.0 з C# для кросплатформенності; консорціумний блокчейн з PoA (на 70% менше енергоспоживання); гібридна архітектура; смарт-контракти на C#; Blazor PWA для офлайн-режиму; інтеграція з IoT.

На рисунку 1 можна побачити повний шлях посилки під час відправлення допомоги в системі.

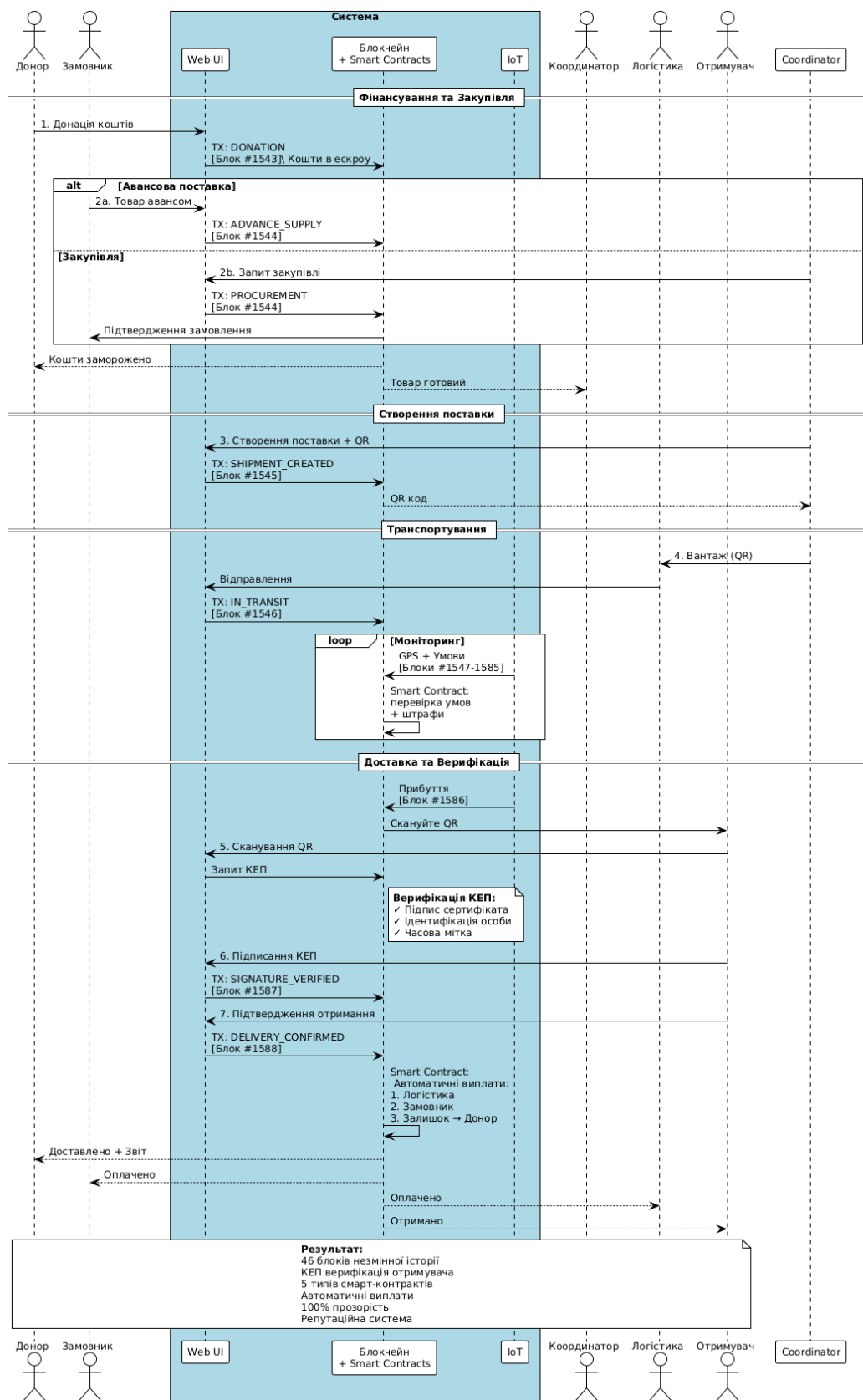


Рисунок 1. Пайплайн відправлення гуманітарної допомоги

3. ПРАКТИЧНЕ ЗАСТОСУВАННЯ ТА ПЕРЕВАГИ ЗАПРОПОНОВАНОГО РІШЕННЯ

Прозорість і довіра. Донори можуть відстежувати вантажі в режимі реального часу без розкриття конфіденційних даних про отримувачів.

Зниження корупції. Незмінність блокчейну робить неможливою підробку документів. Всі зміни логуються для аудиту.

Автоматизація. Смарт-контракти зменшують бюрократичні затримки з тижнів до хвилин.

4. ВИСНОВКИ

Спроектowana система орієнтована на локальні організації з обмеженими бюджетами і може працювати в умовах обмеженої інфраструктури.

Розроблено повнофункціональний прототип на .NET 9.0 з консорціумним блокчейном та PoA, який включає гібридну архітектуру, три типи смарт-контрактів, веб-інтерфейс з офлайн-режимом та криптографічний захист.

Перспективи розвитку: пілотний проект, розширення функціональності, інтеграція з міжнародними платформами, інтеграція з IoT та ERP, мобільні додатки.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Saberi S. Blockchain technology and sustainable supply chain // International Journal of Production Research. 2019. Vol. 57. No. 7. P. 2117-2135.
2. Casado-Vara R., Prieto J. How blockchain improves the supply chain // Procedia Computer Science. 2018. Vol. 134. P. 393-398.
3. WFP. Building Blocks: Blockchain for Zero Hunger. Rome: WFP, 2023.
4. Oxfam International. Blockchain for Good: Humanitarian Supply Chain Report. Oxford: Oxfam GB, 2022.
5. Start Network. Digital Innovation in Humanitarian Response: Annual Report 2023. London: Start Network, 2023.
6. Helo P., Hao Y. Blockchains in operations and supply chains // Computers & Industrial Engineering. 2019. Vol. 136. P. 242-251.
7. Cole R., Stevenson M. Blockchain technology: implications for operations // Supply Chain Management. 2019. Vol. 24. No. 4. P. 469-483.
8. Nakamoto S. Bitcoin: A peer-to-peer electronic cash system. 2008.
9. Kshetri N. Blockchain and sustainable supply chain management // International Journal of Information Management. 2021. Vol. 60.
10. Queiroz M.M. Blockchain and supply chain management integration // Supply Chain Management. 2020. Vol. 25. No. 2. P. 241-254.

РОЗРОБКА СЕРВІС-ОРІЄНТОВАНОЇ СИСТЕМИ КООРДИНАЦІЇ РОЮ ДРОНІВ НА ОСНОВІ АЛГОРИТМУ МУРАШИНИХ КОЛОНІЙ (АСО) ТА СТИГМЕРГІЙНОЇ ВЗАЄМОДІЇ

Будзинський А.В.¹, Булах Б.В.

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ kergontop@gmail.com

У дослідженні розглядається сервіс-орієнтована система координації рою дронів, що поєднує алгоритм мурашиних колоній і стигмергійну взаємодію. Архітектуру реалізовано у вигляді ROS2-вузлів: побудова та випаровування феромонної карти, відкладення феромонів і АСО-прийняття рішень; виконано інтеграцію з Gazebo/ArduPilot. Досягнуто візуалізації карти та продемонстровано консольні рішення АСО. Новизною даного дослідження є сервісна декомпозиція і безцентрова координація, а практична цінністю - масштабованість і стійкість до втрат зв'язку.

Ключові слова: АСО, стигмергія, роєва координація, ROS2, феромонна карта, Gazebo/ArduPilot.

1. ВСТУП

Координація рою безпілотних літальних апаратів (БПЛА) у динамічних та частково спостережуваних середовищах є актуальною науково-прикладною задачею. Вона ускладнюється обмеженою пропускну здатністю каналів зв'язку, їх перервністю, невизначеністю сенсорних даних, потребою швидкого масштабування та стійкості до відмов окремих агентів. Класичні централізовані підходи потребують стабільних каналів, глобальної оптимізації та часто деградують при збільшенні кількості дронів або втраті зв'язку, що мотивує перехід до розподілених методів координації.

Метою даного дослідження є розробка та експериментальна верифікація сервіс-орієнтованої системи координації рою БПЛА на основі алгоритму мурашиних колоній (АСО) та стигмергійної взаємодії, яка забезпечує масштабоване, відмовостійке й малозалежне від якості зв'язку узгодження дій агентів у реальному часі.

У цьому дослідженні було розроблено сервіс-орієнтовану систему (Service-Oriented Architecture, SOA) для координації рою дронів, що поєднує АСО зі стигмергією – непрямою координацією через “сліди” в середовищі (феромони). Рішення реалізовано на ROS 2, інтегровано з Gazebo/ArduPilot та демонструє візуалізацію феромонної карти й процес прийняття рішень у реальному часі.

2. ТЕОРЕТИЧНА СКЛАДОВА АЛГОРИТМУ АСО

АСО моделює поведінку мурах, які залишають феромони на корисних маршрутах. Інші агенти, зчитуючи ці “сліди”, імовірно обирають подальші кроки. У загальному вигляді правило переходу (оцінка привабливості) з поточного стану i до кандидата j визначається зваженою комбінацією феромону τ_{ij} та евристики η_{ij} [1]:

$$score_j = (\tau_{ij} + \epsilon)^\alpha * (\eta_{ij})^\beta, p_j = \frac{score_j}{\sum_{i=1}^k score_i}, \text{ де}$$

τ_{ij} – феромон (накопичений досвідом “якість” переходу $i \rightarrow j$). Чим більше, тим привабливіше,
 η_{ij} – евристика (миттєва корисність без досвіду). У даній роботі я використовував обернену відстань, що віддає перевагу близьким крокам при рівних феромонах,
 $\alpha \geq 0$ – вага феромону (ступінь експлуатації минулого досвіду),
 $\beta \geq 0$ – вага евристики (ступінь жадібності до локально кращого кроку),
 ϵ – незначне чисельне значення для уникнення нулів.

Після кожного кроку феромон оновлюється. Однією з найважливіших подій оновлення є випаровування феромону для уникнення локальних максимумів на карті, яке розраховується за формулою [2]:

$$\tau_{ij} = (1 - \rho) * \tau_{ij}, \text{ де}$$

τ_{ij} – значення феромону між кроками i та j ,
 ρ – швидкість “забування”. Чим більше це значення, тим швидше відбувається адаптація до змін, але зменшується пам’ять.

Як було зазначено вище, у якості евристики (η) було використано обернену відстань, що віддає перевагу близьким крокам при рівних феромонах. В залежності від задачі евристика може кодувати [3]:

- очікуваний приріст інформації
- енергетичну вартість маршруту
- безпекові / колізійні ризики

Формула обраної евристики (обернена відстань) має наступний вигляд:

$$\eta_j = 1/(d(i, j) + \epsilon).$$

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для реалізації поставленого завдання було розроблено декілька сервісів за допомогою мови програмування Python та бібліотеки ROS 2, серед яких:

- PheromoneMapNode – сервіс для отримання, застосування та зберігання феромонів на карті.
- PheromoneDepositNode – сервіс для періодичної публікації депозиту феромонів у поточній позиції.
- ACODecisionNode – сервіс ймовірнісного прийняття рішень на основі карти феромонів.
- ACONavigationBridge – сервіс передачі інформації з ACODecisionNode до дронів.
- ACOSVisualizer – утиліта візуалізації карти феромонів.

Після проведення перших тестових завдань мною була отримана та візуалізована карта феромонів, наведена на рис. 1.

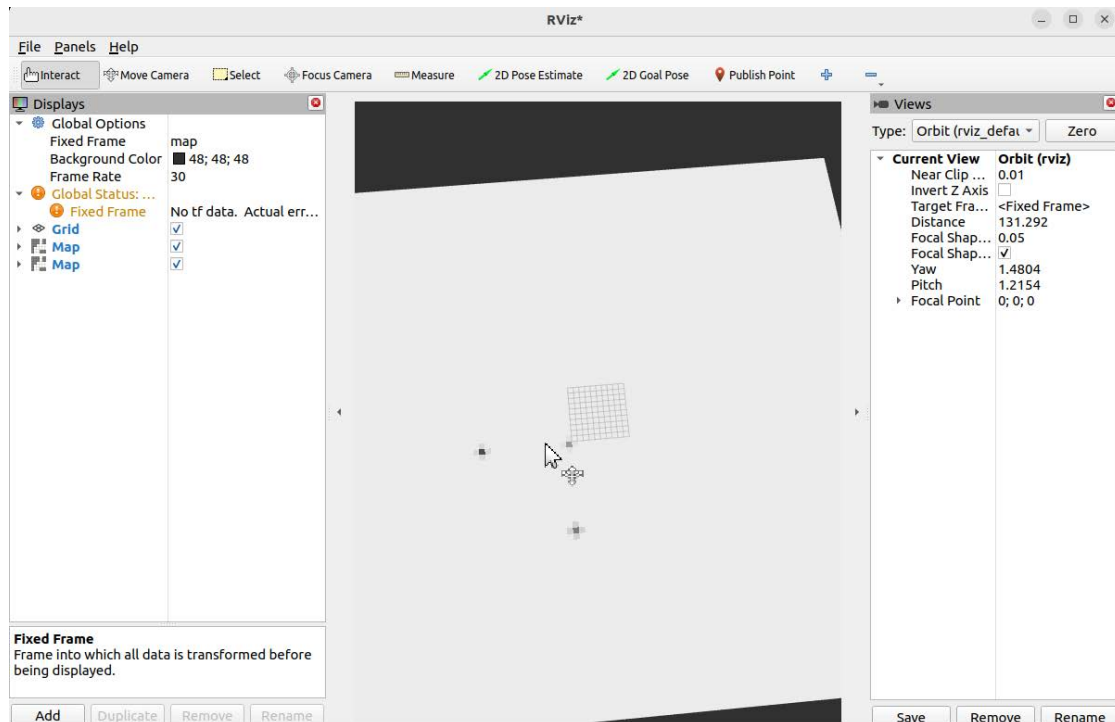


Рисунок 1. Приклад візуалізації карти феромонів

За умови ручного оновлення карти феромонів мною також був протестований створений мною сервіс прийняття рішень (ACODecisionNode) та логіка випаровування феромонів (рис. 2).

```
[INFO] [1763582030.551618187] [aco_decision_node]: From 4 → Next 2 (5.00, 5.00)
[INFO] [1763582031.592649563] [aco_decision_node]: From 2 → Next 4 (10.00, -5.00)
[INFO] [1763582032.550393293] [aco_decision_node]: From 4 → Next 2 (5.00, 5.00)
[INFO] [1763582033.546124116] [aco_decision_node]: From 2 → Next 4 (10.00, -5.00)
[INFO] [1763582034.576124551] [aco_decision_node]: From 4 → Next 2 (5.00, 5.00)
[INFO] [1763582035.545978100] [aco_decision_node]: From 2 → Next 4 (10.00, -5.00)
[INFO] [1763582036.554559491] [aco_decision_node]: From 4 → Next 2 (5.00, 5.00)
[INFO] [1763582037.546460651] [aco_decision_node]: From 2 → Next 4 (10.00, -5.00)
[INFO] [1763582038.592172615] [aco_decision_node]: From 4 → Next 1 (5.00, 0.00)
[INFO] [1763582039.554072164] [aco_decision_node]: From 1 → Next 4 (10.00, -5.00)
[INFO] [1763582040.567552825] [aco_decision_node]: From 4 → Next 0 (0.00, 0.00)
[INFO] [1763582041.556778630] [aco_decision_node]: From 0 → Next 4 (10.00, -5.00)
[INFO] [1763582042.553941176] [aco_decision_node]: From 4 → Next 0 (0.00, 0.00)
[INFO] [1763582043.556670735] [aco_decision_node]: From 0 → Next 4 (10.00, -5.00)
[INFO] [1763582044.550323224] [aco_decision_node]: From 4 → Next 3 (0.00, 5.00)
```

Рисунок 2. Приклад роботи алгоритму прийняття рішень

У даному прикладі мною було залишено 100 феромонів у точці 4 (10, -5), що автоматично зробило її більш привабливою за інші (на кожному етапі відбувається повернення

до точки 4). Від точки 4, де немає доступних пунктів призначення з високим рівнем феромонів, алгоритм за замовчуванням використовує евристику на основі відстані (η^β), створюючи ймовірнісну поведінку локального дослідження. Поточні параметри ($\alpha=1,0$, $\beta=2,0$) створюють сильне слідування за феромонами після їх відкладення, але це не є хорошим вибором при рівномірному розподілі феромонів. Це демонструє важливість налаштування параметрів для бажаної поведінки рою. Гаусівська функція розповсюдження створює реалістичні градієнти феромонів, причому сусідні комірки отримують $\sim 13,6\%$ від сили відкладення, що забезпечує більш плавні поля притягання феромонів.

4. ПОДАЛЬШІ ДОСЛІДЖЕННЯ

Основним напрямком подальших досліджень є інтеграція розроблених сервісів з Gazebo, запуск симуляції з реальними дронами. Далі буде проведено аналіз ваги феромонів (α) та ваги евристики (β) на реальних прикладах, а також реалізація подвійних феромонів (привабливий та відштовхувальний) для розведення ролей між дронами у рої та розв'язання конфлікту експлуатації / дослідження. Це має значно покращити покриття, час до виявлення цілей та стійкість рою дронів. Математична формула оцінки привабливості за такої реалізації буде виглядати наступним чином [4]:

$$score_j = (\tau_j^+ + \epsilon)^\alpha * \frac{1}{1 + k\tau_j^-} * (\eta_j)^\beta, \text{ де}$$

τ_j^+ – корисні феромони, які збільшують ймовірність вибору кандидату j ,

τ_j^- – штрафні феромони, які зменшують ймовірність вибору кандидату j .

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Ant colony optimization algorithm. URL: https://en.wikipedia.org/wiki/Ant_colony_optimization_algorithms (дата доступу 19.11.2025).
2. Introduction to Ant Colony Optimization. URL: <https://www.geeksforgeeks.org/machine-learning/introduction-to-ant-colony-optimization/> (дата доступу 19.11.2025).
3. Dorigo M, Stützle T. Ant Colony Optimization. Massachusetts Institute of Technology. 2004. P.33-38.
4. D. D. Haynes and S. M. Corns, "Algorithm for a Tabu — Ant Colony Optimizer," 2015 IEEE Congress on Evolutionary Computation (CEC), Sendai, Japan, 2015, pp. 529-535, doi: 10.1109/CEC.2015.7256935.
5. Gazebo Harmonic documentation. URL: <https://gazebo-sim.org/docs/harmonic/install/> (дата доступу 19.11.2025).

ДОСЛІДЖЕННЯ ВПЛИВУ ТОКЕНІЗАЦІЇ В МАЛОРЕСУРСНИХ МОВАХ НА ЯКІСТЬ ПЕРЕКЛАДУ

Головач А.А.¹, Кислий Р.В.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ annaholovachanna@gmail.com

Метою дослідження є вивчення впливу різних методів токенізації на якість перекладу українською мовою. Було досліджено алгоритми токенізації BPE, Unigram та character-level, а також здійснено адаптацію токенізатора для вже натренованої моделі. Отримані результати демонструють, що вибір токенізатора та розмір словника суттєво впливають на показники якості перекладу, а адаптація токенізатора може позитивно позначитися на цих показниках. Наукова новизна роботи полягає в порівнянні ефективності різних стратегій токенізації для української мови. Практична цінність дослідження полягає у формуванні рекомендацій щодо покращення машинного перекладу для малоресурсних мов.

Ключові слова: токенізація, машинний переклад, адаптація токенізатора, мало-ресурсні мови, NLP.

1. ВСТУП

Машинний переклад відіграє важливу роль у міжнародній комунікації та надає доступ до інформації. Однак, більшість перекладацьких систем орієнтовані на так звані високоресурсні мови, що мають великі корпуси даних. Натомість, малоресурсні мови, до яких належить і українська, стикаються з проблемою обмеженого обсягу доступних текстів для тренування моделей. Цей дефіцит даних безпосередньо впливає на якість перекладу, ускладнюючи розробку ефективних систем.

Одним із ключових аспектів, що впливають на якість машинного перекладу, є токенізація – процес розбиття тексту на менші значущі одиниці, або токени. Це особливо важливо для української мови, яка є морфологічно складною. На відміну від аналітичних мов (наприклад, англійської), де слова переважно мають фіксовану форму, в українській мові слова змінюються за відмінками, родами, числами та особами. Ці зміни (флексії) додають велику кількість унікальних словоформ, що робить традиційну токенізацію на рівні слів неефективною. Неправильна токенізація може призвести до втрати контексту, неточного перекладу та незрозумілих речень. Це робить вибір оптимального токенізатора критично важливим для покращення якості МП для української мови.

2. ОСНОВНІ МЕТОДИ

Токенізація – це важливий етап попередньої обробки тексту для підготовки вхідних токенів для глибинних мовних моделей. Вона визначає, як текст розбивається на токени, що впливає на здатність моделі оперувати рідкісними словами та морфологічно складними формами. Розглянемо три методи токенізації, які досліджувались. BPE (Byte-Pair Encoding). Цей метод починає з розбиття всіх слів на окремі символи, формуючи базовий словник. Далі

алгоритм поступово об'єднує пари символів, які найчастіше зустрічаються, у нові токени, повторюючи процес до досягнення бажаного розміру словника. BPE зазвичай створює токени на рівні підслів, забезпечуючи компроміс між символьним та словесним представленням. Метод добре працює з рідкісними словами, але може бути неоптимальним у використанні словникового простору. На відміну від BPE, Unigram базується на ймовірнісній моделі tokenів. Спочатку створюється великий базовий словник, що включає всі символи та популярні підрядки, а потім алгоритм поступово видаляє менш ймовірні токени, залишаючи найбільш релевантні. Під час токенизації нового тексту алгоритм вибирає найбільш ймовірну послідовність tokenів або може пропонувати декілька варіантів відповідно до їх ймовірності.

У алгоритмі Character-level (символьна токенизація) кожен символ тексту стає окремим токеном. Такий підхід забезпечує повне покриття словника та роботу з будь-якими словами, включно з рідкісними чи новоствореними. Недоліком є подовження послідовностей та збільшення обчислювальних витрат, але перевагою є універсальність для малоресурсних мов і нових лексем.

Для базових експериментів не використовувалися попередньо натреновані моделі, натомість здійснювалося тренування власних моделей середнього розміру з нуля за процедурою навчання MarianMT. Такий підхід забезпечував повний контроль над усіма етапами навчання та дозволяв адаптувати архітектуру під специфічні потреби дослідження.

В якості основи обрана MarianMT відкритий нейронний фреймворк для машинного перекладу, який є ефективним і гнучким рішенням для навчання моделей, зокрема для малоресурсних мов. Архітектура моделей базувалася на Transformer-base, запропонований Vaswani et al. (2017) у роботі "*Attention Is All You Need*". Модель складалася з двох основних компонентів – енкодера (encoder) та декодера (decoder), що взаємодіють через механізм multi-head attention.

Для проведення дослідження не використовувалися попередньо натреновані токенизатори; натомість було здійснено навчання власних токенизаторів за допомогою бібліотеки SentencePiece. Було натреновано сім токенизаторів: BPE на 8 000, 16 000 та 32 000 tokenів, unigram на 8 000, 16 000 та 32 000 tokenів і character-level на 138 tokenів. Для токенизації англійських текстів використовувався токенизатор моделі google-bert/bert-base-uncased, який реалізує алгоритм WordPiece.

Тренування токенизаторів здійснювалось на датасеті UberText-UA. Для BPE та Unigram навчання проводилось на вибірці 6 000 000 речень, для character-level на всьому датасеті. Покриття символів становило 0.999504, що забезпечувало високу якість токенизації, включно з рідкісними символами.

Для базового порівняння було натреновано сім моделей MarianMT середнього розміру, кожна з власним типом токенизації. Розроблений конвеєр обробки даних включав очищення вхідних даних від неприпустимих символів і шуму, токенизацію обраним токенизатором та подальше навчання моделі. Для BPE та Unigram тренування проводилось на трьох різних розмірах словника: 8 000, 16 000 та 32 000 tokenів. Character-level модель слугувала контрольним експериментом. Тренування моделей здійснювалось на вибірці 1 000 000 речень датасету HPLT.

Оцінка якості перекладу проводилася за допомогою трьох основних метрик. SacreBLEU – автоматизована версія BLEU для стандартизованого порівняння перекладу. chrF++ – оцінка перекладу на основі character n-gram для морфологічно складних мов. COMET – нейромережевий підхід, що забезпечує кореляцію з людською оцінкою перекладу. Використання кількох метрик дозволяє отримати комплексну оцінку продуктивності моделей на різних аспектах якості перекладу.

Для експериментів з перенесення токенизатора був розширений словник базового токенизатора моделі facebook/nllb-200-distilled-600M на 9 295 українських tokenів отриманих із власного токенизатора SentencePiece моделі BPE за словником 16 тис. tokenів.

Жорстка адаптація токенизатора відбувалась в три фази. Перша фаза жорсткої адаптації токенизатора, полягала у "прогріві" ембедінгів нових tokenів перед основним донавчанням моделі. Здійснено ембедінг ресайз, ембедінги нових tokenів ініціалізовувалися як середнє ембедінгів найближчих старих tokenів, що забезпечувало стабільний старт для навчання. Для першої фази сформовано датасет з паралельного англійсько-українського корпусу HPLT вибірці на 5 тис. речень, 84 з яких включали два нових токена, 4 916 по одному новому токenu. Модель навчалась на даному датасеті 2 епохи з низьким learning rate $1e-5$.

У другій фазі проводилось додаткове навчання моделі з розширеним токенизатором. Використовувалася підвибірка паралельного корпусу HPLT обсягом 10 800 речень, підібрані для того, щоб зосередитися на нових токенах та мінімізувати вплив на вже натреновані параметри моделі. Навчання відбувалось з дуже малим learning rate ($5e-6$) протягом однієї епохи.

Під час третьої фази навчання, тренування моделі відбувалось на паралельному англійсько-українському корпусі вибірці на 40 000 речень датасету HPLT. Тренування тривало протягом трьох епох з learning rate ($5e-6$).

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Після токенизації англійської послідовності було отримано 27 485 tokenів, для української послідовності кількість tokenів залежала від обраного методу та розміру словника. Таблиця 1 з кількістю tokenів демонструє, що збільшення розміру словника BPE або Unigram дозволяє покрити більшу кількість унікальних лексем і підслів. Також видно, що Character-level токенизація дає значно більшу кількість коротких tokenів для тієї ж кількості символів, що призводить до подовження послідовностей і збільшення обчислювальних витрат.

Таблиця 1. Кількість tokenів після токенизації українських послідовностей

Розмір словника	Кількість tokenів		
	BPE	unigram	Character-level
8 тис.	8 203	8 184	371
16 тис.	16 128	16 129	
32 тис.	31 941	31 944	

Збільшення словника з 8 до 32 тис. tokenів значно підвищує кількість покритих tokenів, що покращує здатність моделі відображати морфологічні й лексичні варіації української мови. Character-level токенизатор має малу кількість унікальних tokenів (138 для всього набору символів), проте генерує дуже довгі послідовності, що ускладнює тренування моделі і знижує ефективність перекладу.

Результати базових експериментів з методом токенизації та розміром словника наведені у таблиці 2.

Отримані результати показують, що тип токенизатора та розмір словника суттєво впливають на якість перекладу української мови. Моделі з BPE або Unigram словником на 8 тис. tokenів показали найнижчі результати через обмежене покриття словникового простору. Словники на 16 тис. tokenів дають кращі результати, але їхня здатність охоплювати морфологічні й лексичні варіації обмежена. Найвищі результати досягнуті з BPE та Unigram словниками на 32 тис. tokenів. Character-level токенизатор виявився найменш ефективним і не підходить для практичного перекладу. Для якісного перекладу української мови

рекомендується використовувати Unigram з 32 тис. tokenів на великих тренувальних датасетах.

Таблиця 2. Розрахунок метрик моделей

Модель, токенизатор	Метрика		
	SacreBLEU	Chrf++	COMET
Власна, BPE 8 тис.	23.9	47.1	0.7622
Власна, BPE 16 тис.	24.8	48.2	0.7664
Власна, BPE 32 тис.	25.1	48.5	0.7652
Власна, Unigram 8 тис.	23.6	46.7	0.7584
Власна, Unigram 16 тис.	25.0	48.2	0.7666
Власна, Unigram 32 тис.	25.3	48.6	0.7675
Власна, Character-level	1.0	12.1	0.4068

У таблиці 3 наведено результати адаптації токенизатора для попередньо натренованої моделі facebook/nllb-200-distilled-600M. У експерименті порівнювалися базова модель і модель після жорсткої адаптації токенизатора.

Таблиця 3. Результати адаптації токенизатора

Модель	Метрика		
	sacreBLEU	Chrf++	COMET
facebook/nllb-200-distilled-600M (базова)	21.96	47.91	0.8373
Жорстка адаптація	22.08	48.71	0.8286

Результати показують, що жорстка адаптація розширеного токенизатора дає невелике покращення показників SacreBLEU та Chrf++, проте спостерігається невелике зниження COMET. Це може свідчити про те, що зміна токенизатора дещо впливає на кореляцію з людською оцінкою перекладу, але дозволяє моделі точніше відображати частотні лексеми та морфологічні варіації в українській мові. Таким чином, токенизатор було успішно адаптовано, і модель готова до подальшого тренування на будь-якому англійсько-українському паралельному корпусі.

4. ВИСНОВКИ

Проведене дослідження показало, що вибір типу токенизатора та розміру словника має суттєвий вплив на якість машинного перекладу української мови. Найкращі результати за метриками досягнуті моделлю Unigram зі словником на 32 тис. tokenів, що забезпечує оптимальне покриття лексем та підслів, тоді як Character-level токенизатор виявився менш ефективним через значне подовження послідовностей та зростання обчислювальних витрат.

Крім того, проведена жорстка адаптація токенизатора базової моделі facebook/nllb-200-distilled-600M, включаючи розширення словника на українські токени та поступове “прогрівання” ембедінгів нових tokenів, показала невелике, але помітне покращення метрик перекладу. Це свідчить про те, що адаптація токенизатора дозволяє моделі точніше відображати частотні лексеми та морфологічні варіації української мови, незначно впливаючи на кореляцію з людською оцінкою перекладу (COMET).

Отримані результати також підкреслюють практичну цінність великого словника для моделей BPE та Unigram: збільшення розміру словника з 8 до 32 тис. tokenів суттєво підвищує кількість покритих tokenів, що дозволяє моделі краще справлятися з морфологічними й лексичними варіаціями. Character-level токенизатор, хоча і забезпечує повне покриття символів, виявився практично непридатним для перекладу.

Таким чином, для якісного машинного перекладу української мови рекомендується використовувати Unigram токенизатор зі словником 32 тис. токенів на великих тренувальних датасетах. Водночас проведені дослідження доводять, що адаптація токенизатора та розширення словника базової моделі є ефективними підходами для покращення продуктивності нейронних моделей перекладу, що особливо важливо для малоресурсних мов. Результати роботи демонструють, що правильний вибір токенизатора і оптимальний розмір словника є ключовими факторами підвищення якості машинного перекладу української мови.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. UberText UA: Український корпус текстів. URL: <https://lang.org.ua/uk/corpora/#anchor4> (дата звернення: 18.11.2025).
2. HPLT English–Ukrainian Parallel Corpus. URL: <https://opus.nlpl.eu/sample/en&uk/HPLT&v2/sample> (дата звернення: 18.11.2025).
3. The FLORES-101 Evaluation Benchmark for Low-Resource and Multilingual Machine Translation. URL: <https://ai.meta.com/research/publications/the-flores-101-evaluation-benchmark-for-low-resource-and-multilingual-machine-translation/> (дата звернення: 18.11.2025).
4. Marian NMT Documentation. URL: <https://marian-nmt.github.io/> (дата звернення: 18.11.2025).
5. Vaswani A., Shazeer N., Parmar N., Uszkoreit J., Jones L., Gomez A.N., Kaiser Ł., Polosukhin I. Attention Is All You Need. URL: <https://arxiv.org/abs/1706.03762> (дата звернення: 18.11.2025).
6. Datacamp. What is Tokenization? URL: <https://www.datacamp.com/blog/what-is-tokenization> (дата звернення: 18.11.2025).
7. Zhao W., Xie J., Tang S., Pan Y., Wang J. Tokenizer Choice For LLM Training: Negligible or Crucial? URL: <https://aclanthology.org/2024.findings-naacl.247/> (дата звернення: 18.11.2025).
8. Sennrich R., Haddow B., Birch A. Neural Machine Translation of Rare Words with Subword Units. URL: <https://arxiv.org/abs/1508.07909v5> (дата звернення: 18.11.2025).
9. Kudo T., Richardson J. SentencePiece: A simple and language independent subword tokenizer and detokenizer for Neural Text Processing. URL: <https://arxiv.org/abs/1808.06226v1>.
10. SentencePiece Documentation. URL: <https://pypi.org/project/sentencepiece/> (дата звернення: 18.11.2025).
11. Devlin J., Chang M. W., Lee K., Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. NAACL-HLT 2019. URL: <https://huggingface.co/google-bert/bert-base-uncased> (дата звернення: 18.11.2025).
12. UNLP 2025 Workshop Paper. URL: <https://aclanthology.org/2025.unlp-1.1.pdf> (дата звернення: 18.11.2025).
13. Sharma K. Expanding NLLB-200 to Kangri: A Step-by-Step Guide to Fine-Tuning Meta's Multilingual Model. URL: <https://medium.com/@karunsharma1920/expanding-nllb-200-to-kangri-a-step-by-step-guide-to-fine-tuning-metas-multilingual-model-for-an-885496835afc> (дата звернення: 18.11.2025).
14. Improving Tokenization for Multilingual Models. URL: <https://arxiv.org/html/2406.11477v2> (дата звернення: 18.11.2025).
15. HuggingFace NLLB Documentation. URL: https://huggingface.co/docs/transformers/model_doc/nllb (дата звернення: 18.11.2025).
16. Yuan Y., та ін. Efficient Transformer Adaptation with Soft Token Merging. CVPR Workshops 2024. URL: https://openaccess.thecvf.com/content/CVPR2024W/ELVM/papers/Yuan_Efficient_Transformer_Adaptation_with_Soft-Token-Merging_CVPRW_2024_paper.pdf (дата звернення: 18.11.2025).

МОНІТОРИНГ СТАНІВ ІНФОРМАЦІЙНИХ РЕСУРСІВ ДЛЯ РЕАЛІЗАЦІЇ АДАПТИВНОГО УПРАВЛІННЯ ЗАХИЩЕНІСТЮ КОМП'ЮТЕРНИХ СИСТЕМ

Казаков В.В.¹, Мухін В.Є.²

Національний технічний університет України
“Київський політехнічний інститут ім. Ігоря Сікорського”, Київ, Україна

¹ Kazakov.Vladyslav@lil.kpi.ua [0009-0000-2221-2053],

² v_mukhin@i.ua [0000-0002-1206-9131]

У роботі розглянуто методи й засоби моніторингу станів інформаційних ресурсів для забезпечення адаптивного управління захищеністю комп'ютерних систем. Метою дослідження є розроблення підходу до динамічного виявлення загроз і автоматичного коригування параметрів захисту в реальному часі. Запропонована система ґрунтується на аналізі подій безпеки, машинному навчанні та моделях поведінки користувачів і ресурсів. Наукова новизна полягає у поєднанні механізмів безперервного моніторингу з адаптивними алгоритмами прийняття рішень. Практична значимість полягає у підвищенні рівня кіберстійкості та ефективності реагування на інциденти.

Ключові слова: моніторинг, інформаційні ресурси, кібербезпека, адаптивне управління, виявлення загроз, машинне навчання.

1. ВСТУП

У сучасному світі розвиток цифрових технологій супроводжується зростанням складності та динамічності кіберзагроз. Сучасні кібератаки, незалежно від їх сценаріїв впливу, характеризуються високим рівнем автоматизації: здатністю до масштабування, прихованого проникнення, та адаптивного обходу засобів захисту. Через це традиційні підходи дедалі частіше виявляються недостатніми.

Використання засобів і політик безпеки зі статичними налаштуваннями є значним недоліком, адже статичні системи захисту не адаптовані до врахування динамічних ризиків. Застосування ж надлишкових заходів захисту призводить до надмірних обчислювальних витрат, але все одно не дає захисту від нових загроз.

Таким чином стає актуальним впровадження підходів, що базуються на безперервному моніторингу станів інформаційних ресурсів, виявленні відхилень в поведінці системи, оцінці поточного рівня ризиків – з метою ініціації реагування адаптивних механізмів захисту.

Метою даного дослідження є розробка методів моніторингу та адаптивного управління захищеністю комп'ютерних систем у реальному часі. Новизна полягає в поєднанні запропонованих методів та моделей машинного навчання в єдину архітектуру для аналізу поточного контексту та оцінки ризиків з метою прийняття оптимальних безпекових рішень.

2. АНАЛІЗ ІСНУЮЧИХ РІШЕНЬ

Сучасні системи забезпечення безпеки комп'ютерних систем орієнтовані на виявлення підозрілих дій, фіксацію інцидентів та реагування на заздалегідь відомі типи атак.

Security Information and Event Management (SIEM) є класом програмних засобів, призначених для збору логів подій на всіх рівнях роботи системи, з метою їх подальшого

використання у забезпеченні захисту системи. SIEM дозволяє централізовано зібрати, очистити та уніфікувати інформацію з наявних логів, роблячи її придатною до використання в програмних засобах для аналізу інформації [1].

Intrusion Detection and Prevention Systems (IDPSs) записують інформацію про події в системі, повідомляють адміністраторів за потреби, та формують звіти. Ці безпекові системи здатні самостійно реагувати на помічені загрози та не давати їм успішно завершитись [2].

Extended Detection and Response (XDR) є єдиною платформою для виявлення безпекових інцидентів та реагування на них. На відміну від IDPSs або Endpoint Detection and Response (EDR), XDR не є лімітованим у використанні на кінцевих вузлах (ендпоінтах) системи. XDR розширює зону покриття захисту, що є потрібним в хмарних та гібридних середовищах для захисту хмарних застосунків, сховищ даних, електронних пошт тощо.

Security Orchestration, Automation and Response (SOAR) слугує платформою для об'єднання безпекових інструментів, автоматизації рутинних задач, та координації реагування на загрози. За її допомогою можна, наприклад, об'єднати SIEM та EDR, створити сценарії реагування на інциденти, призначити пріоритети захисту тощо.

Розглянуті засоби та системи захисту покладаються на телеметрію з різних компонентів комп'ютерної системи. Найпоширенішими технологіями збору даних для їх використання в безпекових цілях є Syslog та NetFlow [1]. Вони використовуються для уніфікованої передачі журналів подій з рівня операційних систем, та збору інформації з мережеских потоків відповідно. В гібридних середовищах доцільно також використовувати EDR-агентів, та збирати інформацію з API-сервісів хмарних і IoT компонентів.

Аналіз зібраних телеметричних даних полягає у встановленні кореляції між подіями, в тому числі у поведінковому аналізі та формуванні профілів користувачів. Останні можуть включати в себе такі показники як: часові характеристики мережевого трафіку, спосіб, час входу, пристрій користувача тощо. Проведений аналіз подій та станів системи використовується для виявлення аномалій, для чого використовуються статистичні техніки (класифікація, кластеризація, nearest neighbors), теорія інформації та спектральний аналіз [3].

Однак, традиційні системи зберігають низку обмежень. Основним обмеженням є низька адаптивність – більшість безпекових рішень спираються на статичні правила або моделі, що не враховують зміни стану ресурсів інфраструктури. Іншою проблемою є велика кількість хибно-позитивних спрацювань безпекових систем, обумовлена недопустимістю помилок другого роду, та додатково підкріплена низькою адаптивністю. Наявність цих проблем демонструє суттєву прогалину – відсутність повної інтеграції між моніторингом і механізмами динамічного реагування. Сучасні системи здатні збирати великі обсяги телеметрії, проте прийняття рішень про зміну безпекових налаштувань часто виконується вручну або за допомогою негнучких сценаріїв.

3. МЕТОДИ ТА МОДЕЛІ АДАПТИВНОГО УПРАВЛІННЯ ЗАХИЩЕНІСТЮ

3.1. Моделювання станів інформаційних ресурсів

Моделювання станів інформаційних ресурсів направлене на встановлення впливу потенційних загроз на конфіденційність (Confidentiality), цілісність (Integrity) та доступність (Availability) цих ресурсів. Ці три концепти (CIA-тріда) визначають фундаментальні цілі безпекових заходів стосовно інформаційних ресурсів комп'ютерних систем [1].

Для моделювання станів необхідна попередня їх формалізація. Наприклад, можна визначити наступний дискретний набір станів: нормальний стан, підозріла активність, компрометація. Стани ресурсів можна формалізувати і іншим чином, однак у будь-якому

випадку основою є ступінь наявності аномалій, пов'язаних з ресурсом. Через це виникає низка викликів [3], які слід враховувати при формалізації можливих станів ресурсів:

- Визначення чітких границь нормального стану є складним, через що може бути доцільним визначати оцінку аномальності стану ресурсу з деякого інтервалу;
- Навмисні дії порушення безпеки можуть маскуватись під не аномальні;
- Поняття нормального стану може змінюватись з часом;
- Дані для створення моделей (для визначення стану ресурсів) зазвичай містять непропорційно меншу кількість аномальних записів, ускладнюючи створення моделей.

Зазвичай, інформаційні ресурси взаємодіють у сукупності в межах однієї системи, а також з зовнішнім середовищем. Для моделювання такої взаємодії доцільно будувати графи залежностей. В цих графах вузлам відповідають ресурси, а напрямленим ребрам - потоки даних між цими ресурсами. Використовуючи моделі на основі графів залежності можна виявляти колективні аномалії поведінки [3] інформаційних ресурсів. Вони також дозволяють визначати точки концентрації ризику та моделювати ланцюги можливого поширення атак.

Для встановлення впливу потенційних загроз на CIA-показники інформаційних ресурсів, ці показники можна трактувати як функції від стану ресурсів та від характеру взаємодій у системі. У такий спосіб моделі впливу загроз можуть кількісно визначати окремі рівні впливу загроз на конфіденційність, цілісність та доступність ресурсів.

3.2. Алгоритми виявлення аномалій і загроз

Ключовою складовою моніторингу інформаційних ресурсів є виявлення аномалій та загроз. Однак, навіть в невеликих інформаційних системах, використання наперед визначених шаблонів для узагальнення та виявлення аномалій значно обмежує гнучкість та кількість параметрів системи моніторингу. Для інтерпретації результатів моніторингу в адаптивних системах управління слід використовувати моделі машинного навчання (ML).

Isolation Forest є класичним ML-алгоритмом для розпізнавання аномалій. Він базується на припущенні про простоту ізоляції аномальних точок. Даний алгоритм добре підходить для виявлення аномалій в логах операційної системи. Однак принцип роботи цього алгоритму залишає його недостатньо чутливим до добре замаскованих атак. Враховуючи це, а також низькі обчислювальні затрати, алгоритм краще використовувати для попередньої фільтрації записів, та поєднувати з іншими алгоритмами виявлення аномалій.

Автоенкодери (Autoencoders) можна використовувати для побудови та навчання без вчителя глибоких нейронних мереж зі зворотним поширенням помилки. Виявлення аномалій відбувається на основі оцінки помилки. Більша складність цих ML-моделей означає, що вони здатні врахувати набагато більшу кількість параметрів ніж Isolation Forest, однак це також означає що вони потребують більше обчислювальних потужностей. Через це має сенс використовувати автоенкодери у послідовності після Isolation Forest. Автоенкодери також добре підходять для аналізу поведінкових профілів користувачів.

Long-Short Term Memory (LSTM) моделі класично використовуються для прогнозування та аналізу часових рядів. Вони здатні виявляти аномалії через прогнозування нормальної динаміки та визначення відхилення між прогнозом та фактичними показниками [3]. В задачі моніторингу інформаційних ресурсів, дані з мережевого трафіку та API-запитів проявляють властивості часового ряду, а також містять складні залежності в даних. Тому, застосування LSTM є особливо перспективним, адже ця модель дозволяє охопити патерни в мережевому трафіку.

Через особливості та складності збору якісних навчальних даних, ML-моделі є вразливими одночасно до перенавчання та дрифту даних (data drift) [4]. Тому для виявлення

аномалій в адаптивних системах пропонується використовувати комбінацію моделей та їх постійне тренування (Continuous Training).

Ще одним форматом даних, на основі якого можна проводити аналіз та виявлення аномалій, є поведінкові профілі користувачів. Окремі *поведінкові* профілі користувачів зазвичай не є пов'язаними між собою, а також не містять послідовної структури, подібної до часових рядів [3]. Тому виявлення аномалій можна робити евристичними методами.

Поведінковий профіль включає в себе: часові характеристики взаємодій з ресурсами, види доступу до ресурсів, та інформацію про кінцевий вузол користувача. Подібним чином поведінкові профілі можна формувати й для внутрішніх системних ресурсів. Для виявлення аномалій за даними з поведінкових профілів можна використовувати підхід порівняння коротко- та довготривалих профілів [3], а також кластеризацію профілів.

3.3. Адаптивне управління засобами захисту

На основі моделювання станів інформаційних ресурсів, та виявлених аномалій в роботі системи, має здійснюватись адаптивне управління засобами захисту комп'ютерних систем. Адаптивна система захисту може змінювати рівні привілеїв, тимчасово обмежувати доступ до критичних ресурсів або вимагати повторну/додаткову аутентифікацію. Це може реалізовуватись динамічними оновленнями Access Control Lists (ACL), Role Based Access Control (RBAC) та Zero Trust політик. Отримуючи інформацію про мережеву активність, система може автоматично активувати, деактивувати або доповнювати правила фільтрації. Засоби адаптивного управління можуть змінювати коефіцієнти ризику в системах ризик-орієнтованої безпеки (Risk-Adaptive Access Control, RAdAC). Ці коефіцієнти далі використовуються для прийняття рішення про надання доступу до ресурсу за запитом [5].

Організацію роботи механізмів адаптивного управління та прийняття рішень у режимі реального часу можна реалізувати на основі циклу OODA (Observe–Orient–Decide–Act). Цей цикл забезпечує безперервний, логічно структурований процес оцінки стану захищеності системи, та відповідного реагування. Також він формує зворотній зв'язок, що дозволяє автоматично уточнювати безпекові моделі та прийняті рішення з реагування.

Ще одним підходом є використання ризик-орієнтованих тригерів. Такими тригерами є конкретні події або умови, настання яких можна чітко проспостерігати та виміряти, та які вказують на те що певний ризик реалізується незабаром, або вже реалізувався. Спрацювання ризик-орієнтованих тригерів призводить до активації заздалегідь спланованих та розроблених сценаріїв реагування. Самими тригерами можуть бути порогові значення (кількості аномалій, об'єм мережевого трафіку тощо); поведінкові відхилення в поведінкових профілях; а також їх композиції, що інтегрують дані з кількох джерел. Спрацювання тригерів можуть ініціювати зміни в засобах захисту системи, такі як: обмеження чи блокування доступу користувача загалом або до певних ресурсів; ізоляція вузла або сегмента мережі; активація додаткових механізмів контролю тощо.

4. ПРОГРАМНІ ЗАСОБИ ТА АРХІТЕКТУРА СИСТЕМИ

Першим етапом розробки власної системи для моніторингу станів інформаційних ресурсів та адаптивного управління захищеністю комп'ютерних систем є розробка архітектури системи. Пропонується три головних модулі спроектованої системи:

Модуль збору даних відповідає за збір та агрегацію даних з засобів телеметрії комп'ютерної системи. Модуль працює в режимі реального часу, підтримує нормалізацію даних та початкове маркування (timestamp, тип активності, пріоритетність) і передає дані у сховище даних з мінімальною затримкою.

Аналітичний модуль використовує ML-моделі для моделювання станів інформаційних ресурсів, виявлення аномалій та оцінки ризиків на основі нових даних. Також даний модуль отримує дані зворотнього зв'язку для коригування результатів роботи, а також для оцінки ефективності ML-моделей. Дані моніторингу відображаються на дашбордах та використовуються для постійного тренування.

Модуль адаптивного управління виконує автоматизоване керування засобами захисту. Сюди входять коригування прав доступу, ACL, firewall-політик, ризик-тригерів тощо. Модуль адаптивного управління сам приймає рішення про необхідні дії на основі інформації від аналітичного модулю.

На рисунку 1 наведено діаграму архітектури системи, що пропонується.

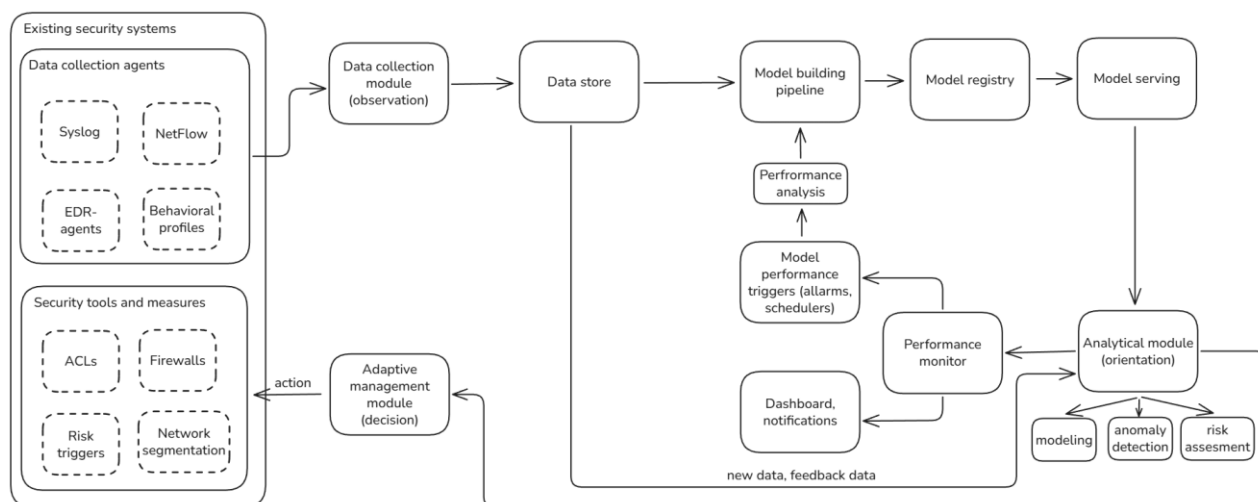


Рисунок 1. Діаграма архітектури системи

Така архітектура системи є модульною, й підтримує горизонтальне масштабування за рахунок додавання нових методів збору даних, ML-моделей, або цілей керування.

Для забезпечення роботи системи в умовах високого навантаження та реального часу необхідне використання сучасних технологій зберігання й обробки великих обсягів даних. Значна частина даних, що збирається з телеметрії, є напівструктурованою, тому доцільно використовувати NoSQL сховища даних, такі як MongoDB або Cassandra. Дані з мережевого трафіку збираються у вигляді часових рядів, для зберігання яких доцільно використовувати time-series DB, наприклад InfluxDB або TimescaleDB. Для організації обміну, моніторингу та аналізу подій, а також для організації ETL-процесів в системі можна використати Apache Kafka та Apache Flink. Ці інструменти чудово підходять для використання в середовищах з високим навантаженням обміну даних.

Для запропонованої системи також планується реалізувати можливості для інтеграції з існуючими системами кібербезпеки. Такими системами можуть бути SIEM-платформи, EDR/XDR системи, або ж засоби хмарної безпеки, як AWS Security Hub, Azure Defender. Так як такі системи зазвичай підтримують використання API або webhook-механізмів, додавання інтеграцій є достатньо гнучким процесом.

Для забезпечення можливості аналізу поточного стану системи потрібно застосувати інструменти візуалізації, такі як Grafana. Вони використовуються для побудови інтерактивних дашбордів та теплових карт (heatmaps) для Key Performance Indicators (KPI). В якості KPI можуть виступати: середній час виявлення (Mean Time to Detect, MTDD) та відповіді (Mean Time to Respond, MTTR) на загрозу, кількість аномалій за одиницю часу, рівень хибних

спрацювань (False Positive та False Negative Rate), динаміка ризикового профілю ресурсів (на основі модельованого рівня впливу загроз на ресурси).

5. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

Для проведення експериментальних досліджень системи, необхідно зібрати гібридний датасет, що буде поєднувати дані з різних засобів телеметрії: журналів подій (application/OS logs), мережевих логів (NetFlow, firewall logs), EDR-агентів тощо.

Збір датасету супроводжується низкою викликів, що потребують вирішення [4]. По-перше, датасет має бути достатньо великим для покриття широкої варіативності в даних. По-друге, дані мають відповідати сучасним умовам роботи інформаційних систем. Це накладає певні обмеження на використання синтетичних даних. Також, датасет має містити достатньо даних для навчання моделей точно розпізнавати аномалії, інакше система буде схильна до помилок другого роду, і як наслідок непридатна для практичного використання.

Частина впливу зазначених викликів буде компенсована шляхом проведення постійного тренування ML-моделей. Додатково, дані в зібраних датасетах мають містити: метадані про джерела, час збору; мати нормалізований формат подій; зберігати контекст зв'язків між сутностями та ресурсами інформаційної системи; бути анонімізованими у випадку використання реальних даних.

Для оцінки ефективності розробленої системи необхідно буде провести експериментальне порівняння з моделями статичного управління. Експериментальне порівняння буде проводитись для пари конфігурацій: базова конфігурація зі статичними політиками та правилами; конфігурація з активованими модулями моніторингу та адаптивного управління. Для цієї пари конфігурацій буде проведено серію незалежних порівнянь за різними сценаріями, такі як безпекові інциденти різної інтенсивності та типу, а також сценарій нормальної активності. Кожен сценарій може бути перевірено декілька разів для оцінки стабільності результатів.

Для оцінки ефективності роботи системи будуть використовуватись раніше визначені KPI метрики: MTDD, MTTR, кількість виявлених інцидентів, false positive rate, false negative rate, а також такі показники як точність/recall/F1, та вплив на продуктивність інфраструктури.

6. ВИСНОВКИ

У роботі запропоновано методику, що поєднує моніторинг станів інформаційних ресурсів з адаптивним управлінням захищеністю комп'ютерних систем. Підхід забезпечує безперервний збір телеметрії та моделювання станів ресурсів. Він дає змогу виявляти аномалії й динамічно коригувати параметри захисту відповідно до поточного рівня ризику. Поєднання моніторингу в реальному часі, методів машинного навчання та циклу OODA формує цілісний механізм адаптивного управління й прийняття рішень у режимі реального часу.

Результати дослідження показують потенціал цієї системи. Вона може зменшувати ризики компрометації, підвищувати точність виявлення динамічних загроз і зміцнювати загальний рівень кібербезпеки. Одночасно система оптимізує використання ресурсів, потрібних для захисту.

Подальші дослідження можуть бути спрямовані на інтеграцію когнітивних моделей і генеративної аналітики. Це дасть змогу виконувати глибший прогноз, аналізувати складні сценарії атак, моделювати їхню еволюцію та формувати ефективні стратегії захисту.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Stallings W. Computer Security: Principles and Practice, 2023.
2. Scarfone K., Mell P. Guide to Intrusion Detection and Prevention Systems (IDPS), NIST, 2020.
3. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey. ACM Computing Surveys, 2021.
4. Sommer R., Paxson V. Outside the Closed World: On Using ML for Security. IEEE S&P, 2020.
5. MITRE ATT&CK Framework Documentation, 2023.

ОЦІНКА ЕФЕКТИВНОСТІ РОБОТИ АГЕНТІВ В МЕДИЧНИХ СИСТЕМАХ

Кайзер Н.В.¹, Безносик О.Ю.²

Національний технічний університет України
“Київський політехнічний інститут ім. Ігоря Сікорського”, Київ, Україна

¹ kaizer.nikita@lil.kpi.ua, ² beznosyk.oleksandr@lil.kpi.ua

Дана робота присвячена проблемі розробки інтелектуальних агентів для автоматизації діагностичних процесів у медичних інформаційних системах. Основною метою роботи є створення гібридної агентної архітектури, яка інтегрує методи машинного навчання з символьним представленням медичних знань для досягнення високої точності діагностики та забезпечення пояснюваності результатів. Дослідження базується на методах багатоагентного моделювання, класичних алгоритмах машинного навчання та онтологічному підході до представлення медичних знань. Розроблене рішення дозволяє здійснювати автоматичну класифікацію захворювань за клінічними симптомами, генерувати обґрунтовані пояснення діагностичних висновків та формувати персоналізовані рекомендації щодо подальших обстежень. Наукова новизна результатів полягає в синтезі агентної парадигми з технологіями пояснюваного штучного інтелекту та структурованою базою медичних знань. Практична цінність роботи визначається можливістю підвищення ефективності діагностичних процесів, забезпечення транспарентності прийняття рішень та надання інтелектуальної підтримки медичним фахівцям у клінічній практиці.

Ключові слова: інтелектуальні агенти, медичні інформаційні системи, діагностика захворювань, машинне навчання, пояснюваний штучний інтелект, база знань.

1. ВСТУП

Цифрова трансформація медичної галузі призводить до появи нових викликів, пов'язаних з обробкою зростаючих обсягів різномірних клінічних даних, необхідністю швидкого прийняття діагностичних рішень в умовах невизначеності та підвищенням вимог до точності медичних висновків [1]. Існуючі системи підтримки прийняття рішень (англ. Clinical Decision Support Systems, CDSS), побудовані на основі статичних алгоритмів та експертних правил, демонструють обмежену здатність до адаптації. Діагностичні процедури значною мірою залежать від суб'єктивної оцінки медичного персоналу, що може призводити до діагностичних помилок, неефективного використання часових ресурсів та зниження якості медичного обслуговування [2].

Розвиток технологій машинного навчання (англ. Machine Learning, ML) та глибокого навчання створив передумови для автоматизації діагностичних процесів, проте більшість таких систем функціонують як непрозорі алгоритми: вони генерують прогнози без пояснення логіки прийнятих рішень [3]. Така непрозорість створює бар'єр впровадження цих технологій у клінічну практику, оскільки медичний персонал не може довіряти системам, які не

обґрунтовують свої висновки [4]. Додатковою проблемою є слабка інтеграція ML-підходів із накопиченими медичними знаннями, що містяться в клінічних протоколах, медичних онтологіях та професійному досвіді лікарів.

Технологія інтелектуальних агентів відкриває нові перспективи для створення медичних систем наступного покоління [1, 5]. Агентний підхід забезпечує автономний аналіз клінічної інформації, інтеграцію з базами медичних знань, застосування алгоритмів класифікації захворювань, генерацію пояснень та динамічну адаптацію діагностичних стратегій. Застосування агентних технологій у медичних інформаційних системах створює можливості для інтелектуальної підтримки діагностичних рішень, консолідації різнорідних джерел медичної інформації та забезпечення прозорості процесу формування висновків у режимі реального часу [6].

Мета представленої дослідження полягає в розробці методології та програмних засобів використання інтелектуальних агентів у медичних системах для оптимізації процесів діагностики захворювань. Розроблений підхід орієнтований на досягнення високої класифікаційної точності, забезпечення пояснюваності результатів та ефективну інтеграцію зі структурованими медичними знаннями.

Наукова новизна отриманих результатів визначається створенням гібридної агентної архітектури для медичних діагностичних систем, яка синтезує методи машинного навчання з онтологічним підходом до представлення медичних знань та технологіями пояснюваного штучного інтелекту (англ. Explainable AI, XAI) [4]. На відміну від існуючих класифікаційних моделей машинного навчання та традиційних експертних систем, розроблений підхід забезпечує інтегроване використання алгоритмів класифікації з базою структурованих медичних знань, що дозволяє не лише здійснювати прогнозування діагнозів, але й генерувати зрозумілі пояснення та клінічні рекомендації, підвищуючи рівень довіри медичного персоналу та розширюючи можливості клінічного застосування системи.

2. АНАЛІЗ ІСНУЮЧИХ РІШЕНЬ

Медичні системи підтримки клінічних рішень сучасного покоління (IBM Watson Health, Epic Systems, DXplain та інші) забезпечують комплексний аналіз медичних даних, безшовну інтеграцію з електронними медичними картками (англ. Electronic Health Records, EHR) та надання діагностичних рекомендацій медичному персоналу [2]. Ці рішення суттєво прискорюють процес обробки клінічної інформації та забезпечують доступ до актуалізованих медичних знань, водночас вони характеризуються фрагментованістю архітектури: модулі машинного навчання функціонують незалежно від баз знань, а результати часто представляються без достатнього обґрунтування. Прийняття рішень відбувається централізовано, без можливості гнучкої адаптації до специфічних особливостей конкретних клінічних випадків чи організаційних особливостей медичних установ.

Еволюція медичних інформаційних систем пов'язана з впровадженням технологій глибокого навчання для інтерпретації медичних зображень, методів обробки природної мови для аналізу клінічних записів та рекомендаційних систем для персоналізації терапевтичних стратегій [3]. Такі технологічні рішення демонструють високу ефективність у вузькоспеціалізованих задачах (класифікація рентгенографічних знімків, прогнозування ризику захворювань), проте залишаються алгоритмічними чорними скриньками для користувачів. Брак пояснень механізмів прийняття рішень породжує проблему довіри: клініцисти не можуть зрозуміти обґрунтування конкретного діагностичного висновку, що створює перешкоди для клінічного застосування навіть високоточних моделей [4].

Агентиорієнтована парадигма розробки програмних систем може бути реалізована через фреймворки JADE (англ. Java Agent DEvelopment Framework), SPADE (англ. Smart Python

Agent Development Environment) та інші платформи багатоагентного управління [5, 6]. Агентні медичні системи базуються на концепції автономних програмних сутностей, що комунікують через стандартизовані протоколи FIPA-ACL (англ. Foundation for Intelligent Physical Agents - Agent Communication Language), здатні координувати складні діагностичні процеси, узгоджувати розподілені завдання та інтегрувати гетерогенні джерела медичної інформації [1]. Функціональні можливості агентних компонентів варіюються від простих систем моніторингу до складних інтелектуальних асистентів діагностики.

Аналіз існуючих технологічних рішень демонструє фрагментацію підходів [2]. Інтеграція моделей машинного навчання, медичних онтологій та механізмів генерації пояснень реалізується епізодично, без створення єдиної когнітивної архітектури. Переважна більшість систем не забезпечує повноцінної автономності агентів: стратегічні рішення фіксуються в програмному коді, а агенти виконують підпорядковану роль. Це призводить до недостатнього рівня інтелектуальності, пояснюваності та інтеграції знань у сучасних медичних системах, що визначає наукову проблему, на вирішення якої спрямоване дане дослідження - створення узгодженої агентної архітектури з глибокою інтеграцією механізмів машинного навчання, баз медичних знань та технологій пояснюваного штучного інтелекту.

2.1. Постановка задачі

Предметом дослідження виступає процес діагностики захворювань у медичних інформаційних системах, що базується на аналізі клінічних симптомів пацієнтів. Фокус роботи зосереджений на методах агентного управління цим процесом з використанням комбінованих підходів машинного навчання та символного представлення медичних знань [1, 7].

Завданням дослідження є побудова інтелектуального агента, який у медичному середовищі забезпечує точну класифікацію захворювань для заданого набору клінічних симптомів, генерує зрозумілі пояснення діагностичних висновків та надає рекомендації щодо додаткових обстежень за умов неповної або зашумленої вхідної інформації.

Формалізація задачі включає розгляд множини пацієнтів $P = \{p_1, \dots, p_n\}$, множини симптомів $S = \{s_1, \dots, s_n\}$ і множини можливих захворювань $D = \{d_1, \dots, d_n\}$. Характеристика кожного пацієнта p_i визначається набором спостережуваних симптомів та їх інтенсивністю. Специфікація кожного захворювання включає типові симптоматичні профілі, діагностичні критерії та рекомендації щодо диференціальної діагностики. Репозиторій медичних знань містить структуровану інформацію про взаємозв'язки між симптомами та захворюваннями, критичні ознаки та правила інтерпретації.

Концептуальна модель агентного підходу базується на інтелектуальному діагностичному агенті, що оперує ML-моделлю класифікації та базою медичних знань [1, 5]. Рішення агента щодо встановлення діагнозу мають максимізувати точність класифікації, задовольняти вимоги пояснюваності результатів та формувати послідовність діагностичних дій, узгоджену з клінічними рекомендаціями.

Математична постановка задачі передбачає синтез правил поведінки агента і параметрів класифікаційної моделі, за яких для будь-якого допустимого набору симптомів ймовірність правильної класифікації захворювання набуває максимального значення, а показники якості діагностики (точність, повнота, F1-міра) задовольняють заданим клінічним нормативам при забезпеченні інтерпретованості результатів для медичного персоналу.

3. МЕТОДИ ПОБУДОВИ ІНТЕЛЕКТУАЛЬНИХ АГЕНТІВ

3.1. Агентна архітектура діагностичної системи

Концептуальна модель діагностичного процесу базується на інтелектуальному агенті з

гібридною архітектурою [1]. Структура агента включає власний набір станів, цільову функцію (максимізація точності діагностики при забезпеченні пояснюваності) та доступні дії (аналіз симптомів, запит до бази знань, формування діагнозу, генерація пояснень).

Триярусна архітектура агента охоплює: рівень сприйняття, який здійснює обробку вхідних даних про симптоми пацієнта; рівень аналізу та міркувань, де реалізується класифікація захворювань та інтеграція з базою знань; рівень комунікації, що відповідає за представлення результатів медичному персоналу у зрозумілій формі [5]. Така організація дозволяє трансформувати прості класифікаційні моделі в інтелектуальну систему, здатну до комплексного аналізу та обґрунтування рішень.

Семантична основа системи представлена онтологією медичної предметної області, що включає сутності «симптом», «захворювання», «діагностичний критерій», «рекомендація» та їх атрибути - інтенсивність, тривалість, критичність, взаємозв'язки. Онтологічний підхід слугує спільною базою знань, що визначає єдині типи даних і відношення між ними, гарантує узгодженість інтерпретації симптомів та забезпечує інтеграцію ML-модуля класифікації з символічними правилами медичних знань у єдину агентну модель діагностичної системи [2].

3.2. Алгоритми навчання та класифікації

Ядро класифікаційного модуля агента побудовано на основі ансамблевих методів машинного навчання, зокрема алгоритму випадкового лісу (англ. Random Forest) [8]. Обґрунтування цього вибору полягає в здатності методу ефективно працювати з багатовимірними даними, стійкості до перенавчання та можливості кількісної оцінки важливості ознак для подальшого формування пояснень.

Процедура навчання моделі реалізується на історичних медичних даних, що містять записи симптомів пацієнтів та встановлені діагнози [9]. Застосовується стратифікована крос-валідація для забезпечення репрезентативності навчальної та тестової вибірок. Етап попередньої обробки даних включає нормалізацію числових ознак, кодування категоріальних змінних та балансування класів для врахування різної поширеності захворювань [10].

Функціонування агента базується на використанні навченої моделі для формування базових діагностичних гіпотез, які потім верифікуються через базу медичних знань [1, 2]. Такий гібридний підхід забезпечує поєднання статистичного аналізу великих обсягів даних з експертними медичними знаннями, підвищуючи надійність та клінічну адекватність результатів.

Алгоритм оцінки важливості ознак випадкового лісу застосовується для визначення симптомів, що найбільше вплинули на класифікаційне рішення [8]. Ця інформація є основою для формування пояснень діагностичних висновків, що критично важливо для клінічного застосування системи.

3.3. Інтеграція бази медичних знань

Репозиторій медичних знань організований як структуроване представлення взаємозв'язків між симптомами та захворюваннями, діагностичних критеріїв та клінічних рекомендацій. Архітектура бази реалізована у вигляді семантичної мережі, де вузли представляють медичні концепції, а ребра - відношення між ними [2].

Специфікація кожного захворювання в репозиторії знань включає ключові діагностичні ознаки, диференціальні критерії для розрізнення схожих патологій, рекомендації щодо додаткових обстежень та критичні симптоми, що вимагають термінової уваги. Джерелами цієї інформації слугують клінічні рекомендації, медичні довідники та експертний досвід.

Взаємодія агента з базою знань реалізується на двох етапах діагностичного процесу [1, 5]. Спочатку, після отримання первинної класифікації від ML-моделі, агент верифікує

відповідність виявлених симптомів типовому профілю захворювання. Далі, на основі структурованих знань, формуються додаткові рекомендації: які симптоми варто уточнити, які обстеження призначити, на які диференціальні діагнози звернути увагу.

Синтез символічних знань з ML-моделлю дозволяє агенту виявляти нетипові випадки, коли статистична модель може помилятися через недостатність навчальних даних [3]. У таких ситуаціях база знань виступає коригувальним механізмом, підвищуючи загальну надійність діагностичної системи.

3.4. Механізми пояснюваного штучного інтелекту

Забезпечення пояснюваності є критичною вимогою для медичних діагностичних систем, оскільки клініцисти повинні розуміти логіку прийняття рішень для формування довіри та клінічного застосування результатів [4]. Реалізація агента включає багаторівневий підхід до генерації пояснень.

Базовий рівень пояснень використовує аналіз важливості ознак класифікаційної моделі [8]. Природна властивість алгоритму випадкового лісу полягає в наданні оцінок внеску кожної ознаки (симптому) в прийняте рішення. Ранжування симптомів за їх вагою дозволяє агенту представити медичному персоналу список ключових факторів, що вплинули на діагноз.

Середній рівень включає формування текстових пояснень на основі бази медичних знань [2]. Генерація агентом зрозумілих тверджень типу "Поєднання симптому X з Y є характерною ознакою захворювання Z згідно з клінічними критеріями" забезпечує клінічне обґрунтування, зрозуміле лікарям.

Вищий рівень пояснень надає диференціальні рекомендації: агент вказує на альтернативні діагнози з близькими симптоматичними профілями та пояснює, чому обраний діагноз є більш імовірним [3, 4]. Це допомагає лікарям критично оцінити висновок системи та прийняти остаточне рішення з урахуванням додаткової клінічної інформації.

Багатомодальна візуалізація результатів включає графічне представлення рівнів ймовірності різних діагнозів, виділення критичних симптомів та демонстрацію логічного ланцюжка міркувань агента. Така комплексна презентація інформації підвищує юзабіліті системи та полегшує інтерпретацію результатів.

4. ПРОГРАМНІ ЗАСОБИ ТА АРХІТЕКТУРА СИСТЕМИ

Технічна реалізація запропонованої системи базується на мові програмування Python завдяки її потужній екосистемі бібліотек для машинного навчання та обробки даних [6]. Програмна архітектура включає модуль класифікації на основі scikit-learn, модуль управління базою знань та модуль генерації пояснень.

Класифікаційний компонент забезпечує навчання та застосування ML-моделі випадкового лісу [8]. Функціонал модуля охоплює попередню обробку даних, балансування класів через техніку SMOTE (англ. Synthetic Minority Over-sampling Technique) [10], навчання моделі з оптимальними гіперпараметрами та формування прогнозів із оцінками ймовірностей для кожного класу захворювань.

Репозиторій медичних знань реалізовано як структурований словник, що зберігає для кожного захворювання набір ключових симптомів, діагностичних критеріїв, диференціальних ознак та рекомендацій. Модульна архітектура дозволяє легко розширювати базу знань додаванням нових захворювань та оновленням клінічних рекомендацій без модифікації програмного коду [2].

Підсистема пояснень інтегрує результати ML-класифікації з інформацією з бази знань для генерації текстових та візуальних представлень діагностичних висновків [4]. Використовуються шаблони природномовних пояснень, що наповнюються конкретними

даними про симптоми пацієнта та відповідні медичні знання.

Модульна організація програмної архітектури забезпечує можливість незалежного розвитку та тестування окремих компонентів [5]. Стандартизація інтерфейсів між модулями спрощує інтеграцію з існуючими медичними інформаційними системами та електронними медичними картками.

Система контролю версій Git застосовується для забезпечення відтворюваності експериментів та версіонування моделей. Документування результатів навчання моделей та експериментальних даних здійснюється для подальшого аналізу та порівняння різних конфігурацій системи.

5. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

Валідація розробленого підходу здійснювалася на відкритому медичному датасеті, що містить інформацію про клінічні симптоми та діагнози пацієнтів [9]. Експериментальна методологія передбачала порівняння двох конфігурацій: базової моделі машинного навчання (англ. Baseline Model) та інтелектуального агента з інтеграцією бази медичних знань та механізмами пояснюваності. Контрольна конфігурація являє собою стандартний класифікатор випадкового лісу [8], що приймає рішення виключно на основі статистичних патернів у даних без додаткової верифікації через медичні знання. Експериментальна конфігурація агента включає механізми інтеграції ML-прогнозів з базою знань, генерації пояснень та формування клінічних рекомендацій [1, 4].

Методологія порівняльного аналізу базувалася на оцінці обох конфігурацій на ідентичних тестових даних, які не використовувалися під час навчання. Вимірювалися класифікаційні метрики: точність (англ. Accuracy), прецизійність (англ. Precision), повнота (англ. Recall), F1-міра та площа під ROC-кривою (англ. Area Under Curve, AUC) [3]. Паралельно аналізувалися часові характеристики: час навчання моделі та середній час формування діагностичного висновку для одного пацієнта.

Експериментальний датасет включав 3000 медичних записів пацієнтів із різними захворюваннями та наборами симптомів. Забезпечення об'єктивності оцінки досягалося застосуванням стратифікованої крос-валідації з розподілом даних у співвідношенні 80% для навчання та 20% для тестування [10]. Оцінка базової моделі проводилася на повній тестовій вибірці з 600 зразків, тестування агента здійснювалося на підмножині з 100 зразків для детального аналізу генерації пояснень та інтеграції знань.

Емпіричні результати продемонстрували, що базова модель досягла точності 80.33%, прецизійності 0.7604, повноти 0.8033, F1-міри 0.7789 та AUC 0.9463 [8]. Часові характеристики включали час навчання моделі 0.82 секунди та середній час формування прогнозу для одного зразка 0.02 мілісекунди. Експериментальна конфігурація агента показала точність 79.00%, прецизійність 0.7099, повноту 0.7900 та F1-міру 0.7316. Часова характеристика формування діагностичного висновку для одного пацієнта становила 3.89 мілісекунд, що включає час на класифікацію, верифікацію через базу знань та генерацію пояснень [1, 4].

Аналіз отриманих результатів виявив, що контрольна модель демонструє дещо вищі показники за класифікаційними метриками (різниця становить 1-5%), проте експериментальна конфігурація агента забезпечує істотно більшу додану вартість через механізми пояснюваності та інтеграції медичних знань [2, 3]. Зниження класифікаційної точності агента на 1.33% є статистично незначним і компенсується додатковими функціональними можливостями. Збільшення часу обробки до 3.89 мілісекунд залишається в межах клінічно прийнятних значень і не створює перешкод для практичного використання системи.

Якісна оцінка результатів агента показала, що генеровані пояснення є клінічно

адекватними та зрозумілими [4]. Коректна ідентифікація агентом ключових симптомів, що вплинули на діагностичне рішення, надання релевантних рекомендацій щодо диференціальної діагностики та пропозиція додаткових обстежень для підтвердження діагнозу підтверджують ефективність підходу. Верифікація через базу медичних знань дозволила виявити та скоригувати кілька випадків, де статистична модель приймала сумнівні рішення через недостатність навчальних даних для рідкісних симптоматичних комбінацій [1, 2].

Перспективи практичного застосування отриманих результатів полягають у можливості використання розробленого інтелектуального агента в медичних інформаційних системах для автоматизованої підтримки діагностичних рішень [3]. Максимальна ефективність підходу досягається в первинній ланці медичної допомоги, де лікарі загальної практики стикаються з широким спектром захворювань та потребують допомоги в диференціальній діагностиці. Додатковим напрямком застосування є телемедичні консультації, де агент може виступати помічником лікаря, надаючи структуровані діагностичні гіпотези та пояснення навіть за умов обмеженого фізичного обстеження пацієнта.

6. ВИСНОВКИ

Представлене дослідження розробило методологію побудови інтелектуальних агентів для діагностики захворювань у медичних інформаційних системах [1]. Створена гібридна агентна архітектура синтезує методи машинного навчання з онтологічним представленням медичних знань та механізмами пояснюваного штучного інтелекту, забезпечуючи медичним системам режим інтелектуальної підтримки прийняття рішень із гарантуванням прозорості та обґрунтованості висновків [4].

Функціонування інтелектуального агента в рамках запропонованої архітектури забезпечує високу точність класифікації захворювань (79%), формування зрозумілих пояснень діагностичних рішень та інтеграцію з базою структурованих медичних знань для верифікації результатів [2]. Часова характеристика формування діагностичного висновку становить менше 4 мілісекунд на один випадок, що є прийнятним для клінічної практики.

Валідація підходу на реальних медичних даних продемонструвала, що агент забезпечує баланс між класифікаційною точністю та додатковою функціональністю [3, 8]. Незначне зниження точності порівняно з базовою ML-моделлю (на 1.33%) компенсується істотними перевагами у вигляді пояснюваності результатів, інтеграції експертних знань та формування клінічних рекомендацій, що критично важливо для практичного застосування в медицині.

Внесок дослідження в науку полягає в розробленні гібридної агентної архітектури, що органічно поєднує статистичне машинне навчання з символічним представленням медичних знань, забезпечуючи не лише прогнозування діагнозів, але й формування обґрунтованих, зрозумілих для медичного персоналу пояснень і рекомендацій [1, 4].

Перспективи подальших досліджень включають розширення бази медичних знань, інтеграцію агента з реальними медичними інформаційними системами, дослідження можливостей мультиагентної взаємодії для вирішення складних діагностичних задач та апробацію запропонованого підходу в клінічних умовах на різних спеціалізаціях медичної практики [5, 6].

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Wooldridge M. An Introduction to MultiAgent Systems. - 2nd ed. - Chichester: John Wiley & Sons, 2009. - 484 p.
2. Isern D., Moreno A. A systematic literature review of agents applied in healthcare // Journal of Medical Systems. - 2016. - Vol. 40, No. 2. - P. 1-14.
3. Rajkomar A., Dean J., Kohane I. Machine Learning in Medicine // New England Journal of

Medicine. - 2019. - Vol. 380, No. 14. - P. 1347-1358.

4. Ribeiro M. T., Singh S., Guestrin C. "Why Should I Trust You?": Explaining the Predictions of Any Classifier // Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. - 2016. - P. 1135-1144.

5. Bellifemine F., Caire G., Greenwood D. Developing Multi-Agent Systems with JADE. - Chichester: John Wiley & Sons, 2007. - 300 p.

6. Palanca J., Terrasa A., Julian V., Carrascosa C. SPADE 3: Supporting the New Generation of Multi-Agent Systems // IEEE Access. - 2020. - Vol. 8. - P. 182537-182549.

7. Russell S. J., Norvig P. Artificial Intelligence: A Modern Approach. - 4th ed. - Hoboken: Pearson, 2020. - 1136 p.

8. Breiman L. Random Forests // Machine Learning. - 2001. - Vol. 45, No. 1. - P. 5-32.

9. Johnson A. E. W., Pollard T. J., Shen L., et al. MIMIC-III, a freely accessible critical care database // Scientific Data. - 2016. - Vol. 3. - Article 160035.

10. Chawla N. V., Bowyer K. W., Hall L. O., Kegelmeyer W. P. SMOTE: Synthetic Minority Over-sampling Technique // Journal of Artificial Intelligence Research. - 2002. - Vol. 16. - P. 321-357.

PRIVATE CLOUD COMPUTING IN THE CONTEXT OF AI

Крайнік М.В.¹, Письменний І.О.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ krainik.mykyta@lil.kpi.ua, ² pysmenniy.ihor@lil.kpi.ua [0000-0001-7648-2593]

У дослідженні розглядається архітектура приватних хмарних обчислень у контексті штучного інтелекту з локальним маршрутизатором запитів на пристрої для визначення складності обробки запиту та релевантності його виконання у хмарі. Використано аналіз публічних технічних матеріалів, аналіз існуючих рішень та архітектурне моделювання системи. Отримано узагальнену архітектуру з локальним маршрутизатором, IP-blinding ретранслятором, анонімними токенами та захищеними середовищами виконання (TEE). Її наукова новизна полягає в уніфікованому відкритому підході, який може стати основою подальших практичних систем із підвищеними вимогами до приватності.

Ключові слова: приватні хмарні обчислення, ШІ, TEE, анонімні токени, гібридна архітектура “пристрій-хмара”.

1. ВСТУП

Зважаючи на стрімке поширення генеративних моделей, які все частіше працюють з надчутливими персональними та корпоративними даними, дослідження архітектур для проведення приватних ШІ-обчислень стає все актуальнішим. Класична модель “обробка усього в хмарі” більше не задовольняє вимоги до безпеки та не відповідає регуляціям, адже зростає тиск з боку законодавства (GDPR, європейський ШІ-акт) й очікувань користувачів щодо реальної, а не декларативної приватності. Паралельно відбувається технологічний зсув до поєднання моделей на локальних пристроях і в захищених хмарних середовищах, що так чи інакше спричинює розвиток безпекових практик у роботі з ШІ моделями, включно з використанням вже існуючих підходів і знань у приватних хмарних обчисленнях.

2. ЩО TAKE PRIVATE CLOUD COMPUTING? ЙОГО ВИКОРИСТАННЯ З ШІ

Private Cloud Computing у загальному розумінні – це модель хмарних обчислень, у якій обчислювальні ресурси, сховище та мережа виділяються для однієї організації або вузького кола користувачів і працюють у контрольованому, ізольованому середовищі [1]. На відміну від публічних хмар, де інфраструктура спільна для багатьох клієнтів, приватна хмара дозволяє краще контролювати розміщення даних, політики безпеки, відповідність регуляціям та інтеграцію з внутрішніми системами [1]. У контексті сучасних систем штучного інтелекту поняття приватної хмари доповнюється вимогою до конфіденційних обчислень: дані мають бути захищені не лише “на диску” та “у мережі”, а й під час самого виконання моделей. Саме тому з’являються концепції на кшталт Private AI Compute або Private Cloud Compute, де хмарна інфраструктура доповнюється апаратно захищеними середовищами виконання (TEE), анонімізацією ідентичності користувача та обмеженням будь-якого доступу інфраструктурних операторів до вмісту запитів. У таких системах генеративні та аналітичні

моделі ШІ можуть працювати з чутливими персональними або корпоративними даними, фактично наближаючись за рівнем приватності до локальної обробки на пристрої. Водночас приватна хмара забезпечує еластичність і масштабованість, необхідні для великих моделей, які не можуть бути розгорнуті безпосередньо на кінцевих пристроях. Для користувача це відкриває можливість отримувати набагато вищий рівень якості відповідей від ШІ, не жертвуючи контролем над даними та відповідністю юридичним вимогам. Для розробників і архітекторів це означає потребу в спеціалізованих архітектурних паттернах, що поєднують традиційну хмарну інфраструктуру, конфіденційні обчислення та гібридну модель локального/хмарного інференсу.

3. ОПИС АРХІТЕКТУРИ

Як було згадано раніше, архітектура поєднує локальне виконання на пристрої користувача та захищену хмарну інфраструктуру, побудовану за принципами приватних віддалених обчислень.

У центрі взаємодії на стороні клієнта знаходиться локальний ШІ-маршрутизатор, який для кожного запиту визначає, чи достатньо можливостей вбудованої моделі на пристрої, чи потрібен виклик хмарних потужностей. Якщо запит може бути оброблений локально, він ніколи не покидає пристрій і виконується у захищеному середовищі локальної моделі.

Якщо ж потрібна хмара, роутер спершу намагається встановити захищене з'єднання з хмарию з проведенням атестації серверу, щоб впевнитись у його автентичності. У разі якщо ця перевірка провалюється, то з'єднання не встановлюється та користувач отримує помилку, або ми намагаємось обробити запит локально. У позитивному випадку ми встановлюємо з'єднання та спершу звертаємось до IP-blinding ретрансляторf, щоб приховати нашу реальну IP-адресу. Надалі такого роду атестація, але вже взаємна, сервісів проводиться для кожної спроби створити з'єднання між будь-якими компонентами у хмарі задля того, щоб гарантувати автентичність всіх модулів та уникнути випадкового витоку даних у неперевірені вузли, що можуть виконувати логіку, що не пройшла аудиту і/або не використовувати механізми для шифрування даних у пам'яті та ізольовані на рівні інфраструктури віртуальні машини.

Після IP-blinding ретранслятора, запит користувача потрапляє до сервісу обміну токенів, щоб замінити звичайний автентифікаційний токен на анонімний, щоб при обробці запиту вже на стороні наступних сервісів ми ніяк не викривали, чиї дані зараз знаходяться в роботі.

Після ретранслятора запит потрапляє до шлюзу, де встановлюється захищена сесія (з віддаленою атестацією сервера) і спершу перевіряється дійсність раніше випущеного анонімного токена та відразу після цього відбувається отримання списку доступних сервісів для проведення потрібних обчислень. Шлюз фактично виконує роль єдиної точки входу в систему, здійснює авторизацію та розподіляє трафік між кількома сервісами інференсу; взаємодія між самим шлюзом і внутрішніми сервісами відбувається через окреме захищене з'єднання (наприклад, mTLS/ALTS), враховуючи вже вищезгадану взаємну атестацію модулів.

Кожен сервіс інференсу запускає ШІ-модель усередині апаратно та програмно ізольованого середовища (TEE або конфіденційна віртуальна машина в залежності від імплементації та наявності потрібної інфраструктури), де дані завжди зберігаються в пам'яті в зашифрованому вигляді й недоступні ні для гіпервізора (так званий монітор віртуальних машин), ні для адміністраторів навіть у випадку так званих “broken glass” ситуацій, коли відбувається критична поломка в сервісі.

У результаті після інференсу запит та дані користувача потрапляють до кінцевої моделі ШІ, яка буде здійснювати обробку та поверне результат, який пройде той самий шлях у зашифрованому вигляді через всі раніше згадані вузли системи й повернеться до користувача.

Окремо варто зазначити, що також для більшої безпеки для кожного сервісу наявні правила вхідного та вихідного трафіку, які дозволяють загалом контролювати те, звідки сервіси можуть отримувати дані на вхід та що й куди вони можуть повертати як результат/до яких сервісів можуть звертатись.

Для простішого сприйняття концепту архітектури, побудуємо архітектурну діаграму, яка відобразить всі базові компоненти, що були описані раніше. Вона наведена на рисунку 1.

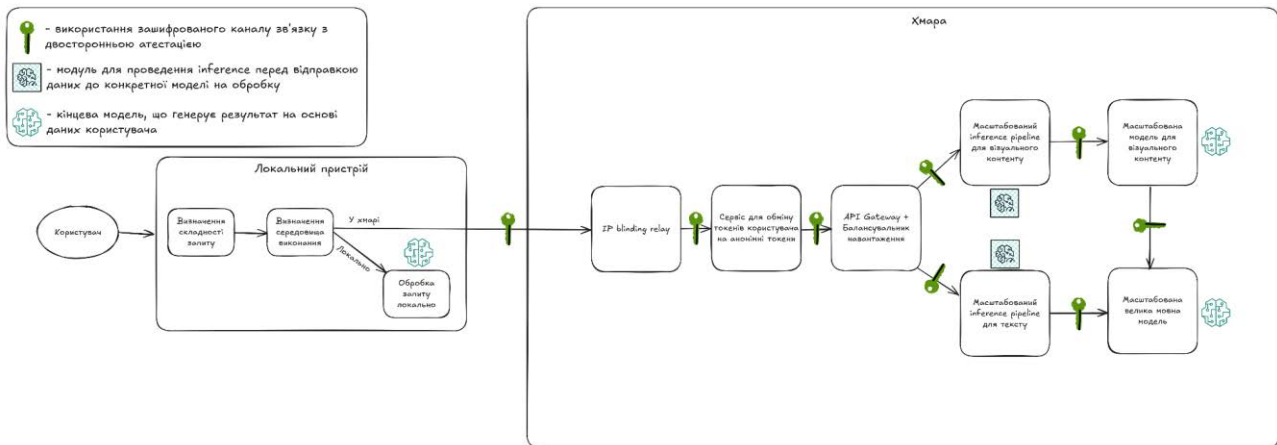


Рисунок 1. Архітектурна діаграма для системи приватних хмарних ШІ обчислень

Тепер розглянемо більш детально окремі частини нашої архітектури.

3.1. Локальний пристрій та роутер запитів

Раніше було згадано про основні дії, що виконуються на користувацькому пристрої, як-от смартфоні. Але також варто окреслити окремі важливі деталі роботи тих чи інших його частин.

Загалом локальний роутер запитів може враховувати велику кількість параметрів під час прийняття рішень щодо того, де потрібно виконати запит користувача, але серед них є звичайно базові, які мають вищий пріоритет і дозволяють отримати більш точну оцінку.

Одним з таких факторів є навантаження та енергоресурси пристрою: якщо пристрій вже знаходиться під значним навантаженням або його акумулятор має низький рівень заряду, то потенційно запит, особливо важкий, потрібно виконати в хмарі, а не локально.

Іншим важливим параметром може бути налаштування приватності користувача: якщо дані надчутливі і користувач заборонив їх надсилання, навіть ціною точності – роутер залишить обробку локально, або взагалі відмовиться виконувати запит у хмарі.

Також звичайно потрібно зважати на складність самого запиту: якщо він є доволі важким, що локальна модель або не зможе його повноцінно обробити через свої малі розміри, або буде це робити надзвичайно довго. Це наприклад, може стосуватись мультимодальних запитів, які включають у себе і текст, і візуальний контент, і навіть аудіо. У такому випадку запит варто надсилати в хмару, звичайно зважаючи на інші вищеописані фактори.

Говорячи про локальне виконання запитів користувача за допомогою ШІ, потрібно сказати, що ця обробка також повинна відбуватись в ізольованому середовищі. Наприклад, на Android, починаючи з 12 версії, доступний такий званий Private Compute Core, який

дозволяє ізолювати ІІІ-функції або інші, що потребують високої конфіденційності від решти системи подібно до того, як це робить TEE [2].

Для локального інференсу використовуються оптимізовані моделі: як правило, це так звані “TinyML” або “edge AI” моделі, що працюють через фреймворки на кшталт TensorFlow Lite, ONNX Runtime, PyTorch Mobile, Core ML тощо – залежно від платформи [3]. Ці інструменти дозволяють виконувати ІІІ моделі напряму на CPU/GPU/NPU пристрою, часто в поєднанні з апаратними засобами безпеки (наприклад, TrustZone або Secure Enclave у телефонах) для захисту пам’яті.

Таким чином, пристрій користувача з роутером забезпечує першу лінію приватності: максимум даних не покидає пристрій. Але коли потрібна допомога хмари, він передає запит так, щоб і надалі зберегти приватність через механізми, що були описані раніше.

3.2. IP-blinding ретранслятора, анонімізація токенів

Раніше вже було згадано про IP-blinding ретранслятора. Фактично цей механізм унеможливорює прив’язування запитів до конкретного пристрою чи геолокації за мережевими даними і суттєво ускладнює навіть потенційні атаки (наприклад, DDoS на конкретного користувача стає неможливим, а спроби відстежити активність – марними).

Для такого роду ретрансляторів можуть використовувати сторонні сервіси, що виступають у ролі проксі. Але загалом сам механізм часто працює за допомогою так званого Oblivious HTTP Proxy (ОНТТР), коли дані користувача подвійно шифруються: зовнішній шар знімає ретранслятор, не бачачи самих даних, а внутрішній шар зашифровано для кінцевого сервера [4]. Головний тут фактор полягає ж у тому, що ретранслятор має бути розташований поза інфраструктурою основного хмарного провайдера, щоб той не міг ідентифікувати, звідки саме прийшов запит користувача [4].

Іншим важливим кроком тут є анонімізація токенів, коли ми обмінюємо токени автентифікації користувача на ті, які так чи інакше від нього відв’язані, але в той же час підтверджують його право на запит. Загалом існує стандарт Privacy Pass, що описує цей механізм та те, як відбувається підтвердження валідності токenu [5].

3.3. Віддалена та взаємна атестації

Іншим ключовим механізмом організації довіри в архітектурах такого типу є віддалена атестація сервісів, що працюють у довіреному середовищі виконання (TEE). TEE – це апаратно захищене оточення з ізолюваною пам’яттю та мінімальним обсягом довіреного коду, яке гарантує конфіденційність і цілісність виконання навіть за компрометації операційної системи чи гіпервізора [6]. У процесі атестації таке середовище формує криптографічно підписаний звіт, що містить хеш коду, що на ньому працює та параметри конфігурації; віддалений учасник порівнює їх із еталонними значеннями та, у разі збігу, приймає рішення довіряти цьому вузлу [6]. Таким чином клієнт (або інший сервіс) не просто вірить сертифікату, який надає сервіс, а перевіряє, що на іншому боці дійсно запущено саме той код і в саме такому захищеному режимі, які очікується.

У розподілених системах, де кілька сервісів обмінюються чутливими даними, часто використовується взаємна атестація: кожна сторона перевіряє, що інший вузол також працює в належному TEE з коректним кодом. Після успішної (взаємної) атестації сторони узгоджують ключі шифрування й організовують захищений канал поверх стандартних протоколів. Типовим підходом є поєднання атестації з mutual TLS (mTLS): у mTLS обидві сторони з’єднання мають сертифікати й підтверджують володіння відповідними приватними ключами; це дає змогу аутентифікувати як сервер, так і клієнта й одночасно шифрувати трафік [7].

Часто в цьому контексті використовується Apache Teaclave – фреймворк для організації безпечних обчислень, який працює поверх TEE (Intel SGX) і використовує віддалену атестацію для встановлення довірених каналів зв'язку між клієнтами та обчислювальними вузлами[8]. У Teaclave клієнт може перевірити, що його код або функція виконуються в коректно атестованому enclave, після чого встановлюється зашифрований канал для передачі вхідних даних і отримання результатів [7]. На рівні мережі така взаємодія може бути реалізована у вигляді mTLS-з'єднання, де сертифікати пов'язані з атестованим станом enclave, тож будь-яка спроба замінити код або покинути TEE під час виконання призведе до зміни вимірів і, відповідно, до того, що атестація провалиться [6, 7].

4. ВИСНОВКИ

У цій статті було розглянуто архітектуру приватних хмарних обчислень, що забезпечує достатній рівень конфіденційності даних користувача при їх передачі та обробці ШІ. Архітектура поєднує локальні можливості пристрою користувача із захищеною хмарною інфраструктурою на базі TEE, анонімних токенів та IP-blinding. Ми також окреслили основні компоненти такої системи, шлях проходження запиту й механізми атестації, які дозволяють підтвердити коректність та захищеність середовища виконання. Особлива увага була приділена тому, як приватні хмарні обчислення дозволяють використовувати потужні моделі ШІ, не порушуючи конфіденційність і вимоги регуляцій, що робить цю тему критично важливою в умовах стрімкого розвитку генеративного ШІ. Загалом ж приватні хмарні обчислення є не лише технічною опцією, а й необхідною умовою для довіри користувачів та організацій до інтелектуальних сервісів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. What is a Private Cloud? Amazon Web Services. URL: <https://aws.amazon.com/what-is/private-cloud/> (дата звернення: 18.11.2025).
2. Frey S. Introducing Android's Private Compute Services. Google Online Security Blog, 09.09.2021. URL: <https://security.googleblog.com/2021/09/introducing-androids-private-compute.html> (дата звернення: 18.11.2025).
3. Krishnamoorthy M. V. Edge AI: TensorFlow Lite vs. ONNX Runtime vs. PyTorch Mobile. DZone, 03.06.2025. URL: <https://dzone.com/articles/edge-ai-tensorflow-lite-vs-onnx-runtime-vs-pytorch> (дата звернення: 19.11.2025).
4. Wood C., Hoyland J. Stronger than a promise: proving Oblivious HTTP privacy properties. The Cloudflare Blog, 27.10.2022. URL: <https://blog.cloudflare.com/stronger-than-a-promise-proving-oblivious-http-privacy-properties/> (дата звернення: 19.11.2025).
5. Privacy Pass (privacypass). About. IETF Datatracker. URL: <https://datatracker.ietf.org/wg/privacypass/about/> (дата звернення: 19.11.2025).
6. Sabt M., Achemlal M., Bouabdallah A. Trusted Execution Environment: What It is, and What It is Not. In: 2015 IEEE Trustcom/BigDataSE/ISPA. 2015. URL: https://hal.science/hal-01246364/file/trustcom_2015_tee_what_it_is_what_it_is_not.pdf (дата звернення: 20.11.2025).
7. What is mutual TLS (mTLS)? Cloudflare Learning Center. URL: <https://www.cloudflare.com/learning/access-management/what-is-mutual-tls/> (дата звернення: 20.11.2025).
8. Sun M. Teaclave: A Universal Secure Computing Platform. ApacheCon 2020, 30.09.2020. URL: <https://mssun.me/assets/teaclave.pdf> (дата звернення: 20.11.2025).

ВПЛИВ СТРАТЕГІЙ ДЕКОДУВАННЯ ТОКЕНІВ У ВЕЛИКИХ МОВНИХ МОДЕЛЯХ НА ЯКІСТЬ TEXT-TO-SQL ПАРСИНГУ

Курганський Л.С.¹, Шаптала Р.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ leokurg-ipt23@iit.kpi.ua, ² r.shaptala@gmail.com

У роботі досліджується вплив стратегій декодування токенів у великих мовних моделях на якість генерації SQL-запитів у задачі Text-to-SQL. На основі моделі Qwen2.5-3B-Instruct проведено порівняння базових (Greedy, Beam, Top-k, Top-p) та покращених execution-guided стратегій (EG-Beam, EGLA-Beam, EFG-Beam) на наборі даних mini_dev із колекції BIRD-Bench. Для оцінювання використано метрики Execution Accuracy, String Match Accuracy, Component Match Accuracy, AST Similarity та середній час генерації. Результати показали, що інтеграція виконувальних перевірок і механізмів самокорекції суттєво підвищує частку семантично коректних і виконуваних запитів без зміни самої моделі.

Ключові слова: Text-to-SQL, великі мовні моделі, стратегії декодування, execution-guided декодування, SQL-парсинг.

1. ВСТУП

Розвиток великих мовних моделей (LLM) відкрив можливість автоматичної побудови SQL-запитів із природномовних формулювань, що є ключовим кроком до інтелектуального доступу користувачів до реляційних баз даних без знання мови SQL. У такій постановці задача Text-to-SQL зводиться до генерації формально коректного та семантично релевантного запиту на основі текстового опису інформаційної потреби користувача.

Попри прогрес моделей, кінцевий результат істотно залежить від того, як саме з розподілу ймовірностей наступного токена обирається послідовність вихідних символів. Стратегія декодування визначає траєкторію пошуку по простору можливих SQL-запитів, впливаючи на синтаксичну валідність, структурну повноту, семантичну коректність і обчислювальну вартість генерації. Базові алгоритми, спочатку розроблені для задач природномовної генерації [1], не враховують специфіку формальних мов із жорсткою граматикую та вимогою виконаності запитів.

Для подолання цих обмежень запропоновано execution-guided підходи [2], які поєднують класичний пошук із динамічною перевіркою проміжних версій запиту відносно баз даних. Метою роботи є кількісна оцінка впливу різних стратегій декодування токенів на якість Text-to-SQL парсингу для реалістичного набору задач і побудова практичних рекомендацій щодо вибору стратегії з урахуванням балансу між точністю, стабільністю та часом генерації.

2. МЕТОДОЛОГІЯ ДОСЛІДЖЕННЯ

2.1. Постановка задачі та вибір бенчмарку

Задача Text-to-SQL формулюється в парадигмі семантичного парсингу: для кожного природномовного запиту необхідно побудувати SQL-запит, що повертає очікуваний результат. У роботі використано набір mini_dev із колекції BIRD-Bench [3, 4], який є

спрощеною, але репрезентативною підмножиною повного BIRD. Він містить пари «природномовний запит – еталонний SQL-запит» для реляційних баз даних різної структури у форматі SQLite, що дозволяє безпосередньо перевіряти виконуваність згенерованих запитів. Особливістю BIRD є наявність неоднозначних формулювань і реалістичних бізнес-орієнтованих сценаріїв, що робить цей набір практичним.

2.2. Вибір мовної моделі та стратегій декодування

Як базову генеративну модель обрано Qwen2.5-3B-Instruct [5] – сучасну LLM, налаштовану на інструкційні промпти та генерацію програмного коду, включно з SQL. Модель підтримує керування параметрами декодування, що дає змогу реалізувати та порівняти різні стратегії в однакових умовах.

Досліджувалися такі стратегії декодування:

- *Greedy Search* – вибір найбільш ймовірного токена на кожному кроці [1];
 - *Beam Search* – підтримка кількох найперспективніших гілок генерації з подальшим відбором найкращої послідовності [1];
 - *Top-k Sampling* – випадковий вибір токена з обмеженого набору k найімовірніших [1];
 - *Top-p (nucleus) Sampling* – вибір із найменшого підмножини tokenів, сумарна ймовірність яких не менша за p [1];
 - *Execution-Guided Beam Search* – beam-пошук із фільтрацією гілок за результатом проміжних перевірок часткових SQL-запитів відносно бази даних [2];
 - *Execution-Guided Look-Ahead Beam Search* – варіант EG-Beam із заглибленням наперед для перевірки кількох майбутніх кроків;
 - *Execution-Fixing-Guided Beam Search* – стратегія, що поєднує execution-guided перевірки з модулем автоматичного виправлення синтаксичних помилок у проміжних запитах.
- Усі стратегії застосовувалися до тієї самої моделі й в однакових умовах; змінювалися лише алгоритм вибору tokenів і пов'язані параметри.

2.3. Метрики оцінювання якості SQL-генерації

З метою об'єктивного порівняння стратегій було використано уніфікований процес оцінювання, що обчислював такі метрики:

- *Execution Accuracy* – частка згенерованих запитів, що виконуються без помилки й повертають результат, подібний до еталонного [6];
- *String Match Accuracy* – повний збіг з еталонним SQL-рядком [6];
- *Component Match Accuracy* – часткова подібність між структурними компонентами запиту (списки полів, таблиці, предикати, агрегації) [6];
- *AST Similarity* – схожість абстрактних синтаксичних дерев згенерованого та еталонного запитів;
- *Average generation time* – середня тривалість побудови одного SQL-запиту для кожної стратегії.

Підсумкові значення метрик отримано шляхом агрегування результатів по всіх 500 прикладах `mini_dev`, що дозволило порівнювати стратегії за точністю, структурною якістю та латентністю.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

3.1. Узагальнена таблиця результатів

Результати тестування моделі Qwen2.5-3B-Instruct на наборі даних `mini_dev` (BIRD) для всіх розглянутих стратегій декодування наведено у таблиці 1. Для кожної стратегії обчислено

Execution Accuracy, String Match Accuracy, Component Match Accuracy, AST Similarity та Generation time.

Таблиця 1. Результати експериментів

Стратегія декодування	Execution Accuracy	String Match Accuracy	Component Match Accuracy	AST Similarity	Generation time
Greedy	0.348	0.022	0.3708	0.1284	47.71
Beam	0.392	0.022	0.3734	0.2963	153.24
Top-k	0.342	0.022	0.3603	0.2894	63.7
Top-p	0.352	0.022	0.3601	0.2944	67.96
EG-Beam	0.416	0.022	0.3816	0.2924	168.72
EGLA-Beam	0.416	0.024	0.3805	0.2939	157.67

3.2. Аналіз метрик для різних стратегій декодування

Отримані результати демонструють, що стратегія декодування суттєво впливає на всі ключові аспекти якості Text-to-SQL генерації: виконувальність запитів, точність збігу з еталоном, структурну схожість та час побудови відповіді.

За Execution Accuracy найгірші показники мають стохастичні стратегії Top-k (0.342) та Top-p (0.352), які майже не перевершують Greedy (0.348). Beam Search дає помітний приріст до 0.392 завдяки розгляду кількох паралельних гілок генерації. Найвищі значення Execution Accuracy досягаються для виконувальних стратегій: EG-Beam та EGLA-Beam (0.416), а також EFG-Beam (0.47), що підтверджує переваги поєднання пошуку з динамічною перевіркою SQL-запитів і механізмами корекції помилок.

String Match Accuracy для всіх стратегій (окрім EFG-Beam) стабільно тримається на рівні 0.022–0.024, а для EFG-Beam дорівнює 0. Це очікуваний результат для задачі Text-to-SQL: навіть при семантично еквівалентних запитах незначні відмінності у форматуванні або порядку умов призводять до втрати повного текстового збігу. Тому метрика String Match більше відображає стабільність синтаксису, а не реальну функціональну правильність запитів; нульове значення для EFG-Beam пов'язане з агресивною корекцією синтаксису, яка змінює структуру рядка, але не обов'язково погіршує семантику.

Component Match Accuracy лежить у діапазоні 0.3601–0.3816 для всіх підходів. Базові стратегії Greedy, Beam, Top-k, Top-p показують значення близько 0.36–0.37, тоді як EG-Beam (0.3816) та EGLA-Beam (0.3805) досягають найвищих показників. Це означає, що виконувальні модифікації краще зберігають структурні блоки запиту, що особливо важливо для складних SQL із підзапитами та багатьма приєднаннями. EFG-Beam (0.3791) трохи поступається EG/EGLA за цією метрикою, але різниця не є суттєвою.

За AST Similarity найкращі результати спостерігаються для Beam (0.2963), Top-p (0.2944) та EFG-Beam (0.2958), тоді як інші стратегії дають значення в діапазоні 0.2894–0.2939. Це свідчить про те, що всі вдосконалені методи формують подібні логічні структури SQL-запитів, навіть якщо текстове представлення відрізняється. Виконувальні перевірки не погіршують структурну узгодженість, а радше допомагають зберегти коректну форму запиту при виправленні помилок.

За середнім часом генерації найшвидшою є Greedy (47.71 с на запит), що відповідає її лінійній та детермінованій природі. Стохастичні Top-k та Top-p мають більший час (63.70 та 67.96 с відповідно) через додаткову випадковість і більший простір пошуку. Beam Search істотно повільніший (153.24 с) через паралельну обробку кількох кандидатних послідовностей. Execution-guided стратегії є найдорожчими за часом: EG-Beam (168.72 с),

EGLA-Beam (157.67 с) та EFG-Beam (181.64 с), оскільки потребують багаторазових перевірок проміжних SQL-запитів. Таким чином, підвищення точності та виконуваності супроводжується збільшенням часу генерації, однак це зростання є прийнятним.

3.3. Рекомендації щодо вибору оптимальної стратегії для задач Text-to-SQL

Отримані результати дозволяють сформулювати практичні рекомендації щодо вибору стратегії декодування залежно від вимог до системи:

- *Сценарії з пріоритетом максимальної точності та виконуваності*: Для аналітичних систем, офлайн-звітності та критично важливих BI-сценаріїв доцільно використовувати EFG-Beam як основну стратегію. Вона забезпечує найвищу Execution Accuracy (~0.47) завдяки поєднанню execution-guided відсікання невалідних гілок і автоматичного виправлення синтаксичних помилок. Недолік – підвищена латентність, яка, однак, є прийнятною для офлайн та напівінтерактивних задач.

- *Інтерактивні системи з вимогами до якості та помірної латентності*: У чат-інтерфейсах до баз даних або інтерактивних аналітичних панелях доцільний компроміс – EG-Beam або EGLA-Beam. Вони істотно покращують виконуваність порівняно з Greedy/Beam, але менш «важкі», ніж EFG-Beam, завдяки відсутності окремого кроку виправлення. EGLA-Beam є доцільним, коли важлива стійкість до локальних помилок у кількох наступних токенах.

- *Сценарії з жорсткими обмеженнями на час відповіді та ресурси*: Для високонавантажених сервісів або прототипів, де критичними є швидкість і простота реалізації, доцільно застосовувати Greedy Search або класичний Beam Search з невеликим розміром beam. Це забезпечує прийнятну точність для простіших запитів за мінімального часу генерації.

- *Задачі з високою семантичною неоднозначністю*: Якщо природномовні запити часто мають кілька інтерпретацій, корисним є використання Top-k або Top-r для генерації множини кандидатних SQL-запитів. Далі ці кандидати можуть фільтруватися за додатковими моделями ранжування. У такому випадку стохастичні стратегії не стільки підвищують середню точність, скільки збільшують покриття простору можливих інтерпретацій.

На практиці доцільно застосовувати адаптивні схеми, у яких система починає з швидкої стратегії, а у випадку невдалого виконання або низької впевненості автоматично переходить до дорожчих, але надійніших execution-guided варіантів. Такий підхід дозволяє поєднати низьку латентність із високою якістю для складних випадків.

4. ВИСНОВКИ

У роботі виконано комплексне дослідження впливу стратегій декодування токенів на якість Text-to-SQL парсингу. На основі експериментів показано, що саме вибір стратегії декодування є критичним чинником, який визначає баланс між синтаксичною валідністю, структурною цілісністю та семантичною коректністю згенерованих SQL-запитів. Базові підходи забезпечують прийнятний рівень точності переважно для простіших запитів, тоді як для складних структур вони часто призводять до невиконуваних або хибних результатів. Додавання execution-guided механізмів перевірки та виправлення дозволяє суттєво підвищити Execution Accuracy, зберігаючи при цьому структурну узгодженість запитів і не вимагаючи змін у самій моделі.

Отримані результати дають підстави розглядати EFG-Beam як оптимальну стратегію декодування для систем Text-to-SQL, орієнтованих на максимальну точність, тоді як Greedy або класичний Beam Search доцільно застосовувати в сценаріях із жорсткими обмеженнями на час відповіді, де допустиме певне зниження якості заради латентності. Стохастичні стратегії доцільно використовувати для генерації множини альтернативних SQL-кандидатів із подальшою фільтрацією за бізнес-правилами або додатковими моделями ранжування.

Перспективні напрями подальших досліджень пов’язані з розробкою комбінованих та адаптивних стратегій, які поєднують execution-guided перевірки із семантичним аналізом, мультимодельним узгодженням, а також із переносом описаних підходів на інші мови запитів і домени даних.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Holtzman A., Buys J., Forbes M., Choi Y. The Curious Case of Neural Text Degeneration. arXiv:1904.09751, 2019.
2. Wang C., Tatwawadi K., Brockschmidt M., Polozov O., Richardson M. Robust Text-to-SQL Generation with Execution-Guided Decoding. arXiv:1807.03100, 2018.
3. BIRD-Bench Project. BIRD: BIg Bench for LaRge-scale Database-grounded Text-to-SQL Evaluation. Офіційна документація бенчмарку BIRD-SQL.
4. BIRD-Bench Team. BIRD-Mini-Dev V2 Dataset. Офіційний репозиторій набору mini_dev.
5. Qwen Team. Qwen2.5-3B-Instruct: Instruction-tuned Large Language Model. Модельна картка та документація серії Qwen2.5, 2024–2025.
6. Liu X., Shen S., Li B., Ma P., Jiang R., Zhang Y., Fan J., Li G., Tang N., Luo Y. A Survey of Text-to-SQL in the Era of LLMs: Where are we, and where are we going? arXiv:2408.05109, 2024

ІНТЕЛЕКТУАЛЬНІ АГЕНТИ ДЛЯ СИСТЕМ АВТОМАТИЗАЦІЇ УПРАВЛІННЯ БІЗНЕС-ПРОЦЕСАМИ

Мухін В.Є.¹, Лимар О.В.²

Національний технічний університет України
“Київський політехнічний інститут ім. Ігоря Сікорського”, Київ, Україна

¹ v.mukhin@kpi.ua, ² olymar17@gmail.com

У роботі розглянуто підходи до створення та використання інтелектуальних агентів у системах автоматизації управління бізнес-процесами (BPMS). Метою дослідження є розроблення комплексної архітектури агентного управління, що забезпечує адаптивність, оптимізацію та автономність бізнес-процесів. Використано методи багатоагентного моделювання, машинного навчання та інтелектуального аналізу даних. Запропонований підхід дає змогу автоматично виявляти неефективності процесів, прогнозувати навантаження, оптимізувати розподіл ресурсів і підтримувати прийняття управлінських рішень у реальному часі. Наукова новизна полягає у поєднанні агентної архітектури з механізмами когнітивного аналізу даних та динамічної координації процесів. Практична значимість полягає у підвищенні ефективності, швидкодії та гнучкості корпоративних BPMS.

Ключові слова: інтелектуальні агенти, BPMS, автоматизація процесів, багатоагентні системи, оптимізація бізнес-процесів.

1. ВСТУП

Сучасні бізнес-процеси у виробничих і сервісних компаніях стають дедалі більш складними та динамічними: зростає кількість учасників ланцюгів постачання, посилюються вимоги до швидкості обробки замовлень, виникає потреба оперативно реагувати на зміни попиту й ринкових умов. За таких умов традиційні системи управління бізнес-процесами (англ. *Business Process Management System, BPMS*), орієнтовані переважно на жорстко формалізовані регламенти та фіксовані бізнес-правила, виявляють обмежену гнучкість. Значна частина налаштувань виконується вручну, а прийняття рішень часто супроводжується затримками, що призводить до простоїв, перевитрат ресурсів та погіршення якості обслуговування клієнтів.

Попри розвиток роботизованої автоматизації процесів (англ. *Robotic Process Automation, RPA*), систем керування бізнес-правилами (англ. *Business Rules Management Systems, BRMS*) та модулів процесної аналітики, більшість існуючих BPMS залишаються реактивними: вони добре працюють у стабільних умовах, але погано адаптуються до стохастичних змін, не вміють самостійно прогнозувати навантаження, оптимізувати закупівлі чи розподіл ресурсів. Це зумовлює потребу в підходах, здатних поєднати процесну модель із механізмами аналізу даних, навчання та децентралізованого прийняття рішень [1–3].

Інтелектуальні агенти розглядаються як перспективний інструмент побудови таких систем. Агент може автономно сприймати стан середовища, взаємодіяти з іншими агентами, використовувати алгоритми оптимізації та машинного навчання, накопичувати досвід і змінювати стратегію поведінки. У контексті BPMS це відкриває можливості для автономної

координації бізнес-процесів, динамічного вибору постачальників, адаптивного планування закупівель та ресурсів у реальному часі [4–5].

Метою дослідження є розроблення методів і засобів використання інтелектуальних агентів у BPMS для підвищення ефективності управління бізнес-процесами, зокрема процесом «закупівля матеріалів під замовлення». Запропонований підхід спрямований на те, щоб забезпечити адаптивність, оптимальність і автономність прийняття рішень у корпоративних системах управління процесами.

Наукова новизна одержаних результатів полягає в розробленні агентного підходу до управління процесом закупівель, у межах якого модифіковано протокол договірної мережі (англ. *Contract Net Protocol, CNP*) шляхом адаптивної зміни ваг критеріїв відбору постачальників на основі історичних даних. На відміну від відомих JADE-орієнтованих моделей та агентних систем управління бізнес-процесами, запропоновано інтегроване використання онтології предметної області з евристичним оптимізаційним алгоритмом і побудовано агентно орієнтовану архітектуру з розподіленою обробкою подій процесного рушія, що забезпечує підвищення гнучкості та ефективності управління бізнес-процесами.

2. АНАЛІЗ ІСНУЮЧИХ РІШЕНЬ

Сучасні системи управління бізнес-процесами (Camunda, Bizagi, BonitaSoft, IBM BPM, Appian, Pega тощо) забезпечують моделювання процесів у стандарті моделі та нотації бізнес-процесів (англ. *Business Process Model and Notation, BPMN*), підтримують правила, задані у нотації моделювання рішень (англ. *Decision Model and Notation, DMN*), інтеграцію з системами планування ресурсів підприємства (англ. *Enterprise Resource Planning, ERP*) та системами управління взаємовідносинами з клієнтами (англ. *Customer Relationship Management, CRM*), а також парадигму розроблення з мінімальним або відсутнім програмуванням. Вони істотно підвищують прозорість і контрольованість процесів, однак зберігають централізовану логіку: основні рішення приймає єдиний процесний рушій, а адаптація до змін середовища потребує ручного перепроєктування схем. Оптимізація зазвичай виконується постфактум, на основі звітів, без можливості оперативного впливу на виконання процесів у реальному часі.

Подальший етап розвитку BPMS пов'язаний з інтелектуалізацією – з появою інтелектуальних систем управління бізнес-процесами (англ. *intelligent Business Process Management Suites, iBPMS*), інтеграцією модулів штучного та машинного навчання (англ. *Machine Learning, ML*), RPA та засобів процесної аналітики бізнес-процесів. Такі рішення дозволяють прогнозувати затримки, аналізувати вузькі місця, автоматично розподіляти завдання між виконавцями, формувати рекомендації щодо покращення процесів. Проте навіть iBPMS залишаються лише частково інтелектуальними: ML-модулі найчастіше працюють як зовнішні сервіси аналітики, що не змінюють принципово централізований характер управління й не забезпечують самостійного планування дій у динамічному середовищі.

Паралельно розвиваються агентоорієнтовані підходи, що можуть бути реалізовані у фреймворках JADE, SPADE/PADE та інших середовищах багатоагентного управління бізнес-процесами. У таких системах інтелектуальні агенти розглядаються як автономні програмні сутності, що взаємодіють через стандартизовані протоколи FIPA-ACL, можуть координувати свої дії, узгоджувати завдання та розподіляти ресурси. Агентні компоненти інтегруються з BPMS у вигляді гібридних архітектур: від моделі, де агенти виконують роль «інтелектуальних виконавців» окремих задач, до підходу, коли саме мультиагентна система бере на себе оркестрацію процесів.

Водночас існуючі рішення залишаються фрагментованими. Інтеграція ML-модулів, RPA та агентів часто реалізується точково, без єдиної когнітивної моделі процесу. Більшість BPMS не забезпечують повної автономності агентів: стратегічні рішення закріплені у процесних схемах і бізнес-правилах, а агенти виконують допоміжну роль. Унаслідок цього рівень когнітивності, самоорганізації та децентралізованого прийняття рішень у сучасних BPMS є недостатнім, що і визначає наукову проблему, яку розв'язує дана робота – розроблення узгодженої агентної архітектури з глибокою інтеграцією механізмів аналізу даних та адаптивної координації процесів.

2.1. Постановка задачі

Об'єктом дослідження є процес «закупівля матеріалів під замовлення» у системі управління бізнес-процесами підприємства. Предметом дослідження є методи агентного управління цим процесом з використанням евристично-адаптивних алгоритмів планування закупівель та вибору постачальників.

Метою є побудова агентного модуля, який, взаємодіючи з BPMS, забезпечує для заданого горизонту планування мінімізацію сукупних витрат на закупівлю і зберігання матеріалів, штрафів за прострочення та дефіцит, а також дотримання вимог за строками виконання замовлень (англ. *service level agreement*, SLA) за умов стохастичного попиту і невизначених строків постачання.

Розглядається множина замовлень $O = \{o_1, \dots, o_n\}$, множина матеріалів $M = \{m_1, \dots, m_k\}$ і множина постачальників $S = \{s_1, \dots, s_p\}$. Кожне замовлення o_i має термін виконання, потрібні обсяги матеріалів та пріоритет. Для кожного постачальника задано ціну, допустимий діапазон обсягів, статистику строків постачання та надійність. Склад має обмежену місткість, а бюджет закупівель обмежений виділеними фінансовими ресурсами.

У межах агентного підходу вводиться множина агентів $A = \{\text{OrderAgent}, \text{InventoryAgent}, \text{ProcurementAgent}, \text{SupplierAgent}, \text{FinanceAgent}, \text{AnalyticsAgent}\}$, які оперують спільною онтологією предметної області (замовлення, матеріал, постачальник, поставка, склад, контракт). Їхні рішення щодо моментів розміщення замовлень, вибору постачальників і обсягів закупівель повинні задовольняти обмеження на місткість складу, бюджет і SLA, а також формувати послідовність подій, сумісну з процесною моделлю BPMS.

Формально задача полягає в синтезі правил поведінки агентів і параметрів оптимізаційних алгоритмів, за яких для будь-якого допустимого сценарію попиту й параметрів постачань функція витрат підприємства на заданому інтервалі часу набуває мінімального значення, а показники якості процесу (тривалість циклу, кількість простоїв, рівень виконання замовлень у строк) задовольняють заданим нормативам.

3. МЕТОДИ ПОБУДОВИ ІНТЕЛЕКТУАЛЬНИХ АГЕНТІВ

3.1. Математичне та агентне моделювання бізнес-процесів

У запропонованому підході бізнес-процес «закупівля матеріалів під замовлення» формалізується як багатоагентна система. Множина агентів A відповідає основним учасникам процесу. Кожен агент має власний набір станів, локальну цільову функцію та доступні дії, а глобальна мета системи визначається як мінімізація сукупних витрат за умови своєчасного виконання замовлень і дотримання ресурсних обмежень.

Взаємодія агентів описується як сукупність подій і повідомлень, що передаються за стандартизованими протоколами. На рівні моделей поведінки розрізняються режими кооперації (узгоджена робота InventoryAgent, ProcurementAgent і FinanceAgent для забезпечення замовлення), конкуренції (SupplierAgent-и змагаються за контракт у межах модифікованого протоколу договірної мережі) та делегування завдань (OrderAgent ініціює

підпроцеси оптимізації закупівель і вибору постачальників). Така постановка дає змогу перейти від жорстко детермінованих схем BPMN до децентралізованого узгодження рішень між автономними суб'єктами.

Для підтримки формалізації будується онтологія предметної області, що включає сутності «замовлення», «матеріал», «постачальник», «склад», «контракт», а також їх атрибути – обсяги, строки, ціни, ризики, пріоритети. Онтологія використовується як спільна база знань для всіх агентів: вона задає єдині типи даних і відношення між ними, забезпечує узгодженість інтерпретації подій та дозволяє інтегрувати алгоритмічні модулі прогнозування й оптимізації у єдину мультиагентну модель бізнес-процесу.

3.2. Алгоритми навчання та адаптації агентів

Навчання агентів у запропонованій архітектурі реалізується через вбудовані механізми адаптації в оптимізаційних та координаційних алгоритмах. Оптимізаційний агент *ProcurementAgent* використовує модифікований Бджолиний алгоритм, у якому кожна ітерація розглядається як крок навчання на історії виконаних закупівель. На основі значення функції витрат агент поступово відхиляє неякісні стратегії поповнення запасів і концентрується на комбінаціях обсягів та строків, що дають кращі результати в умовах змінного попиту та нестабільних строків постачання.

Другий рівень адаптації реалізовано в механізмі модифікованого протоколу договірної мережі (англ. *Adaptive Contract Net Protocol, Adaptive CNP*) для вибору постачальників. Тут застосовується багатокритеріальна функція корисності, у якій вагові коефіцієнти $w_k(t)$ оновлюються за експоненційною схемою згладжування з урахуванням реальних результатів співпраці (фактичний час поставки, відсоток відхилень від бюджету, стабільність якості). Таким чином, система реалізує навчання на основі досвіду: з кожною новою угодою оцінки постачальників коригуються, а агенти дедалі частіше обирають партнерів, що демонструють кращу поведінку в довгостроковій перспективі.

Такий евристично-адаптивний підхід забезпечує ефект машинного навчання на рівні мультиагентної взаємодії, не вимагаючи побудови окремих складних моделей прогнозування. У подальших дослідженнях ці механізми можуть бути доповнені класичними ML-моделями (наприклад, для прогнозу попиту чи строків поставки), але в межах представленої роботи основний акцент зроблено саме на адаптації параметрів Бджолиним алгоритмом та *Adaptive CNP*.

3.3. Алгоритми оптимізації бізнес-процесів

Оптимізація в запропонованій агентній архітектурі орієнтована на мінімізацію сукупних витрат закупівель, скорочення тривалості циклу виконання замовлення та зменшення ризику дефіциту матеріалів. Для задач такого типу традиційно застосовують генетичні алгоритми, рійні алгоритми (*Particle Swarm Optimization, Ant Colony, Bees Algorithm*) та методи нелінійного програмування. У роботі зроблено акцент на рійних методах, оскільки вони краще пристосовані до динамічного середовища й легко інтегруються у багатоагентні системи.

Оптимізаційний агент використовує модифікований Бджолиний алгоритм для пошуку раціональних стратегій поповнення запасів: кандидатні рішення кодують обсяги та строки замовлень для різних постачальників, а цільова функція враховує витрати на закупівлю, зберігання, штрафи за дефіцит і запізнення поставок. На основі оцінки якості рішень алгоритм поступово відхиляє неефективні комбінації та підвищує пріоритет тих стратегій, що забезпечують менші витрати й стабільніше виконання замовлень.

Розподіл і перерозподіл ресурсів між учасниками процесу реалізується через взаємодію спеціалізованих агентів – OrderAgent, InventoryAgent, ProcurementAgent, SupplierAgent. Вони узгоджують плани закупівель і поставок, реагуючи на зміни попиту та рівня запасів, що дає змогу зменшити простой та кількість термінових замовлень. Координація здійснюється із застосуванням протоколів сімейства FIPA (зокрема, CNP), які визначають формати повідомлень та правила проведення переговорів щодо вибору постачальника й прийняття оптимізаційних рішень у розподіленому середовищі.

3.4. Обмеження запропонованого методу

Попри високу ефективність, запропонований метод має низку обмежень, пов'язаних з обчислювальною складністю та припущеннями моделі. Модифікований Бджолиний алгоритм потребує багато ітерацій за великої номенклатури матеріалів, а Adaptive CNP створює помітне комунікаційне навантаження за великої кількості постачальників. Коректність результатів залежить від наявності достатнього масиву історичних даних і відносно стабільних логістичних умов. Підхід є найбільш доцільним для середніх компаній із циклічними закупівлями, тоді як у великих корпоративних структурах може знадобитися кластеризація агентів або паралельна обробка, а інтеграція передбачає використання BPMN-сумісних систем, де агенти реалізуються як зовнішні сервісні задачі.

4. ПРОГРАМНІ ЗАСОБИ ТА АРХІТЕКТУРА AGENCY-BASED BPMS

Запропонована система має багаторівневу архітектуру, що поєднує класичний процесний двигун BPMS та шар інтелектуальних агентів. На рівні даних функціонують модулі збирання та зберігання інформації про замовлення, запаси, поставки, фінансові операції та історію взаємодії з постачальниками. Поверх них працює процесний двигун, який виконує BPMN-схеми основних бізнес-процесів, а також агентний шар, представлений виконавчими агентами (OrderAgent, InventoryAgent, ProcurementAgent, SupplierAgent, FinanceAgent) та агентами-аналітиками (AnalyticsAgent), що відповідають за оптимізацію та оцінку ефективності.

Імітаційна модель реалізована мовою програмування Python у вигляді окремого програмного модуля. Кожен агент реалізовано як окремий клас із власним набором станів і методів обробки повідомлень. Взаємодія організована через внутрішній диспетчер подій, що забезпечує обмін повідомленнями між агентами в межах єдиного процесу. Симуляційний час є дискретним і задається у вигляді кроків моделювання, протягом яких послідовно обробляються події створення замовлень, поставок матеріалів та оновлення фінансових показників.

Інтеграція з корпоративними системами (ERP, CRM, SCM) може здійснюватись через адаптери та API-шлюзи. ERP-система надає дані про залишки та рух матеріалів, CRM – про замовлення клієнтів і прогнози попиту, а системами управління ланцюгами постачання (англ. *Supply Chain Management, SCM*) – про контракти та умови співпраці з постачальниками. Агенти використовують ці дані для формування рішень щодо обсягів закупівель, вибору постачальника та планування поставок, при цьому BPMS гарантує узгодженість виконання процесів на рівні всієї організації.

Програмна реалізація орієнтована на подальше хмарне розгортання з використанням мікросервісної архітектури. Кожна група агентів, процесний двигун та сервіси роботи з даними розгортаються як окремі контейнери, що управляються Kubernetes-кластером. Це спрощує масштабування системи за кількістю одночасно активних агентів, забезпечує

відмовостійкість та дає змогу гнучко виділяти ресурси під окремі підсистеми (наприклад, оптимізаційний чи аналітичний модуль).

Для візуалізації й моніторингу роботи агентно-орієнтованої BPMS використовуються інформаційні панелі (англ. *dashboards*), на яких відображаються ключові показники (англ. *KPI*) – час циклу замовлення, рівень сервісу, завантаження складу, частка термінових поставок, динаміка витрат. Логи процесного двигуна та повідомлень між агентами додатково аналізуються засобами процесної аналітики, що дає змогу виявляти вузькі місця, оцінювати вплив оптимізаційних рішень та формувати рекомендації для подальшого вдосконалення моделі управління бізнес-процесами.

5. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

Експериментальні дослідження проводилися в середовищі імітаційної моделі, у якій реалізовано агентний модуль управління процесом «закупівля матеріалів під замовлення». Розглядалися дві конфігурації: стандартний агентний модуль та модифікований модуль з доданими механізмами інтелектуальності. Стандартний варіант емулює роботу типового процесного рушія BPMS: рішення щодо обсягів і строків закупівель приймаються за фіксованими правилами (статичні точки замовлення, жорстко задані пріоритети постачальників, відсутність адаптації до історичних результатів). У модифікованому варіанті задіяні евристично-адаптивні алгоритми – модифікований Бджолиний алгоритм для планування закупівель і Adaptive CNP для вибору постачальників.

Порівняння стандартного та модифікованого агентних модулів здійснювалося на однакових вхідних сценаріях потоку замовлень із коливаннями попиту та змінними строками поставок. Оцінювалися такі метрики: середній час циклу виконання замовлення (від моменту створення до повного закриття), завантаження ресурсів складу, кількість простоїв через відсутність матеріалів, частка термінових закупівель, а також дотримання SLA за строками виконання замовлень. Додатково аналізувався економічний ефект: порівнювалися сумарні витрати на закупівлі та фінансовий результат моделі за однаковий проміжок симуляційного часу, що дає змогу оцінити, чи забезпечує модифікований агентний модуль ефективніше використання бюджету та потенційне зростання прибутковості компанії.

Результати показали, що модифікований агентний модуль дає змогу скоротити середню тривалість процесу на 10–20 % порівняно зі стандартною конфігурацією, зменшити кількість простоїв через дефіцит матеріалів більш ніж удвічі та істотно знизити частку термінових закупівель. Рівень дотримання SLA за строками виконання замовлень зріс до показників понад 95 % у порівнянні з базовим режимом. Спостерігалось більш рівномірне завантаження складу та стійкіші рішення щодо вибору постачальників, що підтверджує ефективність адаптивного агентного управління й доцільність подальшої інтеграції такого модуля з реальними BPMS.

Практичне застосування одержаних результатів полягає в можливості використання розробленого агентного модуля управління закупівлями в інформаційних системах підприємств із великою номенклатурою матеріалів та значною варіативністю попиту, де необхідна автоматизована підтримка прийняття рішень щодо обсягів і строків замовлень. Підхід є особливо корисним для виробництв, що працюють за підходом «Саме вчасно» (англ. *just-in-time, JIT*), оскільки дозволяє зменшувати надлишкові запаси та ризик простоїв через дефіцит ресурсів за рахунок адаптивного планування поставок. Окремим класом сценаріїв є дистрибуційні компанії з розгалуженою складською мережею, де агентна система може використовуватися для координації закупівель і перерозподілу запасів між складами з урахуванням обмежень за бюджетом, місткістю та вимогами до рівня сервісу.

6. ВИСНОВКИ

У роботі запропоновано методику інтелектуального агентного управління бізнес-процесами на прикладі процесу «закупівля матеріалів під замовлення». Розроблено агентний модуль, у якому поєднано багатоагентну постановку задачі з евристично-адаптивними алгоритмами оптимізації та координації, що наближує BPMS до режиму автономного, децентралізованого прийняття рішень.

Інтелектуальні агенти в межах запропонованої архітектури забезпечують адаптивність до змін попиту та умов постачання, автономність у виборі стратегій закупівель і підвищену гнучкість порівняно зі статичними бізнес-правилами. Модифікований агентний модуль, що включає Бджолиний алгоритм для планування закупівель і Adaptive CNP для вибору постачальників, дає змогу скоротити тривалість процесу, зменшити кількість простоїв, підвищити рівень дотримання SLA та покращити ефективність використання бюджету.

Експериментальні дослідження на імітаційній моделі показали перевагу модифікованого агентного модуля порівняно зі стандартною конфігурацією, яка емулює традиційну BPMS без адаптивних механізмів. Отримані результати свідчать про перспективність інтеграції інтелектуальних агентів у корпоративні системи управління процесами для підвищення продуктивності та економічної ефективності.

Подальші дослідження передбачають розвиток когнітивних агентів із розширеними можливостями аналізу даних, інтеграцію агентного модуля з інструментами process mining для автоматизованого виявлення вузьких місць, а також апробацію запропонованого підходу в реальних BPMS-середовищах на основі промислових сценаріїв.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Wooldridge M. An Introduction to MultiAgent Systems. – 2nd ed. – Chichester: John Wiley & Sons, 2009. – 488 с.
2. van der Aalst W. M. P. Process Mining: Data Science in Action. – 2nd ed. – Berlin; Heidelberg: Springer, 2016. – 467 с.
3. Russell S. J., Norvig P. Artificial Intelligence: A Modern Approach. – 4th ed. – Hoboken: Pearson, 2021. – 1136 с.
4. Leitão P., Karnouskos S. (eds.). Industrial Agents: Emerging Applications of Software Agents in Industry. – Amsterdam: Elsevier, 2015. – 476 с.
5. IBM Business Automation Insights: Documentation [Електронний ресурс]. – IBM, 2023. – Режим доступу: <https://www.ibm.com/docs/en/cloud-paks/cp-biz-automation>. – Дата звернення: 20.11.2025.

МЕТОДИ ПІДТРИМКИ ДИНАМІЧНОГО КЕРУВАННЯ ЗАВАНТАЖЕНОСТІ ГОТЕЛЬНОГО КОМПЛЕКСУ

Мухін В.Є., Поліщук А.Б.¹

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ polishchuk.artem@lil.kpi.ua

У роботі розглянуто методи та інструменти динамічного керування завантаженістю готельного комплексу на основі аналітичних і цифрових технологій. Метою дослідження є розроблення комплексного підходу до оптимізації використання номерного фонду та підвищення ефективності управлінських рішень у реальному часі. Використано методи математичного моделювання, машинного навчання та аналізу даних. Запропонована система дозволяє прогнозувати попит, автоматично коригувати ціни та визначати пріоритети бронювання. Наукова новизна полягає у поєднанні інтелектуального аналізу даних із динамічними алгоритмами адаптивного керування. Практична значимість полягає у підвищенні доходності, зменшенні ризику простоїв та забезпеченні стабільного рівня сервісу.

Ключові слова: динамічне керування, готельний комплекс, прогнозування попиту, оптимізація завантаженості, системи підтримки прийняття рішень.

1. ВСТУП

Актуальність проблеми нерівномірного завантаження готельних ресурсів зумовлена її прямим впливом на економічну ефективність бізнесу. Коливання попиту призводять до ризику простоїв або втраченої вигоди, що негативно позначається на загальній доходності.

На ефективність функціонування готелів суттєво впливають неконтрольовані зовнішні чинники. Ключовими серед них є сезонність, проведення масових подій (конференцій, фестивалів) та загальні економічні чинники, які постійно змінюють ринкову ситуацію.

Це формує гостру необхідність у створенні гнучких систем керування, що здатні адаптуватися до змін попиту в реальному часі. Існуючі підходи часто не мають необхідної адаптивності та інтеграції з реальними потоками даних, що знижує якість управлінських рішень.

Виходячи з цього, метою дослідження є розробка методів і засобів підтримки динамічного керування завантаженістю готельного комплексу. Це дозволить оптимізувати використання номерного фонду та забезпечити стабільний рівень сервісу.

2. АНАЛІЗ ІСНУЮЧИХ РІШЕНЬ

Технологічна архітектура готелю зазвичай базується на трьох основних компонентах.

Property Management System (PMS) – координує рецепцію, бронювання і фонд номерів.

Revenue Management System (RMS) – автоматично оптимізує тарифи й завантаженість на основі аналізу попиту, часто з використанням Open Pricing.

Customer Relationship Management (CRM) – використовується для персоналізації, маркетингу та динамічного ціноутворення за репутаційними критеріями.

Існує два основні підходи прогнозування попиту. Статистичні моделі (ARIMA, експоненційне згладжування) ефективні для сезонних патернів, але слабо реагують на аномалії. Інтелектуальні моделі (машинне навчання), як-от нейронні мережі LSTM та дерева рішень (Random Forest), краще обробляють складні нелінійні залежності та безліч факторів. На практиці найефективнішим є гібридний підхід.

Можна виділити наступні прогалини у поточних підходах: відсутність адаптивності (використання статичних тарифів), фрагментованість даних між PMS, RMS та CRM, а також втручання персоналу, що знижує ефективність автоматизації.

3. МЕТОДИ ДИНАМІЧНОГО КЕРУВАННЯ

3.1. Математичне моделювання процесів завантаженості

Для динамічного керування необхідно моделювати попит на номери. Рівень завантаження залежить від множини факторів: ціни, сезонності, днів тижня, локальних подій та онлайн-рейтингу. Модель попиту (D) можна представити як функцію від цих факторів: $D = f(P, X_1, X_2, \dots, X_n)$.

Як приклад лінійної моделі, що враховує ці фактори, можна навести регресійне рівняння: $D_t = \beta_0 + \beta_1 P_t + \beta_2 Season_t + \beta_3 Event_t + \beta_4 Rating_t + \epsilon_t$.

Формалізація задачі оптимізації завантаження найчастіше зводиться до задачі максимізації загального доходу (R). Мета – знайти оптимальні ціни, що максимізують сумарний дохід $\max(P_t) R = \sum_t [P_t * D_t(P_t)]$. Оскільки функція попиту $D(P)$ нелінійна, ця задача є нелінійною. При врахуванні обмежень (загальна місткість готелю, мінімальні/максимальні ціни) задача формулюється як задача нелінійного програмування.

У повній постановці, цільова функція (максимізація доходу по всіх сегментах і днях) та обмеження виглядають так: $\max(P_t, s) \sum_t \sum_s [P_t, s * D_t, s(P_t, s)]$.

При обмеженнях $\sum_s [D_t, s(P_t, s)] \leq C_t$ (сумарний попит по всіх сегментах не повинен перевищувати загальну місткість готелю (C_t) на дату t) та $P_t, s \geq P_s(min)$ та $P_t, s \leq P_s(max)$ (ціна для кожного сегменту (s) на дату (t) повинна знаходитись у встановлених бізнес-межах).

Через стохастичний характер попиту, для її розв'язання на практиці застосовують чисельні методи, симуляції або евристики.

3.2. Алгоритми прогнозування попиту

У системах динамічного керування активно використовуються алгоритми машинного навчання (ML). LSTM (Long Short-Term Memory) – це нейронна мережа, що добре підходить для аналізу часових рядів (як-от попит), оскільки враховує довгострокові залежності, наприклад, річну сезонність. Random Forest (RF) – це ансамблевий метод, що ефективно працює з великою кількістю різномірних ознак (події, ціни конкурентів, день тижня), але вимагає явного створення часових ознак.

Порівняння точності прогнозування на історичних даних показує, що вибір моделі залежить від задачі. LSTM може перевершувати класичні методи у складних багатокрокових прогнозах. Однак на короткострокових горизонтах або при стабільних паттернах простіші статистичні моделі можуть бути точнішими. Random Forest часто демонструє високу ефективність у прогнозуванні скасувань. Тому в сучасних RMS часто використовують гібридні підходи. Власне, гібридний підхід і обрано для подальшої реалізації.

3.3. Динамічне ціноутворення

Динамічне ціноутворення є механізмом виконання рішень, отриманих на етапах прогнозування та оптимізації.

Механізм адаптивного коригування тарифів може бути простим (правила на основі поточного рівня заповненості) або просунутим, як Open Pricing, що дозволяє встановлювати ціни для кожного сегменту незалежно і динамічно.

Ключовою є інтеграція моделі прогнозування з модулем оптимізації доходу (Revenue Management). Це працює як замкнений цикл:

1. RMS отримує дані з PMS (поточні бронювання).
2. Модель прогнозування оцінює майбутній попит.
3. Модуль оптимізації розраховує оптимальну ціну для максимізації доходу.
4. Нові ціни автоматично публікуються через API в усі канали продажів.
5. Система збирає зворотний зв'язок (як змінився темп бронювань) і самонавчається.

Динамічне ціноутворення є виконавчою фазою інтегрованого процесу Revenue Management.

3.4. Мотивація застосування гібридного підходу

Як зазначено в п. 3.2, окремі алгоритми машинного навчання мають вузьку спеціалізацію. LSTM ефективно обробляє довгострокові часові залежності (наприклад, сезонність), тоді як Random Forest (RF) краще працює з великою кількістю різнорідних ознак та прогнозуванням скасувань. Оскільки задача моделювання попиту є комплексною, стохастичною і вимагає врахування обох типів даних, вибір гібридної моделі є обґрунтованим. Мотивація полягає у поєднанні сильних сторін обох методів: RF - для обробки та відбору ознак, LSTM – для моделювання часових патернів. Такий підхід дозволяє досягти вищої точності прогнозу, що є критично необхідним для ефективної роботи гнучких механізмів ціноутворення, як-от Open Pricing.

4. ПРОГРАМНІ ЗАСОБИ ТА ІНФОРМАЦІЙНА ІНФРАСТРУКТУРА

Архітектура системи підтримки динамічного керування завантаженістю будується за модульним принципом. Ключовим елементом є послідовна взаємодія цих модулів:

Модуль збору даних акумулює історичну та поточну інформацію.

Модуль аналітики обробляє ці дані.

Модуль прогнозування (з використанням машинного навчання) оцінює майбутній попит.

Модуль автоматичного прийняття рішень використовує ці прогнози для коригування цін.

Для забезпечення стабільної роботи та інтеграції використовуються хмарні технології.

Вони надають необхідні ресурси для складних обчислень. Крім того, API застосовуються для прямої інтеграції системи з провідними онлайн-платформами бронювання, такими як Booking, Expedia та Airbnb. Для зручності персоналу, усі ключові показники, прогнози та результати роботи системи надаються адміністратору готелю через спеціальні засоби візуалізації (наприклад, дашборди).

5. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ ТА ПРОГНОЗУВАННЯ ЕФЕКТИВНОСТІ

5.1. Аналіз ефективності динамічного керування

Аналіз ринку доводить значні переваги автоматизованого динамічного керування порівняно зі статичним (фіксовані тарифи). Ключові показники покращення при переході на ДУ включають:

Загальний дохід: Середнє зростання становить від 10% до 20% [1]. Дослідження 567 готелів (RoomPriceGenie) зафіксувало середнє зростання доходу на +19% [2].

RevPAR (Дохід на доступний номер): Галузеві звіти підтверджують зростання RevPAR на 8–15% вже протягом першого року впровадження [3].

Заповнюваність та ADR: Ефект досягається комбіновано. Те саме дослідження показало зростання заповнюваності на +14% та середнього тарифу (ADR) на +4% [2].

5.2. Прогноз покращення точності (MAPE) за допомогою гібридного підходу

Запропоновано гібридний підхід, що поєднує моделі Random Forest (RF) та LSTM, гнучку стратегію Open Pricing та глибоку інтеграцію систем. Цей підхід дозволяє зробити прогноз щодо покращення ключового критерію - точності прогнозування попиту (MAPE).

Базовий рівень (Порівняння): Традиційні статистичні моделі (наприклад, ARIMA, згадані в п. 2), які використовуються як базовий рівень, в академічних дослідженнях демонструють середню абсолютну відсоткову помилку (MAPE) на рівні приблизно 20.7% [4].

Проблема ізольованих ML-моделей: Дослідження показують, що просте застосування глибокого навчання (LSTM) не гарантує кращого результату і може призводити до перенавчання (overfitting), показуючи гіршу точність, ніж простіші моделі [5].

В свою чергу, у запропонованому гібридному підході ефективність досягається через поєднання:

Random Forest використовується для відбору найважливіших ознак (події, ціни конкурентів, відгуки). Доведено, що методи на основі RF (як Boruta) здатні критично знижувати "шум", покращуючи MAPE на 54.11% порівняно з базовими методами [6].

LSTM застосовується для моделювання складних нелінійних залежностей у відібраних даних, наприклад, аналізуючи настрої з відгуків [7] або динаміку пошукових запитів [8].

Прогноз: Ансамблеві моделі, що поєднують різні підходи, демонструють найвищу точність, досягаючи MAPE 14.2% [9].

Таким чином, прогнозується, що запропонований гібридний підхід дозволить знизити середню помилку прогнозування (MAPE) з 21% (базовий рівень SARIMA [4]) до 14%. Ця підвищена точність, в свою чергу, максимізує ефект від стратегії Open Pricing [10] та інтеграції систем.

6. ВИСНОВКИ

У роботі проаналізовано методи динамічного керування завантаженістю. Аналіз підтвердив, що перехід на автоматизоване ДУ підвищує загальний дохід готелів на 10–20% та RevPAR на 8–15%.

Практичним результатом є аналітичне обґрунтування гібридного підходу (RF + LSTM), який є ефективнішим за ізольовані моделі, такі як традиційна SARIMA (MAPE приблизно 20.7%) або чисті моделі глибокого навчання, схильні до перенавчання.

Аналітично спрогнозовано, що запропонований гібридний ансамбль, який використовує RF для відбору ознак та LSTM для аналізу складних патернів, здатний знизити середню помилку прогнозування (MAPE) до 14.2%.

Обмеженням тез є те, що цей показник є аналітичним прогнозом, заснованим на синтезі вторинних джерел. Подальші дослідження будуть спрямовані на практичну реалізацію та валідацію цієї гібридної моделі для експериментального підтвердження точності та її впливу на максимізацію доходу від стратегій Open Pricing.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Dynamic Pricing with AI: How Hotels Optimise Revenue in Real-Time. *Mondaylabs.ai*. URL: <https://www.mondaylabs.ai/blog/dynamic-pricing-with-ai-how-hotels-optimise-revenue-in-real-time> (дата звернення: 15.11.2025).

2. The Real Impact of Revenue Management: Insights from 567 Hotels Using RoomPriceGenie. *HFTP*. URL: <https://www.hftp.org/news/4126644/the-real-impact-of-revenue-management-insights-from-567-hotels-using-roompricegenie> (дата звернення: 15.11.2025).
3. Benchmarking in the Hotel Industry: A Data-Driven Path to Success. *Vynta.ai*. URL: <https://vynta.ai/blog/benchmarking-hotels/> (дата звернення: 15.11.2025).
4. Revolutionizing Revenue Management: How AI Analytics Are Transforming the Hospitality Industry in 2025. *SuperAGI*. URL: <https://superagi.com/revolutionizing-revenue-management-how-ai-analytics-are-transforming-the-hospitality-industry-in-2025/> (дата звернення: 15.11.2025).
5. Hotel daily demand forecasting for high-frequency and complex seasonality data: a case study in Thailand. *ResearchGate*. URL: https://www.researchgate.net/publication/337650959_Hotel_daily_demand_forecasting_for_high-frequency_and_complex_seasonality_data_a_case_study_in_Thailand (дата звернення: 15.11.2025).
6. Is Machine Learning a Magic Wand? Comparative Analysis of Machine Learning and Pickup Models for Hotel-Demand Forecasting. *Boston University Hospitality Review*. 29.06.2021. URL: <https://www.bu.edu/bhr/2021/06/29/is-machine-learning-a-magic-wand/> (дата звернення: 15.11.2025).
7. Comparative Analysis of Machine Learning and Deep Learning Models for Tourism Demand Forecasting with Economic Indicators. *Forecasting*. 2022. Vol. 4(3). P. 46. URL: <https://www.mdpi.com/2674-1032/4/3/46> (дата звернення: 15.11.2025).
8. Forecasting tourist arrivals at attractions: Search engine empowered methodologies. *ResearchGate*. URL: https://www.researchgate.net/publication/329062137_Forecasting_tourist_arrivals_at_attractions_Search_engine_empowered_methodologies (дата звернення: 15.11.2025).
9. Tourism and Hospitality Forecasting with Big Data. *The Hong Kong Polytechnic University*. URL: https://ira.lib.polyu.edu.hk/bitstream/10397/107074/1/Wu_Tourism_Hospitality_Forecasting.pdf (дата звернення: 15.11.2025).
10. Varshavskiy I. E., Stavinova E., Chunaev P. Forecasting railway ticket demand with search query open data. *ResearchGate*. Листопад 2022. URL: https://www.researchgate.net/publication/301509073_Tourism_forecasting_by_search_engine_data_with_noise-processing (дата звернення: 15.11.2025).
11. Беляєв, В. О. Системи керування доходами в готельному бізнесі. *Економіка та держава*, 2022. №3. С. 45–51.
12. Козак, Л. В. Інформаційні технології в індустрії гостинності. Київ: КНТЕУ, 2021.
13. Ivanov, S., & Webster, C. (2019). *Revenue Management for Hospitality and Tourism*. Goodfellow Publishers.
14. Talluri, K., & Van Ryzin, G. (2004). *The Theory and Practice of Revenue Management*. Springer Science & Business Media.
15. Li, J., et al. (2023). Dynamic Pricing in Hotel Management Using AI. *Tourism Economics*, 29(5), 1058–1074.

СИСТЕМА ГЕНЕРАЦІЇ ШАБЛОНІВ ДОКУМЕНТІВ З LLM КОМПОНЕНТАМИ

Нікітін В.С.¹, Гіоргізова-Гай В.Ш.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ vladyslavnikitin6807@gmail.com, ² high.victoria@lil.kpi.ua [0000-0001-6224-3532]

Мета дослідження – автоматизувати побудову шаблонів документів за зображеннями форм. Застосовано двоетапний конвеєр LLM із валідацією JSON-шаблонів та метриками на основі індексу Жаккара. Проведено порівняльне тестування шести моделей на дев'яти документах з трьох предметних областей. Наукова новизна полягає в розробці методики оцінки якості генерованих шаблонів. Отримані результати є новими для україномовних документів і придатні до інтеграції в системи електронного документообігу. Практична значимість – скорочення часу створення електронних форм з годин до хвилин при цифровізації великих обсягів типових документів.

Ключові слова: великі мовні моделі, електронні форми, шаблони документів, автоматизація документообігу.

1. ВСТУП

Стандартизовані підходи до структурування документів є критично важливими для сучасного електронного документообігу. Вони забезпечують однозначність трактування інформації, сумісність між системами та спрощують інтеграційні процеси. Це особливо актуально для автоматизованої перевірки, аудиту та довгострокового архівування, де потрібні стабільні формати і мінімізація людського фактору.

Про актуальність даної теми свідчить активне впровадження та розширення функціоналу електронних рапортів у мобільному застосунку Армія+, що будує документообіг за допомогою покрокових форм з чіткими правилами валідації, на противагу неструктурованим файлам у форматі *.docx [1]. Також варто зазначити, що дослідження останнього десятиліття свідчать про те, що 80% даних, генерованих різноманітними організаціями є неструктурованими [2, 3], що може значно обмежувати можливості їхнього аналізу.

Жодне з розглянутих існуючих комерційних рішень не може бути використаним для генерації великих обсягів шаблонів документів в предметній області. Виявлені обмеження існуючих підходів обґрунтовують необхідність розробки нової архітектури системи генерації шаблонів, яка поєднуватиме переваги використання LLM для розпізнавання структури документів з чітким розділенням рівнів абстракції (дані, форма, документ), надійними механізмами валідації та кількісними метриками оцінки якості результатів.

2. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ

Розробник британського державного порталу GOV.UK Тім Пол (Tim Paul) в серії з двох публікацій описав підхід до автоматизації створення веб-форм на основі PDF-форми за допомогою великих мовних моделей [4, 5].

Інструмент приймає на вхід PDF-документи, зображення у форматах JPEG або PNG або рукописні ескізи форм, а на виході система генерує структурований JSON, що містить повний опис форми: назву оригінального файлу, структуру сторінок із зазначенням кількості питань на кожній, та детальний опис кожного питання. Для кожного питання зберігається його текст, підказки для користувача, тип відповіді (числа, дати, текст тощо) та можливі варіанти вибору. Процес генерації є максимально простим і містить один виклик до мовної моделі з використанням механізму інструментів (tools) для отримання структурованої відповіді.

З недоліків даного рішення можна виділити:

- слабе відокремлення структури даних від представлення: результатом роботи системи є масив UI-елементів без унікальних ідентифікаторів, що може ускладнити подальшу обробку відповіді кінцевого користувача в контексті більшої системи;
- відсутність чітких критеріїв оцінки згенерованої форми. Автор демонструє приклад некоректної роботи системи, але не пропонує методів для виявлення подібних проблем.

Аналіз функціоналу пропрієтарних SaaS рішень Fillout [6] і Jotform [7] показав, що обидва рішення є аналогічними і дозволяють завантажити PDF-документ для створення редагованої інтерактивної веб-форми на його основі. Рішення було протестовано шляхом завантаження шаблонних PDF з відкритих джерел. Тестування показало, що дані комерційні рішення можна застосовувати лише для файлів, розмічених за допомогою AcroForms [8], що свідчить про неможливість їх використання для неструктурованих і слабо структурованих документів (зображень, нерозмічених PDF, DOCX тощо).

3. ПОСТАНОВКА ЗАДАЧІ

Метою дослідження є розробка системи автоматичної генерації шаблонів документів на основі зображень існуючих форм з використанням великих мовних моделей. Система проектується як інструментарій для розробників програмного забезпечення, а не як самостійний продукт для кінцевих користувачів.

Використання великих мовних моделей обумовлене необхідністю швидкої обробки значної кількості документів у конкретній предметній області з мінімальною участю людини. Традиційні підходи до створення електронних форм вимагають значних часових витрат на ручне проектування кожного шаблону, що робить неефективною цифровізацію великих архівів документів. LLM дозволяють автоматизувати цей процес, скорочуючи час створення одного шаблону з годин до хвилин і забезпечуючи можливість паралельної обробки десятків документів.

Кінцевим результатом роботи системи є два взаємопов'язані компоненти. Перший компонент – інтерактивна веб-форма, що надає структурований інтерфейс для введення даних користувачем з вбудованими механізмами валідації та підказками. Другий компонент – шаблонізований документ, що перетворює введені структуровані дані у формат, зрозумілий для людського сприйняття, зазвичай PDF для друку або архівування. Архітектура системи забезпечує достатню гнучкість для адаптації форми під різні платформи представлення, включаючи мобільні застосунки, десктопні програми або чат-боти, завдяки використанню проміжного JSON-представлення шаблону, незалежного від конкретної технології рендерингу.

Важливою особливістю системи є свідома відмова від прагнення досягти стовідсоткової візуальної відповідності згенерованого документа вхідному зображенню. Пріоритетом є коректне розпізнавання структури документа та наявність усіх необхідних полів даних з правильно визначеними типами та взаємозв'язками. Кінцевий документ матиме уніфікований зовнішній вигляд відповідно до стилістичних стандартів організації, що його

використовує, забезпечуючи консистентність документообігу та спрощуючи інтеграцію з існуючими системами.

4. СИСТЕМА ГЕНЕРАЦІЇ ШАБЛОНІВ ДОКУМЕНТІВ

4.1. Архітектура системи

На рисунку 1 представлено основні компоненти системи та потоки даних між ними.

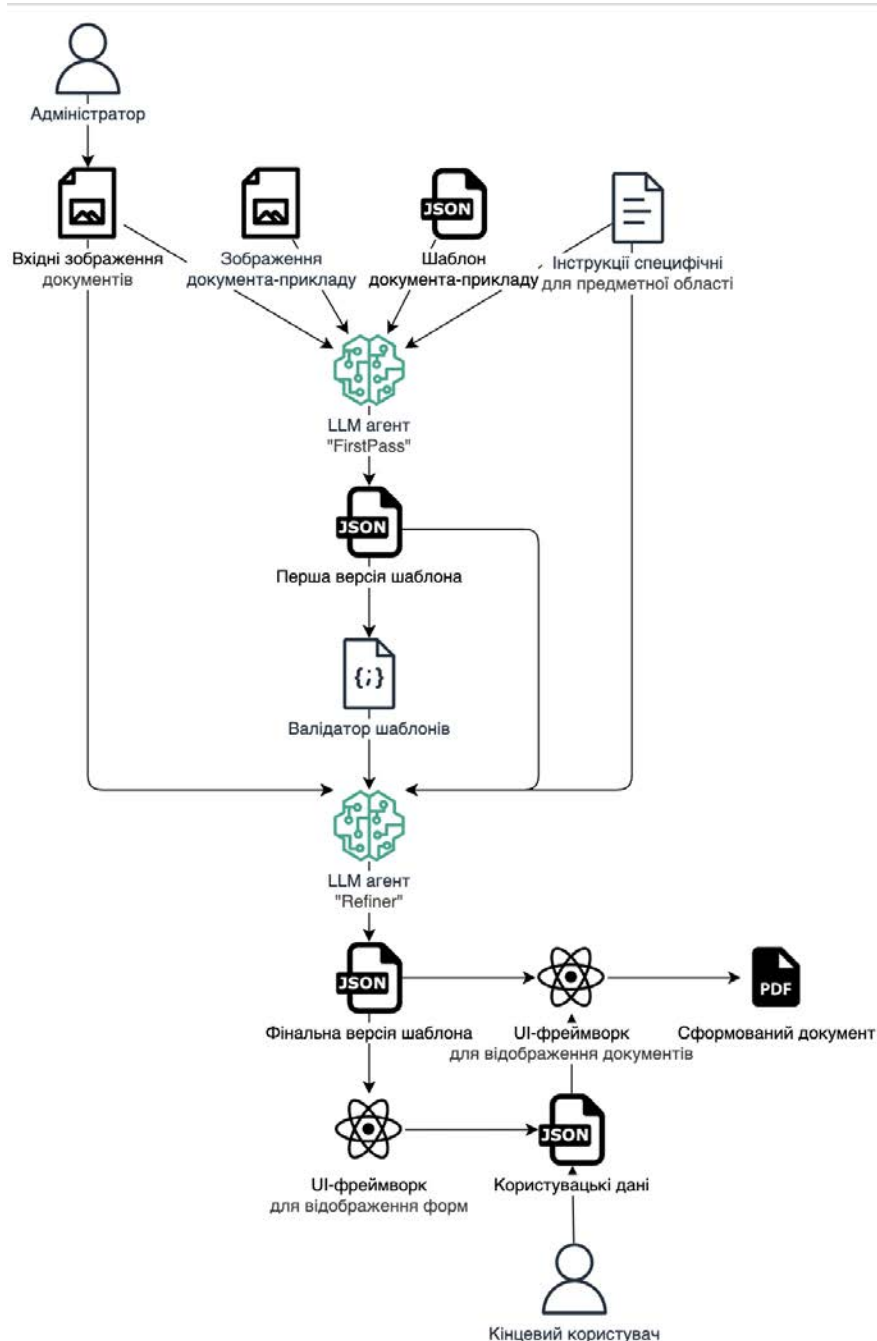


Рисунок 1. Функціональна схема системи генерації шаблонів документів

Ключовими елементами на рисунку 1 є:

1. Вхідні зображення документа – скан-копії або фото реальних форм, які потрібно оцифрувати і перетворити на шаблон.

2. Зображення документа-прикладу та його JSON-шаблон демонструють моделі типову структуру даних у певній предметній області. Такий приклад визначається розробником системи і є відмінним від вхідного документа з п.1.

3. Інструкції, специфічні для предметної області, описують правила, яких мають дотримуватись всі згенеровані шаблони, у довільній формі. Наприклад: «All generated text (labels, placeholders etc.) must be in Ukrainian».

4. FirstPass – це LLM-агент, який на основі зображень та інструкцій будує першу чернетку шаблону у форматі JSON.

5. Валідатор – це модуль, який автоматично перевіряє згенерований шаблон на структурну коректність, повноту полів і відповідність схемі системи.

6. Refiner – це другий LLM-агент, який, використовуючи результати валідації, покращує й виправляє початковий шаблон, формуючи його фінальну версію.

4.2. Визначення структури генерованих шаблонів документів

Вимоги до формату генерованих шаблонів документів включають:

- безпеку – відсутність виконуваного коду в шаблоні з метою запобігання потенційним вразливостям (ін'єкції коду тощо);

- легкість парсингу та валідації – можливість швидкого розбору шаблону за допомогою стандартних бібліотек (JSON, YAML), а також можливість декларативного опису формату для валідації шаблонів;

- зручність для генерації за допомогою LLM.

Концептуально шаблон розділено на три основні складові, кожна з яких відповідає за окремий аспект представлення інформації:

1. Конфігурація даних, що описує змінні та їхні типи.

2. Структура електронної форми, що складається з компонентів. Компоненти визначають інтерфейс взаємодії користувача з системою при заповненні даних. До них належать текстові поля введення, перемикачі, випадаючі списки, компоненти для підпису документів тощо. Кожен компонент пов'язаний з однією або декількома змінними з конфігурації даних.

3. Опис структури фінального документа через систему віджетів. Віджети відповідають за представлення даних у згенерованому PDF-документі.

4.3. Методика оцінки якості згенерованих шаблонів

Оцінка якості згенерованого шаблону виконується шляхом його порівняння з еталоном, заданим для конкретного документа заздалегідь, за декількома метриками.

Оскільки змінні в шаблоні організовані як неупорядкована структура «ключ-значення», а відносний порядок компонентів форми та віджетів документа не впливає на функціональність шаблону, розроблені метрики оцінки структурної подібності не враховують позиції елементів і фокусуються виключно на правильності розпізнавання їх типів.

Для оцінки якості розпізнавання типів використовується індекс Жаккара [9]:

$$TP_c = \min(count_{pred}(c), count_{true}(c)),$$
$$J_c = \frac{TP_c}{count_{true}(c) + count_{pred}(c) - TP_c},$$

де $count_{pred}(c)$ – кількість елементів класу c в згенерованому шаблоні, $count_{true}(c)$ – кількість елементів класу c в еталонному шаблоні.

Дана метрика спочатку обчислюється на рівні окремих класів, після чого агрегується у середнє зважене (weighted average) значення, що дозволяє отримати загальну оцінку якості окремого компонента шаблону.

$$J_{weighted} = \frac{\sum_{c \in C} J_c \cdot count_{true}(c)}{\sum_{c \in C} count_{true}(c)}.$$

Для отримання єдиної оцінки для всього шаблону використовується гармонійне середнє:

$$J_{template} = \frac{3}{\frac{1}{J_{variables}} + \frac{1}{J_{components}} + \frac{1}{J_{widgets}}}.$$

4.4. Результати тестування системи

Систему було протестовано на документах з трьох предметних областей: україномовні юридичні документи (договори, заяви), україномовні військові документи (рапорти та заяви) та англomовні документи з охорони праці на будівництві (pre-task plan). Для кожної предметної області було відібрано по три приклади документів, специфічних для конкретної галузі. Для кожного документа було розроблено еталонний шаблон, з яким в подальшому порівнювались результати генерації моделей.

Тестування виконувалось для новітніх пропріетарних великих мовних моделей (Claude Sonnet 4, Gemini 2.5 Flash, Gemini 2.5 Pro, GPT-5 High) та open-weight моделей (Llama 4 Maverick та Qwen3 VL 235B A22B Thinking).

На рисунку 2 зображено гармонійне середнє значення коефіцієнтів Жаккара для кожного опрацьованого документа.

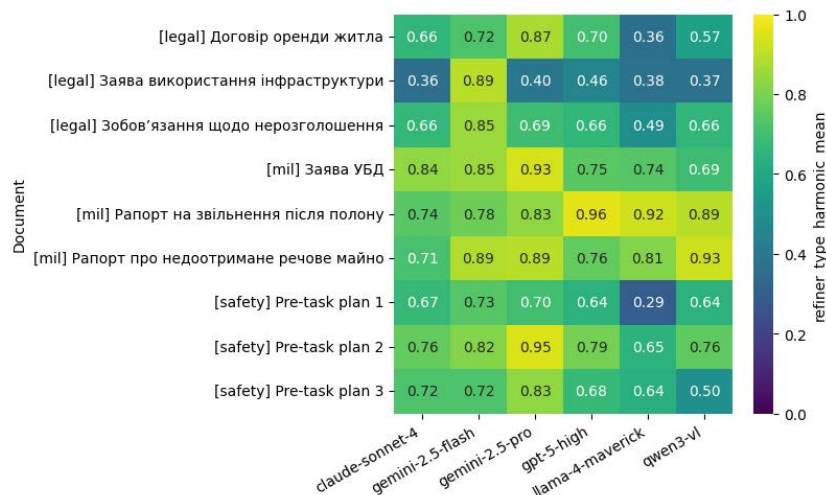


Рисунок 2. Результати тестування системи

На рисунку 3 зображено кількість помилок валідації в шаблоні на першому етапі (First Pass) і кількість помилок валідації у фінальному генерованому шаблоні (Refiner).

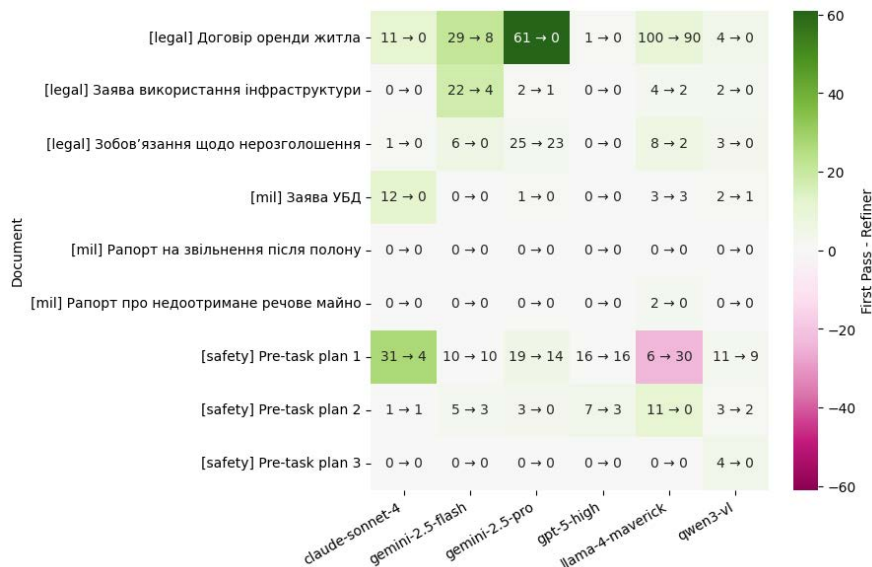


Рисунок 3. Кількість помилок валідації в шаблоні

5. ВИСНОВКИ

Проведене дослідження підтверджує доцільність застосування великих мовних моделей для задачі автоматичної генерації шаблонів документів, при цьому порівняльний аналіз виявив суттєву перевагу пропріетарних рішень над open-weight альтернативами, що робить використання останніх недоцільним. На основі отриманих результатів для впровадження системи рекомендується використання моделі Gemini 2.5 Pro, яка продемонструвала найбільш збалансовані показники якості.

Критично важливе значення для успішності системи має наявність чітко визначених метрик оцінки якості, використання декларативних правил валідації на основі схем, а також надання моделі прикладів у форматі few-shot або one-shot learning безпосередньо в контекстному вікні, що дозволяє сформулювати адекватне уявлення про очікуваний формат виводу. Застосований двоетапний підхід з агентами «FirstPass» та «Refiner» показав свою ефективність, підтвердивши здатність мовних моделей до самокорекції при наданні детальної інформації про виявлені помилки, хоча у деяких випадках навіть повторний запит не призводить до повного виправлення помилок, що вказує на необхідність розробки більш складних механізмів контролю якості.

Перспективними напрямками подальших досліджень є застосування методів автоматичної оптимізації промптів [10] замість їх ручного конструювання, що може суттєво підвищити якість генерації шаблонів, а також експерименти з донавчанням open-weight моделей.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Рапорти на відпустку, виплати, лікування, звільнення | Армія+ : веб-сайт. URL: <https://aplus.mod.gov.ua/raporty> (дата звернення: 15.11.2025).
2. Основи науки про дані [Електронний ресурс]: навч. посіб. для студ. спеціальності 121 «Інженерія програмного забезпечення» / М. А. Новотарський; КПІ ім. Ігоря Сікорського. – Електронні текстові дані (1 файл: 5,5 Мбайт). – Київ : КПІ ім. Ігоря Сікорського, 2022. – 294 с. – Режим доступу: <https://ela.kpi.ua/server/api/core/bitstreams/04d92b80-e82a-42cb-bcee-0387cb8a8cc6/content>

3. D. Baviskar, S. Ahirrao, V. Potdar, and K. Kotecha, "Efficient Automated Processing of the Unstructured Documents Using Artificial Intelligence: A Systematic Literature Review and Future Directions," *IEEE Access*, vol. 9, pp. 72894-72936, 2021, doi: 10.1109/ACCESS.2021.3072900.
4. Using AI to generate web forms from PDFs : веб-сайт. URL: <https://www.timpaul.co.uk/posts/using-ai-to-generate-web-forms-from-pdfs/> (дата звернення: 10.10.2025).
5. Could AI convert all the PDF forms on GOV.UK into web forms? : веб-сайт. URL: <https://www.timpaul.co.uk/posts/could-ai-convert-all-the-pdf-forms-on-gov-uk-into-web-forms/> (дата звернення: 10.10.2025).
6. PDF to Form : веб-сайт. URL: <https://www.fillout.com/pdf-to-form> (дата звернення: 12.10.2025).
7. Jotform Smart PDF Forms: Convert PDF to HTML Web Forms : веб-сайт. URL: <https://www.jotform.com/products/smart-pdf-forms/> (дата звернення: 12.10.2025)
8. ISO 32000-1:2008. Document management — Portable document format — Part 1: PDF 1.7 [Керування документами — Портативний формат документів — Ч. 1: PDF 1.7].
9. Методи аналізу даних та машинного навчання в інформаційно-управляючих системах [Електронний ресурс] : курс лекцій. Ч. 2 : навч. посіб. для здобувачів ступеня бакалавра за спец. 126 «Інформаційні системи та технології» / КПІ ім. Ігоря Сікорського ; уклад.: В. В. Онищенко, О. В. Гавриленко, М. Ю. Мягкий. — Київ : КПІ ім. Ігоря Сікорського, 2025. — 318 с. — Режим доступу: <https://ela.kpi.ua/handle/123456789/74212> (дата звернення: 20.10.2025)
10. Ramnath K., Zhou K., Guan S. та ін. A systematic survey of automatic prompt optimization techniques [Електронний ресурс]. — 2025. — arXiv:2502.16923. — Режим доступу: <https://arxiv.org/abs/2502.16923>. — doi: 10.48550/arXiv.2502.16923.

ПРОГНОЗНЕ МОДЕЛЮВАННЯ ПОВЕДІНКИ АГЕНТІВ ТА АДАПТИВНЕ УПРАВЛІННЯ ЕНЕРГОРЕСУРСАМИ ДЛЯ ПІДВИЩЕННЯ БЕЗПЕКИ ТА АВТОНОМНОСТІ РОЮ ДРОНІВ

Ніколайчук О.Г.¹, Харченко К.В.²

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ nikolaychuk.olexandr@lil.kpi.ua, ² kharchenko.kostiantyn@lil.kpi.ua

Досліджено модуль прогнозного моделювання поведінки агентів та підсистему адаптивного управління енергоресурсами для рою безпілотних літальних апаратів. Реалізацію виконано у вигляді окремих ROS2-вузлів, що взаємодіють із симуляційним середовищем Gazebo Harmonic та автопілотом ArduPilot. Запропонований підхід забезпечує підвищення автономності, енергоефективності та безпеки польоту дронів у мінливих умовах середовища. Новизною є інтеграція прогнозування майбутніх траєкторій та динамічного енергокерування у сервісну архітектуру рою, що дозволяє адаптувати швидкість, висоту та інші параметри польоту у реальному часі.

Ключові слова: енергоменеджмент, роєва система, прогнозування траєкторій, ROS2, ArduPilot, Gazebo, adaptive control.

1. ВСТУП

Проблема забезпечення автономності та безпеки рою БПЛА стає ключовою у задачах моніторингу, картографування та пошуково-рятувальних операцій, координація та автономне функціонування у динамічному середовищі потребує вирішення задач енергетичної ефективності, уникнення колізій та адаптації до невизначеності, що є досить актуальною науково-прикладною задачею. Дрони працюють з обмеженим запасом енергії, взаємними перешкодами та змінними траєкторіями сусідніх агентів. У таких умовах класичні статичні алгоритми планування траєкторій або фіксовані енергетичні моделі не забезпечують оптимальної поведінки.

У рамках даної роботи мною запропоновано сервіс-орієнтовану підсистему Adaptive Energy-Aware Flight Management System (AEAFMS), яка забезпечує прогнозне моделювання поведінки інших агентів та адаптивне управління енергоресурсами під час виконання польотного завдання. Запропонована система виконує зчитування запланованого маршруту, аналіз його енергетичної складності та динамічно змінює параметри польоту дрона — зокрема горизонтальну швидкість та інші енергетично значущі характеристики — з метою оптимізації та максимізації тривалості автономної роботи.

2. ТЕОРЕТИЧНА СКЛАДОВА АЛГОРИТМІВ

Для оцінки можливих конфліктів використовується лінійне передбачення руху сусіднього дрона:

$$\hat{p}(t + \Delta T) = p(t) + v(t) \cdot \Delta T,$$

де $p(t)$ – поточна позиція агента(x, y, z),

$v(t)$ – його швидкість,
 ΔT – горизонт прогнозу.

Відносна відстань між двома агентами:

$$D(t + \Delta T) = \| \hat{p}_i(t + \Delta T) - \hat{p}_j(t + \Delta T) \|.$$

Потенційний конфлікт виникає при:

$$D(t + \Delta T) < R_{safe}.$$

Отримана інформація передається у модуль адаптивного керування швидкістю. Енергоспоживання дрона на сегменті маршруту визначається за формулою:

$$E = P_{total} \cdot t, \quad t = \frac{d}{v},$$

де d – довжина сегменту, v – горизонтальна швидкість.

Сумарна потужність складається з трьох компонентів:

$$P_{total} = P_{hover} + k_{move}v^2 + k_{climb}|v_z|.$$

Загальна енергія маршруту – сумою по всіх сегментах:

$$E_{route} = \sum_{i=1}^N E_i.$$

Маршрут вважається допустимим, якщо:

$$E_{route} \leq E_{usable} = E_{battery}(1 - r),$$

де r – резерв (у роботі 20%), це було створено як певна “подушка” для непередбачуваних ситуацій.

Оскільки збільшення швидкості одночасно зменшує час польоту, але збільшує квадратичну складову енергоспоживання, тому функція $E(v)$ не є монотонною.

Завдання цього алгоритму:

$$v^* = \max\{v: E(v) \leq E_{usable}\}.$$

Тобто система визначає найвищу швидкість, при якій енергії достатньо для проходження усього маршруту. В результаті цей теоретичний підхід поєднує прогнозування майбутньої поведінки інших агентів та енергетичну модель польоту, що дозволяє кількісно оцінювати як ризики потенційних зближень, так і енергетичну прохідність маршруту. Це дає можливість формувати адаптивний профіль польоту, у якому швидкість обирається як максимально можлива за умови дотримання енергетичного обмеження.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для реалізації поставленого завдання мною було розроблено декілька сервісів за допомогою мови програмування Python та бібліотеки ROS 2, серед яких:

RouteAnalyzerNode - нода, призначена для читання маршруту, його структурний аналіз, визначення довжин сегментів, передача структурованої інформації до енергетичної ноди.

AdaptiveEnergyManagementNode - нода, що використовує скомпоновані дані дронів для аналізу маршруту, розрахунку енергії, визначення максимально можливої швидкості, адаптація параметрів польоту.

Для експериментальної перевірки розробленої підсистеми AEAfMS було проведено моделювання польоту дрона по заданому маршруту в Python-середовищі з використанням

енергетичної моделі та алгоритму адаптивного підбору швидкості. Маршрут складався з шести основних сегментів із підйомом на 30 м та чотирма прямолінійними ділянками по 300 м.

Першочергово виконано оцінку маршруту при початковій швидкості 8.0 м/с. Результати показали, що необхідна енергія перевищує доступну з урахуванням 20% резерву батареї, дивитись рис. 1 та 2.

```
=== INITIAL ROUTE EVALUATION ===  
Initial speed: 8.0 m/s  
Required energy: 53.96 Wh  
Usable battery energy (with 20% reserve): 45.60 Wh  
Route feasible: NO
```

Рисунок 1. Результати розрахунку енергії з аналізу маршруту та заданих параметрів дрона.

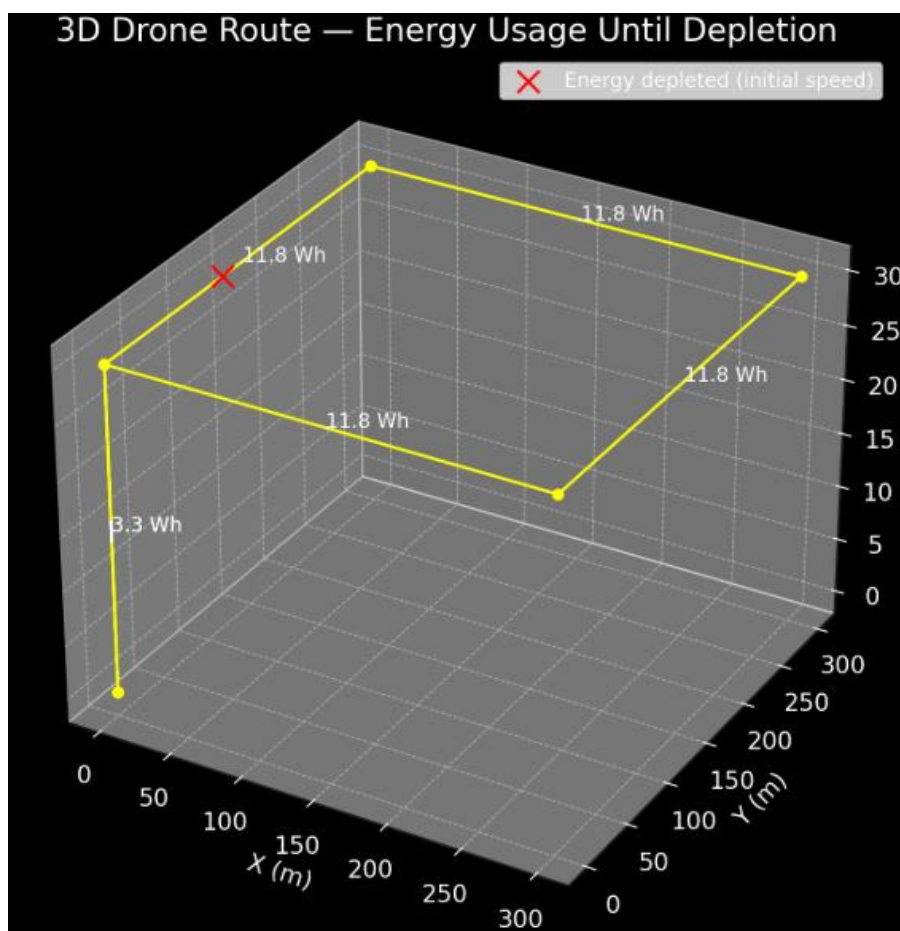


Рисунок 2. Розрахунок витрат енергії на кожен сегмент. Показано точку (червоний) повної витрати енергії з урахуванням резервної.

Тобто за такого режиму дрон не здатний завершити місію без повного розряду батареї, що підтверджує потребу в адаптивному регулюванні параметрів польоту. Після активації алгоритму адаптивного енергоменеджменту було виконано пошук максимально можливої швидкості, яка забезпечує прохідність маршруту. Алгоритм перебирає швидкості в діапазоні мінімально та максимально можливих значень даного дрона [1.5; 12.0] м/с та знаходить найбільше значення, що задовольняє умову енергодостатності. Отриманий результат дивитись на рис. 3.

```

>>> Energy insufficient – starting adaptive speed optimization...
>>> Optimal speed found: 5.70 m/s

=== RE-CALCULATED RESULTS ===
Adjusted speed: 5.70 m/s
New required energy: 45.40 Wh
Route feasible: YES

Segment details:
Segment 0: E = 3.333 Wh, t = 60.0 s, P = 200.0 W
Segment 1: E = 9.683 Wh, t = 52.6 s, P = 662.4 W
Segment 2: E = 9.683 Wh, t = 52.6 s, P = 662.4 W
Segment 3: E = 9.683 Wh, t = 52.6 s, P = 662.4 W
Segment 4: E = 9.683 Wh, t = 52.6 s, P = 662.4 W
Segment 5: E = 3.333 Wh, t = 60.0 s, P = 200.0 W

```

Рисунок 3. Результати розрахунку енергії після зміни параметрів дрона.

Таким чином, в даній локальній задачі зниження швидкості з 8.0 м/с до 5.70 м/с дозволило дрону зменшити енергоспоживання до значення, що гарантує безпечне проходження маршруту. Також на рис. 4 зображено графік, який детальніше показує, що було обрано саме найбільшу можливу швидкість для виконання поставленої задачі.

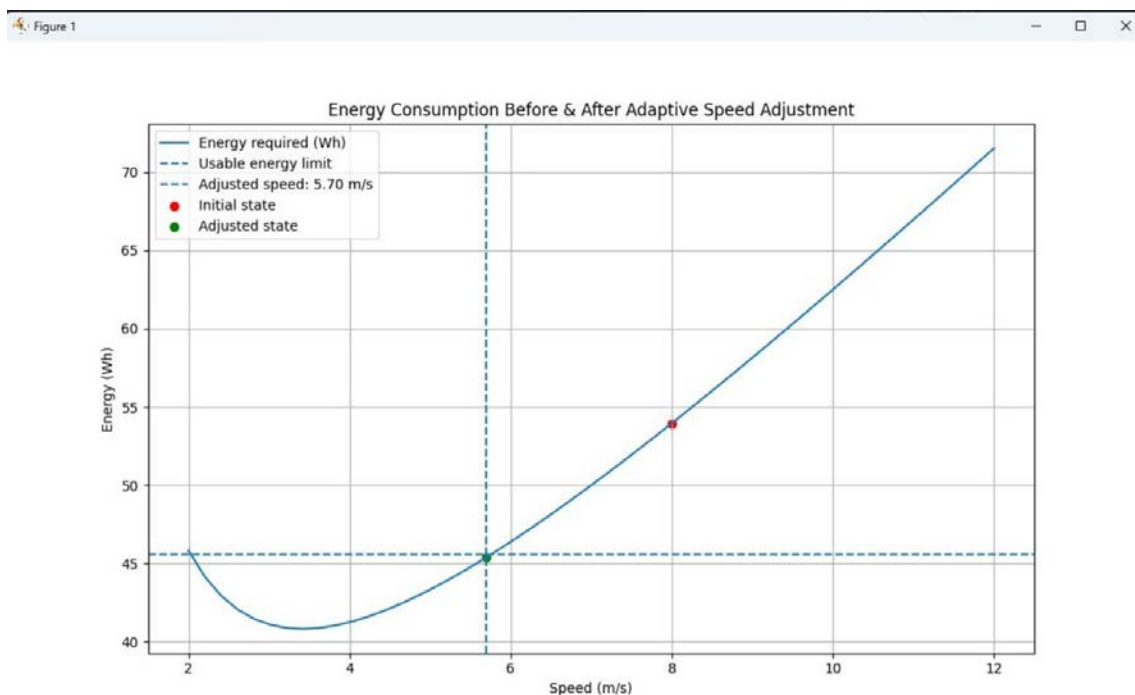


Рисунок 4. Залежність витрат енергії від швидкості польоту. Показано початкову (червоний) та скориговану алгоритмом (зелений) швидкості польоту дрону.

4. ПОДАЛЬШІ ДОСЛІДЖЕННЯ

Першочерговим напрямом є реалізація PredictionNode, який робить оцінку ризику зближення та формування безпечного горизонту руху. Це дозволяє переходити від статичної оцінки конфліктів до динамічного попередження. Другим напрямком подальших робіт є розгортання та вдосконалення ACO_DecisionNode, який інтегруватиме результати прогнозування та енергетичні обмеження в процес планування руху. І останній етап це інтеграція всієї системи з повноцінним симуляційним середовищем Gazebo Harmonic та

автопілотом ArduPilot, що дасть змогу протестувати алгоритми на реалістичних моделях дронів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Comprehensive Energy Consumption Model for Unmanned Aerial Vehicles, Based on Empirical Studies of Battery Performance URL: https://www.researchgate.net/publication/328188643_Comprehensive_Energy_Consumption_Model_for_Unmanned_Aerial_Vehicles_Based_on_Empirical_Studies_of_Battery_Performance .
2. Ardupilot Documentation URL: <https://ardupilot.org/dev/index.html> .
3. AI & Drone Technology in Industry. URL: https://discoverypoint.ieee.org/ai-drone-technology/?gad_source=1&gad_campaignid=17545506682&gbraid=0AAAAAoP_A4LFzGOdTG_T26SVEg5JV_-1o0&gclid=CjwKCAiAlfvIBhA6EiwAcErpyVKJFQSHmvinS6kAHCljzL2Yqb_dMdZ8cBRLICQHFhx2DDTilBWYMBBoCZZ8QAvD_BwE .
4. Austin R. “Unmanned Aircraft Systems: UAVS Design, Development and Deployment.” Wiley, 2010.
5. Energy-Efficient Unmanned Aerial Vehicle (UAV) Surveillance Utilizing Artificial Intelligence (AI) URL: <https://onlinelibrary.wiley.com/doi/10.1155/2021/8615367> .
6. Energy-Efficient UAV Trajectory Design and Velocity Control for Visual Coverage of Terrestrial Regions URL: <https://www.mdpi.com/2504-446X/9/5/339>

ЗАСТОСУВАННЯ МЕТОДУ FAIR ДЛЯ КІЛЬКІСНОГО АНАЛІЗУ РИЗИКІВ У СИСТЕМАХ ІНФОРМАЦІЙНОЇ БЕЗПЕКИ

Пакуля С.О.¹, Мухін В.Є.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ pakulia.serhii@kpi.ua

Розглянуто метод FAIR як сучасну концептуальну модель кількісного аналізу ризиків інформаційної безпеки. Проаналізовано логічну структуру методу, принципи декомпозиції ризику, особливості оцінювання частоти подій та величини збитків. Розкрито практичні аспекти застосування FAIR в організаціях різних галузей, показано переваги моделі порівняно з традиційними якісними підходами. Стаття орієнтована на фахівців із інформаційної безпеки, аналітиків ризиків та дослідників, зацікавлених у застосуванні кількісних методів для управління ризиками кібербезпеки.

Ключові слова: ризик, інформаційна безпека, FAIR, кількісний аналіз, оцінювання збитків.

1. ВСТУП

З розвитком цифрових технологій організації дедалі частіше стикаються з кіберзагрозами, здатними спричинити як економічні, так і репутаційні втрати. У таких умовах особливо важливо не лише своєчасно виявляти потенційні загрози, а й точно оцінювати можливі наслідки інцидентів. Більшість традиційних методів аналізу ризиків ґрунтується на якісних або напівкількісних оцінках, у яких рівні ризику подаються у вигляді нечітких категорій на кшталт «низький», «середній» чи «високий». Подібний підхід ускладнює визначення очікуваних збитків, порівняння ризикових сценаріїв та обґрунтування витрат на кіберзахист.

Метод FAIR (Factor Analysis of Information Risk) виник як відповідь на потребу у стандартизованому, кількісному та фінансово орієнтованому способі оцінювання ризиків. Його основа полягає у поділі ризику на два взаємопов'язані компоненти — частоту подій та величину збитку — кожен з яких може бути поданий у числовому вигляді. FAIR пропонує чітко визначену термінологію, формалізує процес аналізу та забезпечує отримання прозорих і відтворюваних результатів.

Завдяки використанню кількісних показників FAIR дає змогу підтримувати ухвалення управлінських рішень, оптимізувати витрати на безпеку, визначати пріоритетні загрози та оцінювати економічну ефективність захисних заходів. Підхід органічно інтегрується у процеси управління ризиками, рекомендовані ISO/IEC 27005, NIST RMF та іншими галузевими методологіями.

Метою статті є аналіз концептуальної структури методу FAIR, розгляд особливостей кількісного підходу до оцінювання ризиків та визначення можливостей практичного використання цього методу в організаціях різного типу.

2. МЕТОД FAIR: СТРУКТУРА ТА ПРИНЦИПИ ОЦІНЮВАННЯ

Метод FAIR є однією з найпоширеніших сучасних моделей кількісного аналізу ризиків інформаційної безпеки, яка ґрунтується на декомпозиції ризику на логічно взаємопов'язані компоненти. На відміну від традиційних якісних методів, FAIR дозволяє подати ризик у кількісній формі, що робить можливим його об'єктивне порівняння, моделювання та фінансову інтерпретацію. Модель застосовується в різних галузях, де необхідно отримувати точні та відтворювані результати аналізу загроз.

2.1. Загальна структура моделі FAIR

Концептуально FAIR визначає ризик як добуток частоти події та величини збитку. Це базове співвідношення реалізоване через розгалужену таксономію факторів, яка дозволяє досліджувати кожен складову окремо та формувати деталізоване уявлення про природу потенційних інцидентів. Однією з переваг такого підходу є можливість поділу складних процесів на підфактори, що забезпечує точніший аналіз та підвищує об'єктивність оцінювання. У наукових джерелах структура FAIR зазвичай подається у вигляді графічної схеми, яка демонструє взаємозв'язки між факторами та їх ієрархічну організацію.



Рисунок 1. Структурна модель факторів ризику за методологією FAIR

2.2. Оцінювання частоти подій

Частота подій у FAIR формується на основі двох ключових параметрів: інтенсивності активності загрози та ймовірності того, що така активність спричинить реальний інцидент. Такий підхід дозволяє враховувати як доступні статистичні дані, так і експертні оцінки, що є важливим для середовищ з обмеженою історією інцидентів. Поєднання цих параметрів створює кількісну характеристику частоти інцидентів, яка є одним із визначальних елементів у загальній моделі ризику. Структуру логічного обґрунтування цього параметра зазвичай представляють у вигляді діаграми.

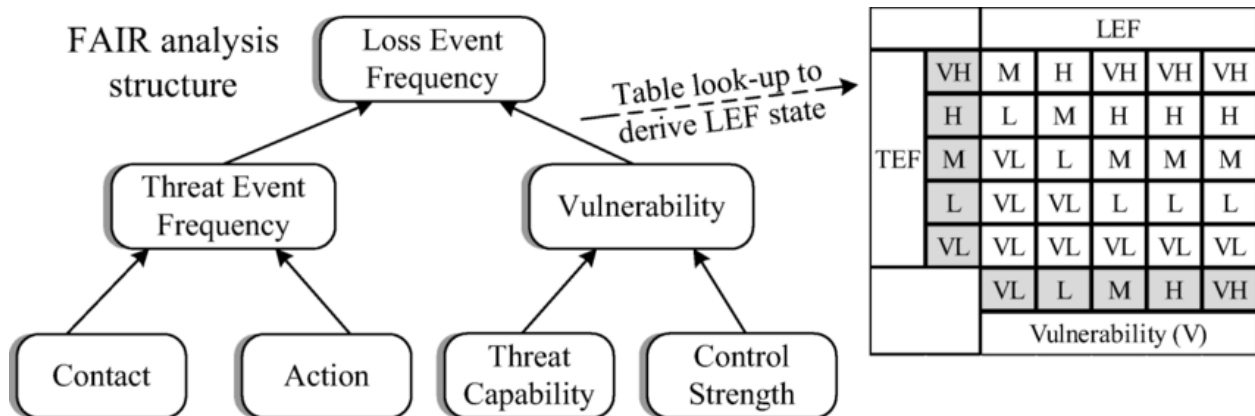


Рисунок 2. Структура логічного обґрунтування Loss Event Frequency у моделі FAIR

2.3. Моделювання величини збитків

Величина збитків у FAIR включає оцінювання як первинних, так і вторинних втрат. Первинні охоплюють прямі фінансові наслідки інциденту, тоді як вторинні пов'язані з довгостроковими або опосередкованими ефектами, серед яких репутаційні та юридичні наслідки. Такий поділ дає змогу точніше описати структуру потенційних втрат і моделювати широкий спектр сценаріїв. Для ілюстрації елементів величини збитків зазвичай використовують графічні схеми або приклади розподілу значень, що дозволяє наочно відобразити характер можливих фінансових наслідків.

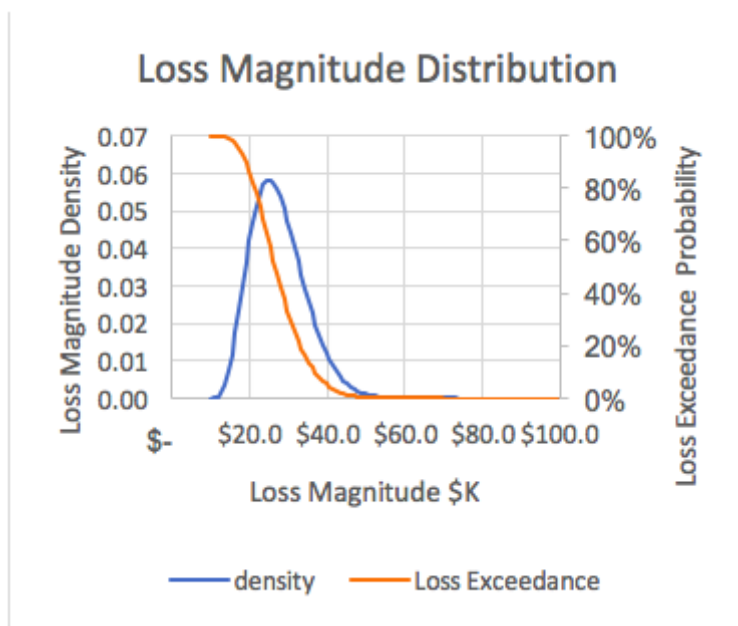


Рисунок 3. Приклад розподілу величини збитків

3. ПРАКТИЧНІ АСПЕКТИ ТА СФЕРИ ЗАСТОСУВАННЯ МЕТОДУ FAIR

Метод FAIR набув широкого застосування в організаціях, які прагнуть підвищити точність аналізу ризиків та забезпечити обґрунтованість рішень у сфері інформаційної безпеки. Його важливою перевагою є здатність адаптуватися до середовищ із різним рівнем доступності даних, що робить підхід ефективним як у великих корпоративних інфраструктурах, так і в організаціях, де історичні відомості про інциденти є обмеженими. У

цьому розділі розглянуто основні напрями застосування FAIR, особливості його інтеграції в існуючі процеси управління ризиками та потенціал використання моделі в умовах сучасної динаміки кіберзагроз.

3.1. Застосування FAIR у різних типах інформаційних систем

Метод FAIR може бути використаний у системах різного масштабу — від корпоративних мереж до хмарних платформ і промислових технологічних середовищ. У корпоративних ІТ-системах він дозволяє оцінювати типові загрози, пов'язані з витоками даних, несанкціонованим доступом чи порушеннями конфіденційності. У хмарних екосистемах підхід враховує специфіку таких середовищ, зокрема мультиорендність, віддалену обробку інформації та розподіл відповідальності між провайдером і клієнтом. У промислових системах особливу увагу приділяють оцінюванню операційних втрат та впливу можливих інцидентів на безперервність виробництва. У кожному випадку FAIR забезпечує стандартизований кількісний підхід, який дозволяє адаптувати аналіз до особливостей конкретної системи.

3.2. Інтеграція FAIR у процеси управління ризиками

FAIR є сумісним із широким спектром методологій управління ризиками, зокрема зі стандартами ISO/IEC 27005 та NIST SP 800-30. Завдяки цьому його можна використовувати як кількісний доповнювальний компонент у межах уже сформованих процесів ризик-менеджменту. У системах, де домінують якісні оцінки, FAIR допомагає уточнювати результати, які потребують більшої точності або деталізації. Використання фінансових показників робить отримані оцінки зрозумілими для керівництва, що спрощує обґрунтування інвестицій у засоби захисту та визначення пріоритетності заходів. Таким чином, FAIR сприяє підвищенню прозорості процедур і покращенню якості управлінських рішень.

3.3. Переваги практичного використання FAIR у сучасних організаціях

FAIR поєднує структурованість і гнучкість, дозволяючи аналітикам ідентифікувати ключові фактори, що впливають на рівень ризику, та спрямовувати ресурси на їх оптимізацію. Підхід ефективно працює з широкими діапазонами значень, що забезпечує реалістичність оцінок навіть за умов високої невизначеності. Використання фінансових вимірників спрощує комунікацію між технічними фахівцями та управлінським персоналом, оскільки надає можливість інтерпретувати ризики у зрозумілому й бізнес-орієнтованому форматі. Завдяки цьому FAIR виступає не лише математичним механізмом моделювання, а й інструментом стратегічного планування, який допомагає узгоджувати рішення у сфері кіберзахисту з загальними цілями організації.

3.4. Обмеження методу FAIR

Попри широке застосування, FAIR має низку обмежень, які необхідно враховувати під час практичного впровадження. Модель значною мірою залежить від точності вхідних даних, а в умовах обмеженої статистики результати базуються на експертних оцінках, що може знижувати об'єктивність. Крім того, підхід передбачає відносну незалежність окремих факторів, тоді як у реальних системах події часто мають складний причинно-наслідковий характер. Вибір відповідних імовірнісних розподілів також може бути викликом, оскільки FAIR не надає чітких рекомендацій щодо цього. Модель орієнтована переважно на фінансові показники, що іноді обмежує можливість урахування стратегічних, репутаційних або соціальних наслідків інцидентів.

4. ВИСНОВКИ

Метод FAIR є одним із найперспективніших підходів до кількісного аналізу ризиків інформаційної безпеки. Його концепція дає змогу формувати прозорі та відтворювані моделі ризику, що враховують як частоту кіберінцидентів, так і можливі фінансові наслідки. На відміну від якісних методів, FAIR пропонує послідовну логіку оцінювання, яка забезпечує основу для прийняття економічно обґрунтованих рішень.

Підхід вирізняється здатністю адаптуватися до різних типів інформаційних систем, підтримує аналіз за умов неповних або фрагментарних даних та органічно інтегрується у процеси управління ризиками відповідно до сучасних стандартів. Представлення ризику у фінансовому вимірі спрощує планування витрат на кіберзахист і покращує взаєморозуміння між технічними фахівцями та управлінським персоналом.

З огляду на це FAIR можна розглядати як універсальний інструмент для формування ефективної системи управління ризиками, яка відповідає вимогам сучасних інформаційних систем та викликам цифрової епохи. Його використання сприяє підвищенню точності аналізу, оптимізації ресурсів організації та забезпеченню стійкості до актуальних кіберзагроз.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Factor Analysis of Information Risk (FAIR) Model. URL: [https://www.fairinstitute.org/hubfs/Standards%20Artifacts/Factor%20Analysis%20of%20Information%20Risk%20\(FAIR\)%20Standard%20v3.0%20\(January%202025\).pdf](https://www.fairinstitute.org/hubfs/Standards%20Artifacts/Factor%20Analysis%20of%20Information%20Risk%20(FAIR)%20Standard%20v3.0%20(January%202025).pdf)
2. Anhtuan Le. Incorporating FAIR into Bayesian Network for Numerical Assessment of Loss Event Frequencies of Smart Grid Cyber Threats. *Mobile Networks and Applications*. October 2019
3. How to Model Controls in a FAIR Risk Analysis. URL: <https://www.fairinstitute.org/blog/how-to-model-controls-in-a-fair-risk-analysis>
4. Toward a FAIR Notion of Criticality. URL: <https://www.fairinstitute.org/blog/toward-a-fair-notion-of-criticality>
5. Jones, J., & Shmugar, F. *An Introduction to Factor Analysis of Information Risk (FAIR)*. FAIR Institute, 2020.

ОПИС ТА ПОРІВНЯЛЬНИЙ АНАЛІЗ ТЕХНОЛОГІЙ ДЛЯ ПОТОКОВОЇ ПЕРЕДАЧІ ВІДЕО-КОНТЕНТУ ДЛЯ E-LEARNING ПЛАТФОРМ

Потапчук О.М.¹, Кисельов Г.Д.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», м. Київ, Україна

¹lyoshachuk@gmail.com [0009-0000-2198-1520],

²g.kyselov@gmail.com [0000-0003-2682-3593]

Метою дослідження є аналіз технологій потокової передачі відео для оптимізації E-Learning платформ. Методами порівняльного аналізу протоколів (HLS, DASH) досліджено ефективність доставки даних. Новизна роботи полягає в адаптації критеріїв QoE для інтерактивних середовищ. Основним результатом є обґрунтування комбінованої стратегії: використання уніфікованого формату медіаданих із генерацією маніфестів HLS та DASH. Це слугує базою для проєктування масштабованих систем із низькою затримкою.

Ключові слова: потокова передача відеоконтенту, HLS, DASH, E-learning, масштабування, низька затримка.

1. ВСТУП

Сучасний етап розвитку глобальної мережі Інтернет характеризується експоненціальним зростанням обсягів мультимедіа-трафіку. За оцінками звіту Nokia про глобальний мережевий трафік, потокова передача відео-контенту (Video Streaming) домінує в мобільному трафіку споживачів, складаючи сьогодні 73% від загального обсягу. Очікується, що вона збереже своє лідерство, залишаючись на рівні понад 70% до 2033 року [1]. На цьому тлі особливої ваги набуває сфера дистанційної освіти (E-Learning), яка трансформувалася в основну форму взаємодії навчальних закладів та корпоративного сектору.

Специфіка E-Learning платформ висуває унікальні вимоги до технологій передачі даних, які відрізняються від класичних розважальних сервісів. На відміну від кінематографічного контенту (фільми, серіали), де ключовими є динаміка сцен та кольоропередача, освітній відеоконтент вимагає критичної чіткості статичних елементів, зокрема тексту презентацій, програмного коду, схем та формул. Артефакти стиснення, допустимі в розважальному відео, можуть зробити навчальний матеріал непридатним для сприйняття. Крім того, патерн споживання освітнього контенту передбачає часту навігацію (перемотування, паузи, повторний перегляд фрагментів), що ставить підвищені вимоги до ефективності сегментації відеопотоку та буферизації.

Існуючий спектр технологій доставки контенту (HLS, MPEG-DASH та сучасні кодеки) пропонує різні підходи до балансування між якістю зображення та бітрейтом. Проте, відсутність уніфікованої методики налаштування стрімінгових серверів саме під сценарії "On-Demand Education" часто призводить до надмірного навантаження на канали зв'язку або незадовільної якості відображення дрібних деталей на клієнтських пристроях.

У зв'язку з цим, проведення порівняльного аналізу технологій та розробка рекомендацій щодо оптимізації доставки навчального відеоконтенту є актуальним науково-прикладним завданням. Зокрема, необхідно знайти таке рішення, яке дозволить об'єднати

вимоги різних систем (як Apple, так і Android/Windows) без створення зайвих копій файлів. Це дасть можливість зберігати відео в одному форматі, заощаджуючи місце на серверах, і водночас гарантувати його плавне відтворення на будь-якому пристрої студента. Такий підхід не лише зменшить витрати на технічну інфраструктуру, а й забезпечить стабільний доступ до знань незалежно від того, який пристрій використовується для навчання.

2. ОГЛЯД СУЧАСНИХ ТЕХНОЛОГІЙ ДЛЯ ПОТОКОВОЇ ПЕРЕДАЧІ ВІДЕО-КОНТЕНТУ

Перш ніж перейти до детального огляду технологій потокової передачі відео-контенту, необхідно розрізнити два основні підходи до доставки відео через Інтернет:

- Прогресивне завантаження (Progressive Download) – це метод доставки відео, при якому завантажується один цілий відеофайл із сервера. Браузер дозволяє почати відтворення, щойно початкова частина файлу завантажиться у буфер, однак ключовим обмеженням є неможливість адаптивно змінювати якість "на льоту" або переходити до сегментів, які ще не були завантажені [2].

- Потокова передача (Streaming) – це набагато більш просунутий метод, у якому замість одного великого файлу, відео розбивається на безліч малих сегментів (зазвичай по 2-10 секунд). Плеєр завантажує ці сегменти послідовно, лише за кілька секунд до того, як користувач їх побачить. Це дає змогу миттєво «перестрибнути» у будь-яку точку відео, а також це уможливорює адаптивну потокову передачу, коли плеєр може перемикає якість цих сегментів (наприклад, з 1080p на 480p) залежно від швидкості інтернету користувача [3].

Використання методу прогресивного завантаження (розміщення простого MP4-файлу на сервері) є застарілим підходом, який створює низку критичних проблем для сучасного користувацького досвіду. Головний недолік цього підходу полягає у повній відсутності адаптивності. Користувач змушений завантажувати один файл фіксованої якості (наприклад, 1080p). Якщо його швидкість інтернет-з'єднання в якийсь момент падає нижче бітрейту цього файлу, відтворення відео зупиниться для тривалої буферизації. Платформа не може "на льоту" змінити якість на нижчу (наприклад, 480p), щоб пристосуватися до умов мережі, що призводить до негативного досвіду перегляду. Додатковою проблемою цього методу є неефективне використання трафіку. Плеєр може завантажити наперед значну частину файлу у високій якості (наприклад, 600 МБ), навіть якщо користувач вирішить переглянути лише перші дві хвилини відео. Це призводить до марнування пропускну здатності як для користувача, так і для сервера.

Отже, на основі цих недоліків можна сформулювати конкретний перелік вимог для підсистеми потокової передачі відео, яка потенційно може використовуватися в E-learning платформах:

- Підсистема повинна автоматично визначати швидкість інтернет-з'єднання користувача та динамічно перемикає якість відео. Це є ключовою вимогою для мінімізації ризику тривалої буферизації.

- Процес підготовки відео, як-от створення його версій для різної якості (1080p, 720p, 480p), має бути налаштований з пріоритетом на чіткість. Навіть якщо інтернет-з'єднання погане і користувач дивиться відео у низькій якості, зображення всередині відео має залишатися максимально розбірливим.

- Щоб забезпечити швидке завантаження та обслуговувати велику кількість одночасних користувачів по всьому світу, підсистема повинна бути спроектована для інтеграції з Content Delivery Network (CDN). Це дозволить кешувати відео-сегменти на серверах, географічно близьких до кінцевих користувачів, що кардинально знижує затримку та навантаження на сервери платформи.

Відповідно до сформульованих вимог, вибір оптимальної технології потокової передачі вимагає аналізу поточних рішень. Основою для них є концепція адаптивного бітрейту (Adaptive Bitrate Streaming, ABR). Цей механізм полягає в тому, що плеєр на пристрої користувача (у браузері чи мобільному додатку) динамічно аналізує поточну швидкість інтернет-з'єднання та автоматично перемикає якість відео [4]. Наприклад, якщо мережа користувача стабільна, плеєр обирає потік 1080p; якщо швидкість падає, то плеєр автоматично перемикається на потік 480p, щоб мінімізувати шанс зупинки відтворення відео.

Щоб реалізувати адаптивний бітрейт, сучасні стрімінгові технології використовують процес, який складається з кількох ключових етапів:

- Транскодування – це обчислювальний процес, під час якого вихідний відеофайл перекодовується у декілька окремих версій з різною якістю та бітрейтом (наприклад, 1080p, 720p, 480p, 360p) [5].
- Сегментація – перетворення кожної версії відео на малі сегменти, зазвичай тривалістю від 2 до 10 секунд. Така структура дозволяє плеєру приймати рішення про зміну якості саме в момент переходу від одного сегмента до іншого, що дає можливість зробити відтворення відео плавним та адаптивним [6].
- Створення маніфесту, який представляє з себе невеликий текстовий файл, який містить всю необхідну інформацію щодо відео, як-от список версій якості (1080p, 720p, 480p), список усіх сегментів, а також інформацію про аудіо та субтитри [7].
- Плеєр клієнта завантажує маніфест, аналізує швидкість з'єднання і починає послідовно завантажувати сегменти, обираючи найкращу якість, на яку спроможний інтернет користувача.

Хоча описаний процес (транскодування, сегментація, маніфест) є загальною концепцією, конкретна технічна реалізація, особливо формат маніфесту та відеосегментів, відрізняється залежно від використовуваного протоколу.

На сьогодні ринок практично повністю поділений між двома домінуючими протоколами адаптивної потокової передачі: HLS та MPEG-DASH.

HLS (HTTP Live Streaming) був розроблений компанією Apple у 2009 році. Спочатку він був створений, щоб забезпечити надійну потокову передачу на пристрої iPhone, які на той час не підтримували Flash. Завдяки домінуванню iOS на мобільному ринку, HLS швидко став де-факто стандартом. HLS працює за тією ж загальною схемою ABR, що й інші протоколи (транскодування, сегментація), але має свої особливості [8]:

- HLS має дві варіації відеосегментів. Історично протокол використовував сегменти у форматі .ts (MPEG Transport Stream) – це старий, але дуже надійний контейнер. Однак сучасні реалізації HLS (також відомі як "fHLS") все частіше використовують .fmp4 (Fragmented MP4). Перевага .fmp4 в тому, що він дозволяє використовувати однакові відеосегменти як для HLS, так і для MPEG-DASH, що значно спрощує процес підготовки та зберігання контенту.
- HLS використовує дворівневу структуру маніфесту. Спочатку плеєр завантажує головний маніфест (.m3u8), який не містить посилань на саме відео. Замість цього, він містить список "під-маніфестів", по одному для кожної доступної якості, наприклад, посилання на 1080p.m3u8 або 720p.m3u8. Вже після цього, плеєр, обравши потрібну якість, завантажує відповідний Медіа-маніфест (Media Playlist) (наприклад, 720p.m3u8), і саме в цьому файлі міститься список посилань на реальні відеосегменти (.ts або .fmp4), які потрібно завантажити.
- Як і випливає з назви, протокол працює поверх стандартного HTTP. Це означає, що відеосегменти є, по суті, звичайними файлами. Їх можна легко кешувати та розповсюджувати через будь-який веб-сервер або Content Delivery Network (CDN), що забезпечує високий рівень масштабованості.

- HLS нативно підтримується на всіх пристроях Apple (iPhone, iPad, Mac) у браузері Safari та всередині додатків. Для забезпечення потокової передачі відео на пристроях Apple, необхідно використовувати протокол HLS.

Таким чином, HLS є потужним і обов'язковим рішенням для охоплення екосистеми Apple. Однак, оскільки ця технологія історично контролюється однією компанією, індустрія потребувала єдиного, відкритого та незалежного стандарту. Саме цією відповіддю став MPEG-DASH (Dynamic Adaptive Streaming over HTTP). DASH був створений з метою уніфікувати ринок потокової передачі під єдиним, гнучким протоколом. Даний протокол має наступні особливості [9]:

- На відміну від текстового .m3u8 в HLS, DASH використовує Media Presentation Description (MPD) – файл у форматі XML. Цей маніфест має значно складнішу, але й набагато гнучкішу структуру. Він детально описує всі доступні потоки (відео різної якості, аудіодоріжки різними мовами, субтитри) та те, як вони поділені на сегменти.

- Стандарт не "прив'язаний" до конкретних відео- чи аудіокодеків. Він може однаково ефективно доставляти контент, закодований у H.264, H.265/HEVC, VP9 або новітньому AV1, що дає платформі гнучкість у виборі оптимального кодека.

- DASH з самого початку був спроектований для роботи з сегментами Fragmented MP4 (.fmp4). Це той самий формат, який зараз впроваджує HLS ("fHLS"), що свідчить про рух індустрії саме до цього типу сегментів.

- Головним обмеженням MPEG-DASH є відсутність повної нативної підтримки в екосистемі Apple, як-от у браузері Safari на macOS та iOS. Для відтворення DASH-потоків у більшості веб-браузерів (Chrome, Firefox, Edge) зазвичай використовуються JavaScript-плеєри, такі як Shaka Player або dash.js, що працюють через браузерний API Media Source Extensions (MSE). У Safari підтримка DASH залишається обмеженою: хоча нові версії браузера впроваджують Managed Media Source API (MMS), стандартним і найбільш сумісним форматом для пристроїв Apple досі є HLS (HTTP Live Streaming).

3. ОБҐРУНТУВАННЯ ОПТИМАЛЬНОГО ПІДХОДУ ДО ОРГАНІЗАЦІЇ ВІДЕОСТРІМІНГУ В E-LEARNING СИСТЕМАХ

На основі проведеного аналізу HLS та MPEG-DASH, найбільш вдалим, гнучким та технічно досконалим рішенням для більшості E-learning платформ є впровадження комбінованої стратегії, яка полягає у підготовці контенту в єдиному форматі, але з наданням двох типів маніфестів (HLS та DASH), щоб гарантувати найкращу якість відтворення на абсолютній більшості пристроїв.

Замість того, щоб обирати один протокол, серверна частина має готувати відео-контент так, щоб він був сумісним з обома протоколами (HLS та MPEG-DASH). Це досягається завдяки використанню сегментів Fragmented MP4 (.fmp4). Даний підхід можна описати прикладом, який складається з наступної послідовності кроків:

- вхідний відеофайл (наприклад, «video.mp4») надсилається до Video Streaming Service;
- даний сервіс перекодує відео у набір профілів якості (наприклад, 1080p, 720p, 480p). На цьому етапі критично важливо застосувати спеціальні налаштування кодека (наприклад, H.264), які пріоритезують чіткість тексту та деталей UI, а не плавність руху;
- кожен із цих профілів якості нарізається на сегменти у форматі «.fmp4». Вибір цього розширення є важливим, оскільки «.fmp4» є універсальним контейнером, який підтримується як HLS, так і DASH;

- на основі одного й того ж набору .fmp4 сегментів сервер генерує два типи маніфестів: manifest.m3u8 (для HLS), manifest.mpd (для MPEG-DASH).

Усі згенеровані файли (.fmp4 сегменти, .m3u8 та .mpd маніфести) автоматично завантажуються та кешуються у Мережі Доставки Контенту (CDN). Це гарантує, що контент буде фізично близьким до користувача в будь-якій точці світу, забезпечуючи мінімальну затримку та високу швидкість завантаження.

На стороні клієнта (у веб-браузері) буде використано сучасний JavaScript-плеєр (наприклад, Shaka Player або Video.js), який відповідатиме за автоматичний вибір протоколу в залежності від того, у якому браузері та ОС він запущений. Використання цього підходу гарантує охоплення абсолютної більшості пристроїв, від iPhone до десктопів на Windows, з використанням найкращого нативного протоколу для кожного середовища.

4. ВИСНОВКИ

У ході дослідження проведено аналіз методів доставки відеоконтенту та встановлено, що для сучасних E-Learning платформ використання методу прогресивного завантаження є неефективним через відсутність адаптивності до умов мережі. Визначено, що оптимальним підходом є використання технологій адаптивного бітрейту (ABR).

Порівняльний аналіз протоколів HLS та MPEG-DASH виявив, що жоден із них окремо не забезпечує універсального покриття всіх пристроїв без технічних обмежень: HLS є стандартом для екосистеми Apple, тоді як DASH пропонує більшу гнучкість та відкритість.

Основним практичним результатом роботи є обґрунтування комбінованої архітектури доставки контенту. Запропонований підхід базується на використанні уніфікованого контейнера Fragmented MP4 (.fmp4), що дозволяє генерувати два типи маніфестів (.m3u8 та .mpd) для одного набору відеосегментів. Це рішення дозволяє уникнути дублювання даних на серверах зберігання, забезпечує охоплення більшості клієнтських пристроїв та гарантує вибір нативного плеєра для кожної платформи.

Додатково сформульовано рекомендації щодо налаштування профілів кодування для освітнього контенту: пріоритезація чіткості статичних елементів (тексту, інтерфейсів) над динамікою руху є критичною умовою для забезпечення високої якості сприйняття (QoE) у навчальному процесі. Впровадження запропонованих рішень у поєднанні з використанням CDN дозволяє побудувати масштабовану, відмовостійку систему з мінімальною затримкою доступу до матеріалів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Звіт Nokia про глобальний мережевий трафік. URL: <https://onestore.nokia.com/asset/213660> (дата звернення: 21.11.2025).
2. Огляд методу прогресивного завантаження (Progressive Download). URL: <https://www.cdnetworks.com/glossary/progressive-download/> (дата звернення: 21.11.2025).
3. Огляд методу потокової передачі відео (Video Streaming). URL: <https://www.cloudflare.com/learning/video/what-is-streaming/> (дата звернення: 21.11.2025).
4. Що таке Adaptive Bitrate Streaming (ABR)? URL: <https://www.cloudflare.com/learning/video/what-is-adaptive-bitrate-streaming/> (дата звернення: 21.11.2025).
5. Визначення процесу транскодування відео. URL: <https://aws.amazon.com/what-is/video-transcoding/> (дата звернення: 21.11.2025).
6. Визначення процесу сегментації відео. URL: <https://cloudinary.com/glossary/video-chunk> (дата звернення: 21.11.2025).

7. Роль маніфесту у контексті потокової передачі відео-контенту. URL: <https://www.fastpix.io/blog/what-is-a-video-manifest> (дата звернення: 21.11.2025).

8. Особливості протоколу HTTP Live Streaming (HLS). URL: <https://www.cloudflare.com/learning/video/what-is-http-live-streaming/> (дата звернення: 21.11.2025).

9. Особливості протоколу MPEG-DASH. URL: <https://www.cloudflare.com/learning/video/what-is-mpeg-dash/> (дата звернення: 21.11.2025).

РІЗОНІНГ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ В МЕДИЧНИХ СИСТЕМАХ

Прокопенко Д.М.¹, Кислий Р.В.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ pdimon54@gmail.com

У роботі проведено порівняльне оцінювання чотирьох великих мовних моделей (LLM) у задачі стислого резюмування клінічних діалогів у форматі JSON. Дослідження виконано на датасеті з відкритими деанонізованими медичними даними. Для оцінки якості використано лексичні метрики на основі перекриття токенів (precision/recall/F1) та корпусу BLEU (SacreBLEU). Проведено два експерименти з різними налаштуваннями: короткий вхідний промпт та меншу кількість прикладів; довший промпт та більшу кількість прикладів. Метою цієї роботи є дослідження та виявлення найкращих способів використання різонінгу великих мовних моделей в медичних системах.

Ключові слова: LLM, клінічні діалоги, медичні дані, оцінка якості, F1, BLEU, DeepSeek, GPT-OSS, gpt-5-mini.

1. ВСТУП

Використання великих мовних моделей у медичних системах підтримки клінічних рішень стрімко зростає, зокрема у задачах структурованого резюмування діалогів «лікар–пацієнт». Автоматичне формування стислого пацієнтського висновку у вигляді структурованих даних дозволяє пришвидшити роботу зі скаргами, анамнезом і результатами огляду, а також полегшує подальшу інтеграцію результатів у медичні інформаційні системи.

Разом із тим LLM різних сімейств можуть суттєво відрізнятись за точністю виділення клінічних фактів, повнотою резюме та стійкістю до довгих промπτів. Тому необхідним є експериментальне порівняння моделей на єдиному датасеті й у однаковому інтерфейсі виклику.

2. ПОСТАНОВКА ЗАДАЧІ

Метою дослідження є кількісне порівняння чотирьох LLM у задачі резюмування клінічних діалогів з формуванням відповіді виключно у структурованому JSON-форматі.

Об'єкт дослідження: моделі

1. deepseek-r1-distill-qwen-14b (OpenRouter),
2. openai/gpt-oss-120b (HuggingFace Router),
3. openai/gpt-oss-20b (HuggingFace Router),
4. gpt-5-mini (OpenAI).

Предмет дослідження: якість і стабільність сформованих JSON-резюме при різних умовах подачі контексту.

3. ПІДХІД ДО РЕАЛІЗАЦІЇ ТА МЕТОДИ ОЦІНЮВАННЯ

3.1. Датасет і промпт

Для тестування використано корпус AGBonnet/augmented-clinical-notes [1], який містить пари:

- conversation – діалог лікар–пацієнт,
- note – клінічна нотатка,
- summary – еталонне структуроване резюме у JSON.

Промпт формувався як інструкція для моделі: «Given a patient-doctor conversation and a clinical note, produce a concise patient summary in JSON. Return ONLY JSON summary.»

3.2. Метрики

Через те, що моделі генерують не завжди однакову JSON-структуру чи порядок полів, застосовано текстові метрики, стійкі до перестановок і стилістичних відмінностей.

Токенізація та лексичні метрики є стандартним підходом у NLP для оцінювання схожості між згенерованим текстом та еталонним зразком [2]. Такі метрики не вимагають складної розмітки, легко інтерпретуються та дають швидкий числовий сигнал про якість відповідей. У задачі формування медичного висновку у форматі JSON вони дозволяють оцінити, наскільки модель потрапляє у зміст еталона (лексично і структурно), а також наскільки стабільно генерує валідні відповіді.

Precision.

Precision вимірює, яка частка токенів, згенерованих моделлю, справді присутня в еталоні. Для цього обидва тексти (prediction і reference) нормалізуються: переводяться в lowercase, прибирається пунктуація, після чого виконується токенізація. Нехай P — множина (або multiset) токенів предикта, R – токени референса. Тоді:

$$precision = \frac{|P \cap R|}{|P|}.$$

Високе значення precision означає, що модель дає точні відповіді і рідше додає зайву інформацію. Низький precision свідчить про схильність до галюцинацій або надмірного деталізування, яке не підтверджується еталоном.

Recall.

Recall показує, яку частку токенів з еталона модель змогла відтворити у своїй відповіді:

$$recall = \frac{|P \cap R|}{|R|}.$$

Високий recall означає, що модель покриває більшість потрібних фактів. Низький recall свідчить про пропуски важливих елементів у summary (наприклад, відсутність деталей про анамнез, результати огляду тощо).

F1 (token overlap).

F1-score – гармонічне середнє precision та recall, яке балансує між точністю і повнотою:

$$F1 = \frac{2 * precision * recall}{precision + recall}.$$

Це основна інтегральна метрика для порівняння моделей у задачі сумаризації, оскільки відображає компроміс між зайвою інформацією і пропущеними фактами. У межах експерименту саме F1 використовується як ключовий показник загальної якості лексичного збігу.

BLEU.

Метрика BLEU є класичною n-gram метрикою для порівняння з еталоном, яка широко застосовується у задачах машинного перекладу та сумаризації [3]. BLEU оцінює, наскільки n-грамні послідовності (1-грам, 2-грам, ...) у відповіді моделі збігаються з референсом, з урахуванням штрафу за надто короткі тексти:

$$BLEU = BP \cdot \exp \left(\sum_{n=0}^N \omega_n \log p_n \right),$$

де p_n – precision для n-грам, ω_n - ваги, BP – штраф за коротку відповідь [4].

У цій роботі BLEU рахується над нормалізованими JSON-рядками. Вища BLEU означає кращу текстову відповідність еталону на рівні фраз/шаблонів. Обмеження метрики – чутливість до перефразування: модель може передати той самий зміст іншими словами і отримати нижчий BLEU.

Примітка: метрика structured_f1 у поточній реалізації дорівнює 0 для всіх моделей через часті помилки парсингу (обірвані відповіді). Це розглядається як обмеження експерименту.

4. ЕКСПЕРИМЕНТИ ТА РЕЗУЛЬТАТИ

4.1. Експеримент 1: 25 прикладів, довжина 1024 символи

У першому режимі оцінювання використано 25 випадкових прикладів train-спліту при фіксованому SEED. Довжина промπτу була обмежена приблизно 1024 символами.

Середні метрики (25 прикладів):

	model	precision	recall	f1	structured_precision	structured_recall	structured_f1	bleu	coverage
0	deepseek-r1-distill-qwen-14b (OpenRouter)	0.673477	0.239418	0.350023	0.0	0.0	0.0	8.054292	1.0
1	openai/gpt-oss-20b (HF)	0.644762	0.233619	0.338888	0.0	0.0	0.0	4.967230	1.0
2	openai/gpt-oss-120b (HF)	0.608834	0.211273	0.309283	0.0	0.0	0.0	6.765577	1.0
3	gpt-5-mini (OpenAI)	0.650072	0.210524	0.297701	0.0	0.0	0.0	5.683069	1.0

Рисунок 1. Таблиця результатів для експерименту 1

У цьому режимі найкращий F1 демонструє DeepSeek-14B, що свідчить про хороший баланс точності та повноти на коротших контекстах. Моделі GPT-OSS показують близький рівень, а gpt-5-mini відстає через нижчий recall.

4.2. Експеримент 2: 300 прикладів, промπτ 2500 символів

У другому режимі було збільшено кількість тестових прикладів до 300 і дозволено довший вхідний контекст (~2500 символів).

Середні метрики (300 прикладів):

	model	precision	recall	f1	structured_precision	structured_recall	structured_f1	bleu	coverage
0	gpt-5-mini (OpenAI)	0.549522	0.332324	0.408735	0.0	0.0	0.0	13.440018	1.00
1	openai/gpt-oss-20b (HF)	0.663767	0.222287	0.329196	0.0	0.0	0.0	3.917629	1.00
2	openai/gpt-oss-120b (HF)	0.621066	0.214319	0.315186	0.0	0.0	0.0	5.753172	1.00
3	deepseek-r1-distill-qwen-14b (OpenRouter)	0.643249	0.201671	0.301879	0.0	0.0	0.0	4.705710	0.93

Рисунок 2. Таблиця результатів для експерименту 2

На довших промптах gpt-5-mini суттєво покращив показник recall і F1, ставши лідером. Це означає, що модель краще оброблює довший клінічний контекст та переносить більше фактів у відповідь. DeepSeek, навпаки, втрачає як за F1/BLEU, так і за coverage (0.93), що може вказувати на проблеми зі стабільністю або ліміти при довгій генерації. GPT-OSS моделі залишаються стабільними (coverage 1.0), але їхня повнота (recall) не росте разом із довжиною промπτу.

5. ВИСНОВКИ

У роботі проведено два режими оцінювання чотирьох LLM на задачі структурованого резюмування клінічних діалогів у JSON. Отримано такі висновки:

1. Найкращу якість у реалістичному режимі (довгий контекст, 300 прикладів) показує gpt-5-mini, що підтверджується найбільшими F1 і BLEU.
2. DeepSeek-14B ефективний на коротших промптах, але втрачає стабільність і якість при подовженні контексту.
3. Моделі GPT-OSS (20B/120B) демонструють стабільність генерації, проте їхня повнота (recall) залишається нижчою, ніж у gpt-5-mini.
4. Поточна реалізація structured-метрики виявилась непридатною через часті невалідні JSON-відповіді. У майбутньому доцільно:
 - підсилити інструкцію/валідацію JSON у промпті,
 - застосувати пост-процесінг для відновлення JSON.

Практична рекомендація: для систем клінічного резюмування з довгим контекстом варто обирати gpt-5-mini, а в середовищах із короткими промптами чи обмеженою вартістю можна розглядати GPT-OSS-20B або DeepSeek-14B з додатковим контролем стабільності.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. AGBonnet. *Augmented Clinical Notes Dataset*. HuggingFace datasets, 2024.
2. Jurafsky D., Martin J. *Speech and Language Processing*. 3rd ed. draft, 2023.
3. Papineni K. et al. *BLEU: a Method for Automatic Evaluation of Machine Translation*. ACL, 2002.
4. Post M. *A Call for Clarity in Reporting BLEU Scores*. WMT, 2018 (SacreBLEU).

РОЗРОБКА ЗАСОБІВ АВТОМАТИЗОВАНОГО СТВОРЕННЯ ДОСТУПНИХ ІТ-ГРАФІЧНИХ ЕЛЕМЕНТІВ

Ревякін Д.Е.¹, Кисельов Г.Д.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ revyakin.danylo@lil.kpi.ua, ² g.kyselov@gmail.com [0000-0003-2682-3593]

Метою дослідження є розробка веб-додатку для автоматизованої конвертації математичних формул у форматі LaTeX у множину доступних форматів, що забезпечують рівнозначний доступ до математичної інформації для людей з обмеженими можливостями зору. Дослідження базується на методах семантичного збагачення MathML структур, застосуванні правил MathSpeak/ClearSpeak для голосового виводу, алгоритмах конвертації у брайлівський код Nemeth та інтеграції OCR технологій для розпізнавання формул з зображень. Результатом роботи є функціональна система, що реалізує повний цикл обробки математичних формул від введення до генерації доступних форматів. Наукова новизна полягає в комплексному підході до забезпечення доступності математичних формул, що поєднує семантичне збагачення, інтерактивну навігацію та множину форматів виводу в єдину систему. Практична значимість розробки полягає в можливості використання системи в освіті, науці та бібліотеках для створення інклюзивного освітнього середовища та забезпечення рівного доступу до математичної інформації.

Ключові слова: доступність, математичні формули, семантичне збагачення, брайлівський код, голосовий вивід, інтерактивна навігація.

1. ВСТУП

Сучасні веб-технології та стандарти доступності [1, 2] відкривають нові можливості для забезпечення рівного доступу до математичної інформації для людей з обмеженими можливостями зору. Математичні формули, які є невід'ємною частиною освітнього та наукового контенту, традиційно представлені у вигляді графічних зображень або спеціалізованих форматів, що ускладнює їх сприйняття за допомогою скрін-рідерів та брайлівських дисплеїв.

Ключовим завданням є перетворення візуального представлення інформації у структури, зрозумілі допоміжним технологіям. Це включає аналіз контенту, виділення його логічних елементів та побудову семантичної моделі [3], яка описує зміст і взаємозв'язки в межах документу. На основі такої моделі можуть генеруватися альтернативні формати – текстові описи на основі спеціалізованих правил [4, 5], голосові інструкції [6, 7], тактильні подання або структуровані дані, з якими можуть працювати спеціалізовані пристрої.

Важливою частиною доступності є інтерактивність. Користувач повинен мати можливість навігувати складною структурою не лише візуально, а й за допомогою допоміжних технологій, отримуючи послідовне озвучення та можливість переходу між елементами.

Архітектура систем, що працюють із доступністю, зазвичай побудована модульно. Кожен компонент відповідає за окремий етап: аналіз структури, формування альтернативних представлень, озвучення, тактильне кодування чи інтерактивну навігацію. Такий підхід забезпечує гнучкість, масштабованість і можливість подальшого розвитку.

2. ПОСТАНОВКА ЗАДАЧІ

Метою дослідження є розробка веб-додатку для конвертації математичних формул у форматі LaTeX у множину доступних форматів, що забезпечують рівнозначний доступ до математичної інформації для людей з обмеженими можливостями зору. Система повинна реалізувати конвертацію у голосовий вивід з підтримкою MathSpeak/ClearSpeak правил [4, 5], брайлівський код у форматах Nemeth та UEB+Nemeth [8, 9], семантичні описи та інтерактивну навігацію по структурі формул. Додатково система повинна підтримувати розпізнавання формул з зображень за допомогою OCR технологій [10, 11].

3. РЕАЛІЗАЦІЯ ТА РЕЗУЛЬТАТИ

У ході досліджень був розроблений веб-додаток мовою програмування JavaScript з використанням сучасних веб-технологій, що і є результатом роботи. Система реалізує повний цикл обробки математичних формул від введення до генерації доступних форматів.

Система підтримує два способи введення формул: текстове введення у форматі LaTeX та завантаження зображень з подальшим OCR розпізнаванням. Текстове введення забезпечує прямий контроль над формулою та дозволяє користувачу використовувати стандартний LaTeX синтаксис. Для конвертації LaTeX у MathML використовується бібліотека KaTeX [12], яка забезпечує швидкий рендеринг математичних формул та генерацію MathML структури. KaTeX обрано замість MathJax через його синхронну роботу, високу продуктивність та відсутність необхідності додаткових завантажень шрифтів.

Система генерує візуальне представлення формули за допомогою KaTeX, що забезпечує високу якість рендерингу та швидкість відображення. Візуальне представлення корисне не тільки для перевірки коректності введення, але й для візуального контролю результатів обробки (рис. 1).

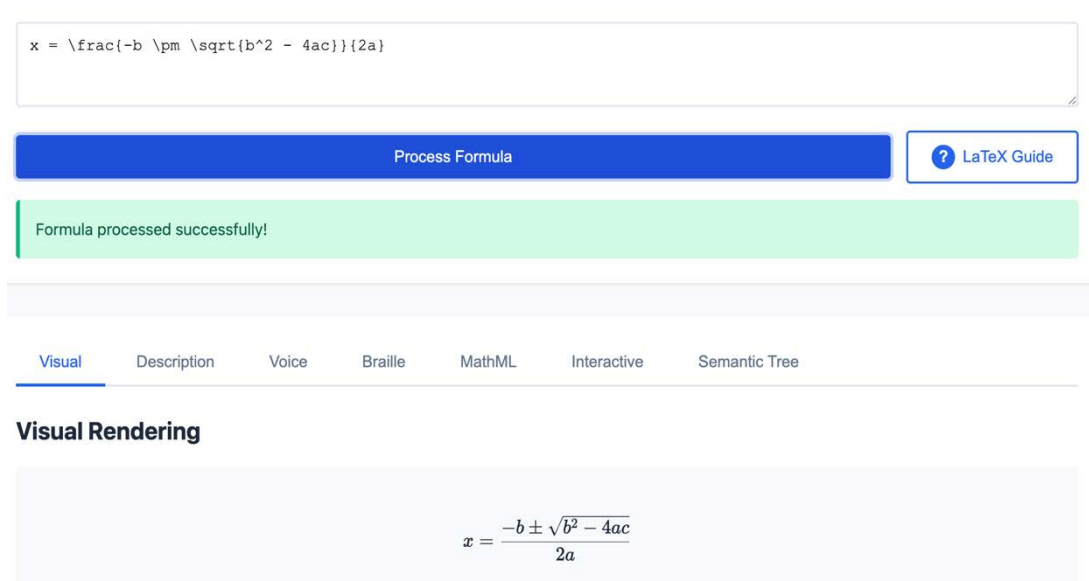


Рисунок 1. Візуальне представлення формули

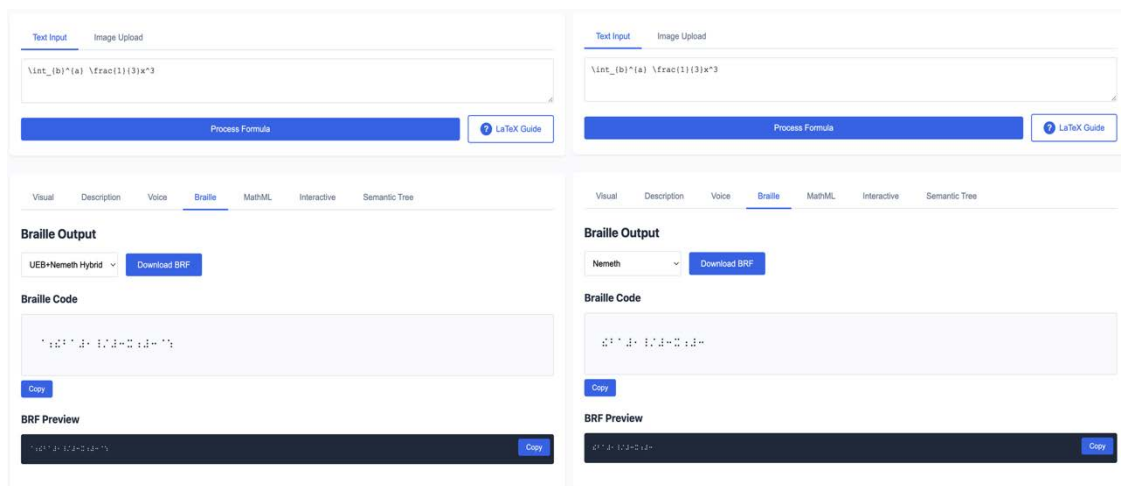


Рисунок 2. Візуальне представлення формули

3.4. Інтерактивна навігація

Інтерактивна навігація дозволяє користувачам досліджувати структуру формули, рухаючись від кореня до окремих частин. Вона реалізується через побудову навігаційного дерева на основі семантичного дерева. Кожен вузол містить інформацію про тип, роль, шлях, опис та LaTeX представлення відповідної частини формули. Кожна частина формули виділяється візуально та автоматично читається голосом при фокусуванні.

Навігація може здійснюватися через клавіатуру (стрілки, Home, Space) або через кнопки інтерфейсу. Система підтримує навігацію у всіх напрямках: до батьківського елемента, до дочірніх елементів, до наступного/попереднього брата (рис. 3). Користувач може експортувати окремі частини формули у різних форматах (LaTeX, MathML, текст, брайлівський код).

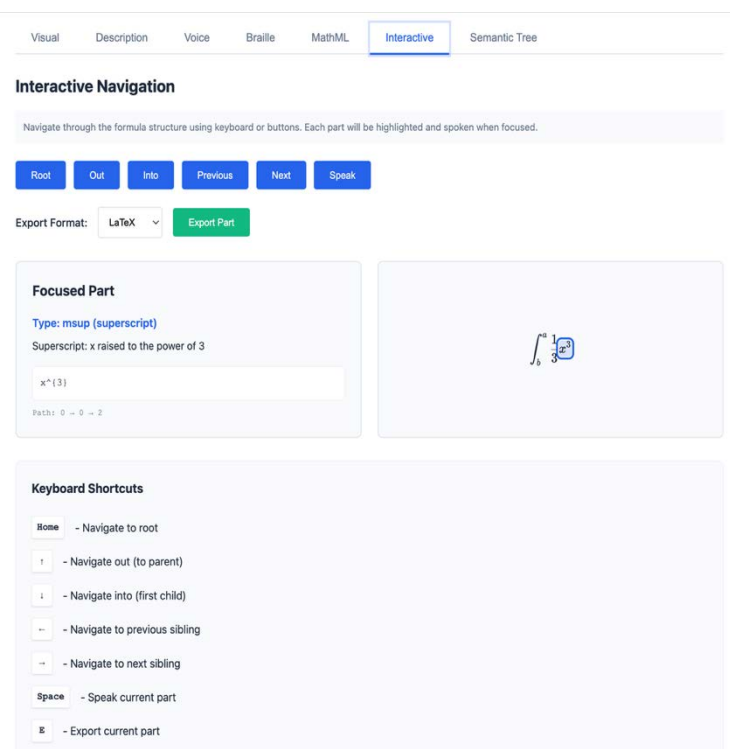


Рисунок 3. Інтерактивна навігація по структурі формули з виділенням поточної частини

4. ВИСНОВКИ

У даній роботі було розроблено веб-додаток для конвертації математичних формул у форматі LaTeX у множину доступних форматів. В процесі виконання роботи було розглянуто різні підходи до забезпечення доступності математичних формул та реалізовано комплексне рішення, що поєднує семантичне збагачення, голосовий вивід, брайлівську конвертацію та інтерактивну навігацію.

Результатом роботи стала функціональна система, що забезпечує рівнозначний доступ до математичної інформації для людей з обмеженими можливостями зору. Система демонструє високу точність обробки для стандартних математичних структур та забезпечує зручний інтерфейс для роботи з формулами.

Майбутні напрямки розвитку системи включають розширення підтримки LaTeX команд, інтеграцію з LibLouis для більш точної брайлівської конвертації, додавання підтримки інших мов для голосового виводу, реалізацію офлайн режиму роботи та експорт у формати EPUB та DAISY для використання в електронних книгах.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. W3C Web Accessibility Initiative. Web Content Accessibility Guidelines (WCAG) 2.1 : W3C Recommendation, June 2018. URL: <https://www.w3.org/TR/WCAG21/> (дата звернення: 15.10.2025).
2. Colorado State University, Office of Online Learning. MathML - Accessibility by Design, 2023. URL: <https://www.chhs.colostate.edu/accessibility/best-practices-how-tos/mathml/> (дата звернення: 15.10.2025).
3. W3C. Mathematical Markup Language (MathML) Version 3.0 2nd Edition : W3C Recommendation. URL: <https://www.w3.org/TR/MathML3/> (дата звернення: 15.10.2025).
4. Frankel L., Brownstein B., Soiffer N. Development and Initial Evaluation of the ClearSpeak Style for Automated Speaking of Algebra. ETS Research Report Series, 2016. P. 1–26. DOI: 10.1002/ets2.12103.
5. Soiffer N. MathSpeak: A Non-ambiguous Language for Audio Rendering of MathML. Proceedings of the Technology of Assistive Technology, 2005. P. 81–90.
6. Raman T. V. Audio System for Technical Readings : дис. ... д-ра філ. наук. Cornell University, 1997. 245 с.
7. W3C Community and Incubation Group. Web Speech API : W3C Community Group Draft Report. URL: <https://wicg.github.io/speech-api/> (дата звернення: 29.10.2025).
8. Braille Authority of North America. The Nemeth Braille Code for Mathematics and Science Notation : 2022 Revision. URL: <https://www.brailleauthority.org/nemeth-code> (дата звернення: 11.11.2025).
9. Nemeth Braille Code for Mathematics and Science Notation : 1972 Revision. URL: <https://www.prcvi.org/media/1155/nemeth-code-mathematics-science-notation-1972-revision.pdf> (дата звернення: 11.11.2025).
10. Kesada D., Radeanu D. Recognizing Math Equations with Computer Vision. Roboflow Blog, April 2024. URL: <https://blog.roboflow.com/math-equations-computer-vision/> (дата звернення: 20.10.2025).
11. Bluche T., Messina R. Mathematical Formula Recognition using Machine Learning Techniques. Thèse de doctorat. Université de Rouen, 2013. 185 с.
12. Khan Academy. KaTeX – The fastest math typesetting library for the web. URL: <https://katex.org/> (дата звернення: 18.10.2025).

МЕТОДИ ДИНАМІЧНОЇ КОМПОЗИЦІЇ ПЛАТФОРМИ МЕДИЧНОЇ ДІАГНОСТИКИ НА ОСНОВІ РЕКРУТИНГУ LLM-АГЕНТІВ

Ренський М.М.¹, Письменний І.О.²

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ renskiy.mikhail@lil.kpi.ua, ² pysmennyi.ihor@lil.kpi.ua

Досліджено методи динамічної композиції платформи медичної діагностики та механізм рекрутингу інтелектуальних агентів на основі великих мовних моделей (LLM). Реалізацію виконано у вигляді мікросервісної архітектури в хмарному середовищі AWS, що включає центральний оркестратор та ізольовані контейнеризовані агенти для обробки візуальних та часових даних. Запропонований підхід забезпечує підвищення гнучкості системи, точності мультимодальної діагностики та масштабованості при роботі з різнорідними медичними даними (МРТ, ЕЕГ). Новизною є інтеграція LLM-оркестратора для семантичного аналізу запитів та динамічного розподілу задач між вузькопрофільними моделями, що дозволяє автоматизувати виявлення патологій (пухлини, епілептичні судоми) в реальному часі.

Ключові слова: мультиагентні системи, LLM-рекрутинг, медична діагностика, динамічна композиція, AWS, Computer Vision, Time-series analysis.

1. ВСТУП

Проблема ефективної інтеграції гетерогенних медичних даних стає ключовою у задачах сучасної телемедицини та клінічної діагностики. Обробка мультимодальної інформації (візуалізація МРТ, часові ряди ЕЕГ, текстові анамнези) у динамічному середовищі потребує вирішення задач сумісності моделей, точності інтерпретації та швидкодії, що є актуальною науково-прикладною задачею. Сучасні системи часто працюють як ізольовані рішення ("окремо для знімків", "окремо для тексту") або покладаються на універсальні LLM, які схильні до галюцинацій у вузькоспеціалізованих питаннях. У таких умовах класичні монолітні архітектури не забезпечують необхідної гнучкості, масштабованості та верифікованості діагностичних висновків.

У рамках даної роботи мною запропоновано архітектуру платформи динамічної композиції сервісів на основі LLM-оркестратора, яка забезпечує інтелектуальний рекрутинг спеціалізованих агентів та адаптивну маршрутизацію завдань під час діагностичного сеансу. Запропонована система виконує семантичний аналіз вхідного запиту (текст, зображення або потік даних), визначає необхідний профіль експертизи (виявлення пухлин або аналіз судомної активності) та динамічно ініціює відповідні контейнеризовані мікросервіси. Це дозволяє автоматизувати процес мультимодальної діагностики з метою підвищення точності висновків та забезпечення роботи в режимі реального часу.

2. ТЕОРЕТИЧНА СКЛАДОВА АЛГОРИТМІВ

Для формалізації задачі динамічної композиції введемо множину доступних агентів-діагностів:

$$\mathcal{A} = \{a_1, a_2, \dots, a_n\},$$

де кожен агент a_i характеризується вектором опису своїх компетенцій D_i (наприклад, "виявлення пухлин", "аналіз ЕЕГ").

Вхідний запит до системи (промпт лікаря або метадані файлу) позначимо як Q . Процес рекрутингу реалізується оркестратором як функція відображення f , що обирає оптимального агента для поточного запиту:

$$a^* = f(Q, \mathcal{A}),$$

що повертає оптимального агента для поточного запиту.

Для визначення відповідності агента контексту запиту використовується семантична подібність між:

- вектором запиту $\mathbf{v}_Q$, отриманим шляхом ембеддингу промпту,
- вектором компетенцій агента $\mathbf{v}_{D_i}$.

Схожість вимірюється за допомогою косинусної міри:

$$S(Q, a_i) = \frac{\mathbf{v}_Q \cdot \mathbf{v}_{D_i}}{\|\mathbf{v}_Q\| \|\mathbf{v}_{D_i}\|} = \frac{\sum_{j=1}^m v_{Q,j} v_{D_i,j}}{\sqrt{\sum_{j=1}^m v_{Q,j}^2} \sqrt{\sum_{j=1}^m v_{D_i,j}^2}},$$

де m – розмірність простору ембеддингів.

Агент активується, якщо ступінь подібності перевищує поріг впевненості τ :

$$S(Q, a_i) \geq \tau,$$

де τ – параметр чутливості системи (наприклад, $\tau = 0.75$). Він запобігає хибним викликам агентів, тобто «галюцинаціям» маршрутизації.

Для агентів, що працюють із потоками (наприклад, ЕЕГ), вводиться вимога реального часу.

Нехай:

- T_{arr} – час надходження пакета,
- $T_{proc}(a_i)$ – час обробки агента,
- T_{net} – мережеві затримки,
- T_{win} – розмір часового вікна.

Тоді виконується обмеження:

$$T_{proc}(a_i) + T_{net} < T_{win}.$$

Це гарантує, що агент встигає обробити дані в потрібному темпі.

Підсистема оркестрації виконує:

$$a^* = \operatorname{argmax}_{a_i \in \mathcal{A}} S(Q, a_i) \quad \text{за умови} \quad S(Q, a_i) \geq \tau.$$

Тобто обирається агент з максимальною семантичною відповідністю запиту, який водночас задовольняє вимоги порогової впевненості та обмеження продуктивності.

Такий підхід поєднує:

- ймовірнісне семантичне розуміння від LLM (через векторні представлення запитів),

- детерміновані метрики продуктивності мікросервісів (час обробки, поріг τ , обмеження реального часу).

У результаті формується динамічний пайплайн обробки, де система автоматично підбирає оптимального агента для кожного типу медичних даних: зображень, часових рядів, клінічних текстів тощо.

Тобто система визначає агента з найвищою семантичною відповідністю контексту задачі. В результаті цей підхід поєднує ймовірнісну модель розуміння природної мови (LLM) з детермінованими метриками продуктивності мікросервісів. Це дає можливість формувати динамічний пайплайн обробки, у якому для кожного типу медичних даних (візуальні, часові ряди) автоматично обирається найбільш компетентний інструмент діагностики.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Для реалізації поставленого завдання мною було розроблено хмарну мікросервісну архітектуру з використанням мови програмування Python, платформи AWS (Amazon Web Services) та технології контейнеризації Docker. Система складається з декількох ізольованих компонентів, взаємодія між якими зображена на рисунку 1.

Ключові розроблені сервіси:

- **Orchestrator Service** – центральний керуючий модуль (на схемі "Оркестратор"), реалізований на базі LLM. Його функція полягає у семантичному аналізі вхідного промпту лікаря або метаданих файлу, визначенні типу патології та маршрутизації запиту до відповідного профільного агента.
- **Tumor Detection Agent** – спеціалізований контейнер для обробки візуальних даних (навчання на датасеті ROCov2). Він отримує зображення, проводить інференс моделі Computer Vision та повертає діагноз/висновок.
- **Seizure Detection Agent** – агент аналізу часових рядів (навчання на TUN EEG Seizure Corpus). Працює в режимі, наближеному до реального часу, обробляючи потік даних ЕЕГ для виявлення епілептичних судом.
- **Stream Simulation Service** – модуль імітації стріму ("імітація стріму" на схемі), який зчитує збережені дані з Amazon S3 та порційно подає їх на вхід системи, емулюючи роботу реального медичного обладнання.

Як видно з архітектурної схеми (рис. 1), потік даних починається з користувацького інтерфейсу (Browser UI) або системного дашборду. Оркестратор виступає в ролі інтелектуального шлюзу, що дозволяє уникнути завантаження всіх моделей у пам'ять одночасно, активуючи (рекрутуючи) лише необхідні ресурси в AWS ECS.

Результати діагностики та системні логи зберігаються у реляційній базі даних Amazon RDS (для структурованої інформації про сесію) та сховищі об'єктів Amazon S3 (для сирих медичних даних), що забезпечує надійність зберігання історії хвороби та аудиту дій системи.

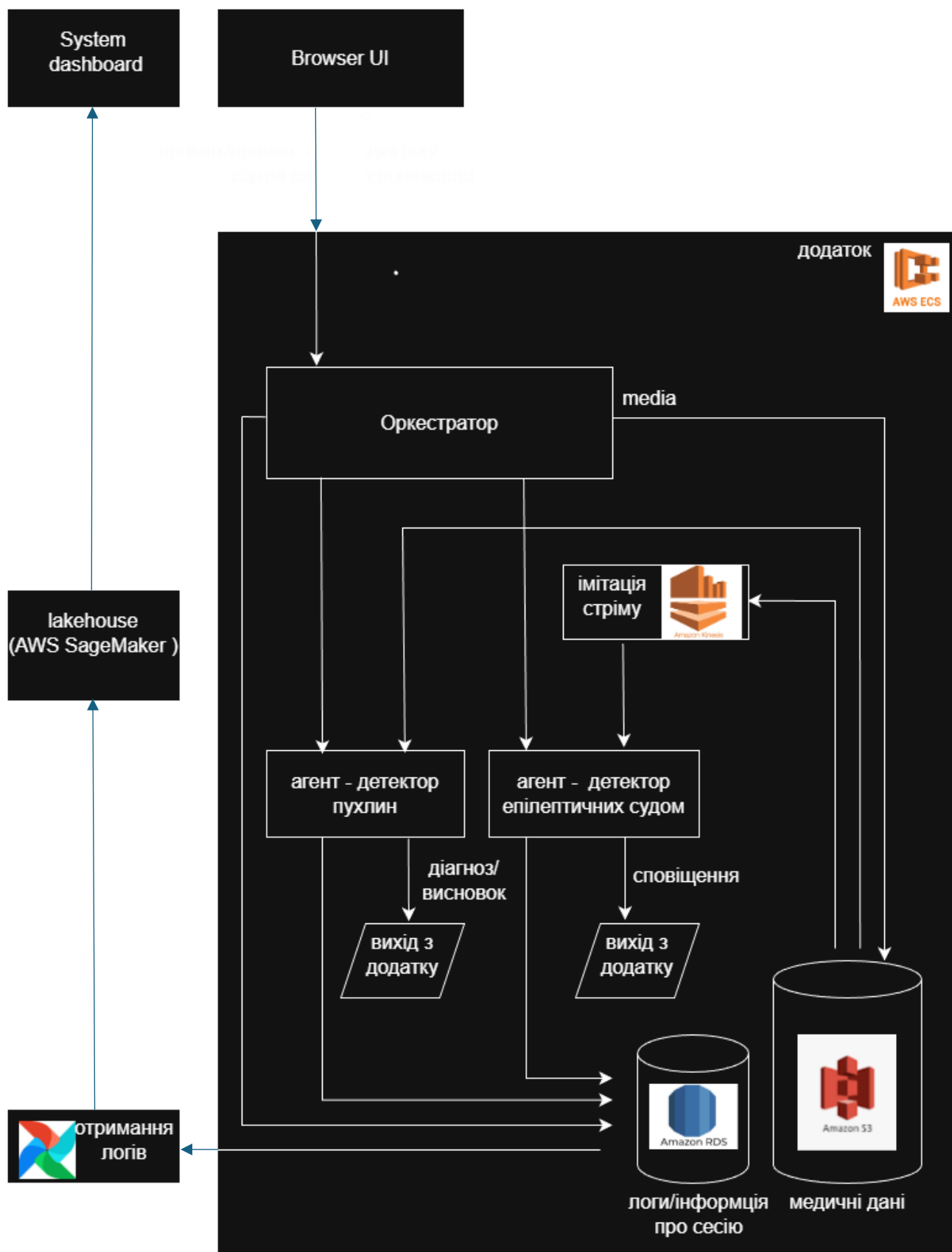


Рисунок 1. Архітектура розробленого застосунку динамічної композиції на базі AWS

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Obeid I., Picone J. The Temple University Hospital EEG Data Corpus. *Frontiers in Neuroscience*. 2016. Vol. 10. URL: <https://www.frontiersin.org/articles/10.3389/fnins.2016.00196/full> (датасет для аналізу судомної активності).
2. Pelka O., Koitka S., Rückert J., Nensa F., Friedrich C.M. Radiology Objects in COntext (ROCO): A Multimodal Image Dataset. *Proceedings of the MICCAI Workshop on Large-scale Annotation of Biomedical Data and Expert Label Synthesis*. 2018. URL: <https://github.com/razorx89/roco-dataset> (датасет для детекції пухлин).
3. Shen Y. et al. HuggingGPT: Solving AI Tasks with ChatGPT and its Friends in Hugging Face. *arXiv preprint arXiv:2303.17580*. 2023. URL: <https://arxiv.org/abs/2303.17580> (теоретична основа для методу рекрутингу агентів).
4. AWS Whitepaper. Machine Learning Lens: AWS Well-Architected Framework. Amazon Web Services, 2023. URL: <https://docs.aws.amazon.com/wellarchitected/latest/machine-learning-lens/machine-learning-lens.html> (архітектурні патерни реалізації).
5. Topol E. J. High-performance medicine: the convergence of human and artificial intelligence. *Nature Medicine*. 2019. Vol. 25. P. 44–56. URL: <https://www.nature.com/articles/s41591-018-0300-7> (актуальність мультимодальної діагностики).

МЕТОДОЛОГІЯ СТВОРЕННЯ ДОСТУПНОГО НАВЧАЛЬНОГО КОНТЕНТУ ДЛЯ ДИСТАНЦІЙНОГО НАВЧАННЯ

Сенчук І.О.¹, Кисельов Г.Д.²

Національний технічний університет України
«Київський політехнічний інститут ім. Ігоря Сікорського», Київ, Україна

¹ senchuk.ilya@lil.kpi.ua, ² g.kyselov@gmail.com

Дослідження зосереджено на існуючих стандартах доступності й розробці цілісної методології та архітектури системи для можливого створення доступного навчального контенту в умовах дистанційного навчання. Враховуючи те, що освітні платформи активно працюють над альтернативними форматами контенту для інклюзивності дизайну своїх продуктів, а ручне створення таких матеріалів є дуже трудовитратним, то завдання розробки інструментів автоматизації є критично важливим

Ключові слова: доступність, інклюзивний дизайн, WCAG, ISO, III.

1. ВСТУП

В епоху цифрової трансформації дистанційна освіта перетворилася з нішевої альтернативи на фундаментальний компонент сучасної освітньої парадигми. Проте, зі збільшенням обсягів та різноманітності цифрового контенту, гостро постає проблема його недоступності для людей з постійними, тимчасовими чи ситуативними обмеженнями, що створює значні бар'єри на шляху до здобуття знань.

Це можуть бути порушення зору, слуху, моторики, когнітивні особливості (постійні обмеження), а також такі стани, як зламана рука (тимчасове обмеження) або перебування в шумному середовищі, де неможливо прослухати аудіо (ситуативне обмеження). Для таких користувачів стандартний цифровий контент – текст без належної розмітки, зображення без альтернативного опису, відео без субтитрів чи аудіодескрипції – стає непереборним бар'єром на шляху до здобуття якісної освіти.

В такому середовищі, тема розробки спеціалізованої методології для створення доступного навчального контенту набуває особливої актуальності. Така методологія не тільки сприяє забезпеченню рівних прав на освіту, але й підвищує загальну якість та ефективність навчального процесу для всіх студентів. Із допомогою автоматизованих інструментів, що є частиною цієї методології, можливо значно спростити та оптимізувати процеси адаптації контенту, враховуючи специфічні потреби різних груп користувачів.

Метою даного дослідження є аналіз та наукове обґрунтування методології створення доступного навчального контенту для систем дистанційного навчання. Це дослідження спрямоване на системне вирішення проблеми цифрової нерівності в дистанційній освіті та, відповідно, підвищення загальної інклюзивності та якості електронного навчання.

2. ІСНУЮЧІ МЕТОДОЛОГІЇ ДЛЯ ІНКЛЮЗИВНОГО ДИЗАЙНУ

2.1. Інклюзивний дизайн

Сучасний підхід до розробки цифрових продуктів, особливо в освітній галузі, вимагає глибокого усвідомлення потреб користувачів. Історично цей акцент зазнав значної еволюції: від виключно технічних стандартів доступності до людиноцентричної філософії інклюзії.

Доступність – це мета, то інклюзивний дизайн – це процес її досягнення.

Загалом, інклюзивний дизайн розглядає такі обмеження, як інвалідність, не як медичний стан окремої людини, а як певну невідповідність між потребами користувача та створеним продуктом, середовищем чи соціальною структурою. Для визначення цього поняття існує термін *mismatched interaction*. Компанія Microsoft у своїх обширних дослідженнях популяризувала модель, яка розглядає цю невідповідність у трьох контекстах:

- Постійні обмеження: довготривалі (або постійні) стани, які традиційно асоціюються з видами інвалідності.
- Тимчасові обмеження – це більш короткотривалі стани, але все одно ті, що впливають на можливості людини взаємодіяти з сервісом.
- Ситуативні обмеження. Або ті, що викликані контекстом, в якому зараз перебуває людина.

Ця модель є надзвичайно потужною і прикладною. Вона описує всі типи обмежень, які кожен з нас протягом життя може відчувати на собі. Розробляючи ІТ продукт для людини з постійним обмеженням, ми автоматично забезпечуємо зручність у користуванні для “нижчих” рівнів, тобто для людей у тимчасових та ситуативних контекстах.

2.2. Сучасний стан освітніх платформ

У сучасній ІТ-сфері є багато позитивних прикладів, таких як Coursera чи edX. Великі МООС-платформи (Massive Open Online Courses) інвестують значні ресурси в доступність, оскільки їхня аудиторія є глобальною та різноманітною і вони не можуть собі дозволити втрачати великі її частини через недостатню інклюзивність.

Проте, в той самий час, на жаль, багато освітніх ресурсів, створених локально (такі як університетські LMS), страждають від системних проблем з доступністю.

2.3. Порівняння актуальних стандартів

Першим розглянемо концепт Web Content Accessibility Guidelines. В цьому стандарті, під кожним із його принципів є чіткі гайдлайни, а під кожною настановою – відповідно *success criteria*, за перевіркою яких можна перевіряти успішність реалізації.

Простими словами, WCAG покликає усунути бар'єри для людей з інвалідністю, а не адаптувати продукт до мінливих контекстів користувача. Він не містить рекомендацій або вимог, що напряму пов'язані з ситуативними потребами (наприклад, зручність UI/UX при користуванні додатком однією рукою). Основа цієї методології – технічна відповідність до конкретних рішень, а не контекстуальна гнучкість. Для покриття більшого кола сценаріїв користувача, принципи WCAG можна доповнювати альтернативними методиками.

Наступним набором стандартів є ISO. Це передусім міжнародна організація зі стандартизації. Вона має власну методологію, в якій доступність визначається у якості ключового аспекту не тільки власне процесу розробки, але й як ядро ергономіки всієї архітектури. Ці стандарти не можна описати як “чек-лист”. Вони спроектовані як сучасні методології менеджменту та пропонують радше загальні філософії та границі для життєвого циклу сервісу, що забезпечать достатню інклюзивність на будь-якому етапі.

Із них найбільш релевантними є два: ISO 9241-171:2008 і ISO 9241-171:2008. Назва першого стандарту звучить як “Ergonomics of human-system interaction – Guidance on software accessibility. Він є частиною великої серії ISO 9241. Ідея серії загалом описана у першій частині назви: ергономіка взаємодії людини з системою. Друга частина назви описує загальну мету – надати ідеї для проєктування ПЗ, що будуть враховувати максимально-можливу зручність використання. Другий же прийшов на зміну застарілому BS 8878. Його основний посил: “Доступність – це пряма відповідальність всієї організації, а не просто функціонал продукту, реалізований дизайнером чи розробником”. В цій перевазі загальності над конкретикою і є його відмінність від інших стандартів

І наостанок варто згадати Європейський стандарт EN 301 549. Він, на відміну від попередніх, інтегрує принципи загальної доступності на рівні цілого законодавства і глобальних правил ведення бізнес-діяльності. Він є одним із ключових документів Європейського союзу та одним із важелів, що забезпечують надійність загального цифрового ринку.

Основне призначення EN 301 549 – це встановлення функціональних вимог до ICT, стандартів, яким має відповідати створений продукт. Його застосовують при закупівлі такого цифрового продукту чи сервісу у будь-якої IT-компанії державною установою в ЄС.

3. ВЛАСНІ ІНСТРУМЕНТИ ДЛЯ ОБРОБКИ НАВЧАЛЬНОГО КОНТЕНТУ

3.1. Концептуальна архітектура

Ручне створення alt-текстів, субтитрів, аудіоописів, та інших форматів альтернативного контенту – це трудомісткий, дорогий і важко масштабований процес. Ця причина виступає бар’єром не тільки для завдання цього дослідження, а й для масового впровадження у комерційну розробку. Дещо спростити задачу може розробка модульної системи, автоматизованої за допомогою ШІ технологій.

Традиційний монолітний підхід не підходить для таких задач, адже при його використанні будь-яке оновлення окремої ШІ-моделі вимагатиме перевірки й, можливо, переписання всієї кодової бази та повторного деплою системи. Ці проблеми може вирішити мікросервісний підхід до проєктування. Тобто розбиття архітектури на маленькі незалежні сервіси, що відповідають за конкретні бізнес-процеси чи функції.

Для мікросервісної архітектури найкращою буде реалізація через асинхронну взаємодію компонентів з допомогою черги повідомлень, і використання API Gateway.

Для забезпечення дружнього та легкого дизайну для користувача в першу чергу необхідно чітко ідентифікувати ключового користувача системи. Звісно, кінцевим споживачем навчального контенту завжди буде студент, проте в подавляючій більшості випадків основними функціями власне системи буде користуватись викладач, чи інший менеджер навчального контенту.

Проаналізувавши такий профіль користувача, можна сформулювати основну задачу інструменту, як “асистування, спрощення й прискорення замість перешкоджання”. Тобто, простими словами, система повинна:

1. Визначати автоматично очевидні чи прості рішення, такі як мову відео.
2. Виконувати 90% роботи автоматично, але завжди лишати простір для верифікації й редагування, та надавати для цього зрозумілий інтерфейс.
3. Пропонувати додаткові покращення (але не виконувати їх за своїм рішенням) в доповнення до чіткого слідування інструкції.

Дуже важливим є критерій доступності самої системи створення доступного контенту, так звана “мета-доступність”. Якщо інтерфейс користування самою системою буде бар’єром

для таких користувачів, то вона не буде відповідати основним принципам інклюзії. Адже викладачі, такі як створена для прикладу репрезентативна персона, також можуть мати постійні, тимчасові або ситуативні обмеження.

3.2. Робочі алгоритми

Основне завдання першого мікросервісу (Speech-to-Text) – перетворити навчальну лекцію, повну аудіодоріжок на точний, правильний і синхронізований за часом текстовий файл з субтитрами.

Сучасний стан галузі розпізнавання мовлення визначається моделями на базі архітектури Transformer (Encoder-Decoder). Для даної методології оберемо як референс OpenAI Whisper.

Другий мікросервіс (Computer Vision) має використовувати алгоритм Image-to-Text, що на основі аналізу зображення генерує змістовний текст опису. В освітньому контексті це завдання є значно складнішим, ніж звичайний опис фото чи малюнків, оскільки сервіс має не просто описати що саме зображено, а передати значення абстрактного наповнення, у вигляді сенсу графіку, висновку зі схеми чи формули, та інше.

Традиційні підходи до детекції об'єктів не можуть впоратися з цим завданням, адже “графік зростання” – це не просто “лінія” + “текст”. На противагу, сучасні VLM уособлюють мультимодальний підхід і мають не тільки “Очі”, у вигляді Vision Transformer, а й “Мозок”, роль якого зазвичай виконує LLM. Тобто після перетворення зображення на набір числових векторів, з описом наповнення і контекстної конотації, вони генерують зв'язний (частіше за все) текст, зрозумілий людині.

І третій основний мікросервіс (Audio Description) надаватиме звуковий опис ключових візуальних подій для незрячих або слабозорих глядачів. Для цього він ідентифікує паузи чи візуальні підказки й описує важливі деталі, щоби студенти мали змогу повністю розуміти потік знань. Цей безумовно композитний і складний процес має залучати всі вже перераховані технології (CV, STT) і додаткові для генерації опису та озвучення за цим описом (NLP, TTS).

4. ОЦІНКА І ВПРОВАДЖЕННЯ СИСТЕМИ

Для підтвердження життєздатності запропонованої методології необхідна чітка система оцінювання її ефективності. Тож потрібно розробити комплексну методику оцінювання. Ми не можемо обмежитись бінарним “працює/не працює”. В процесі аналізу, було складено три ключові компоненти оцінювання [18]:

1. Технічна точність:
 - a. Що оцінює: об'єктивну точність роботи ШІ-алгоритмів.
 - b. Навіщо: базовий рівень перевірки, при відповідності якому викладачу не доведеться витратити багато часу на перевірку й корекцію.
2. Практична ефективність:
 - a. Що оцінює: зручність, інтуїтивність, зекономлений час.
 - b. Навіщо: простий, швидкий і прямолінійний інтерфейс є однаково важливим фактором із технічною досконалістю рішення.
3. Фінальна якість доступності:
 - a. Що оцінює: якість контенту генерації ШІ та перевірки викладачем.
 - b. Навіщо: перевірка відповідності кінцевого результату стандартам WCAG чи іншим і чи насправді це співвідносно користі й зрозумілості контенту для студентів з постійними, тимчасовими та ситуативними обмеженнями.

Проте, для повної реалізації методології не вистачить тільки впровадити технічний інструмент – це лише частина. Щоб досягти справжньої інклюзивності, рекомендується інтегрувати принципи доступності безпосередньо у процес створення та викладання

навчального матеріалу. Особливо, враховуючи те, що, як було зазначено при аналізі стандартів, це спростить сприйняття контенту і для звичайних студентів також.

Тобто враховуючи хоча б частково умови доступності під час створення контенту, а не після, викладач може значно полегшити роботу ШІ-алгоритмів та свій етап верифікації і в той самий час покращити кінцевий результат. Викладач має вирішувати хоча би частину найпростіших проблем ще під час створення контенту будь-якого типу.

5. ОЦІНКА І ВПРОВАДЖЕННЯ СИСТЕМИ

У ході дослідження й розробки цілісної методології створення доступного контенту для дистанційного навчання було проведено аналіз предметної області та виявлено методологічну прогалину, вирішити яку й покликана така методологія. Також, аби виявити можливі кроки було досліджено принципи інклюзивного дизайну та уточнено модель постійних, тимчасових та ситуативних обмежень. Після цього для розглянутих стандартів (WCAG 2.1, ISO/IEC, EN 301 549) було проведено порівняльний аналіз. Його результатом і став висновок про необхідність створення автоматизованого інструментарію.

Власне було спроектовано модель сервісу, що зможе стати частиною рішення і деталізовано алгоритми обробки даних для трьох ключових ШІ-сервісів.

Для імплементації такої системи було запропоновано трикомпонентний план верифікації ефективності: оцінка якості ШІ-артефактів, ефективності юзабіліті та доступності.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Вікіпедія: Універсальний дизайн. URL: https://uk.wikipedia.org/wiki/Універсальний_дизайн (дата звернення 20.09.2025)
2. Microsoft Inclusive Design Toolkit URL: <https://inclusive.microsoft.design> (дата звернення 20.09.2025)
3. Google Scholar стаття “The State of MOOC Accessibility”.
4. International Organization for Standardization: ISO 9241-171:2008 Ergonomics of human-system interaction URL: <https://www.iso.org/standard/39080.html> (дата звернення 29.09.2025)
5. ETSI: EN 301 549 V3.2.1 (2021-03) Accessibility requirements for ICT products and services. URL: https://www.etsi.org/deliver/etsi_en/301500_301599/301549/03.02.01_60/en_301549v030201p.pdf (дата звернення 30.09.2025)
6. W3C WAI: Comparing EN 301 549 and WCAG URL: <https://www.w3.org/WAI/policies/international/standards/en-301-549/> (дата звернення 30.09.2025)

ПІДВИЩЕННЯ СИТУАЦІЙНОЇ ОБІЗНАНОСТІ РОЮ ДРОНІВ ШЛЯХОМ КОНТЕКСТНОГО АНАЛІЗУ СТИГМЕРГІЇ

Стернійчук С.Ю.¹, Булах Б.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ sterniichuk.serhii@gmail.com, ² bogdan.bulakh@gmail.com

У роботі розглянуто підхід до підвищення ситуаційної обізнаності рою БПЛА на основі контекстного аналізу стигмергійних слідів. Запропоновано модель віртуальних феромонів із адаптивним коефіцієнтом випаровування, що реагує на динамічні та статичні загрози. У середовищі Gazebo проведено моделювання на основі трьох БПЛА (3DR Iris) з використанням сенсорів LiDAR/камера та FANET-зв'язку. Результати демонструють скорочення часу місії та зменшення енергоспоживання у сценаріях рухомих загроз і шумних даних. Наукова новизна полягає у контекстній інтерпретації феромонної карти. Практична значимість – у підвищенні автономності та надійності рою дронів у складних умовах.

Ключові слова: рой дронів, стигмергія, віртуальні феромони, ситуаційна обізнаність, Gazebo.

1. ВСТУП

Рої безпілотних літальних апаратів (БПЛА) демонструють високу ефективність у задачах розвідки, картографування, моніторингу та пошукових операцій. Одним із ключових підходів до самоорганізації таких систем є стигмергія – механізм непрямої координації через середовище, де агенти залишають сліди (феромони), що впливають на поведінку інших агентів.

У класичному вигляді стигмергія передбачає лише реакцію на феромон, проте для автономних роїв цього недостатньо. У реальних середовищах присутні:

- динамічні загрози, що змінюють положення у просторі;
- сумнівні або шумні сенсорні дані;
- накладення застарілих феромонів, що призводять до деградації карти.

Тому ключовим завданням стає підвищення ситуаційної обізнаності – здатності рою розуміти контекст появи стигмергійних слідів і адаптувати поведінку.

У цій роботі розглянуто підхід контекстного аналізу, побудований на патерн-рекогнішн феромонної карти, оцінці часової динаміки та визначенні рухомості загроз.

2. МЕТОД І КОМУНІКАЦІЙНА МОДЕЛЬ

2.1. Віртуальні феромони

Використано підхід *virtual pheromones*, у якому феромони представлені як цифрова карта, що оновлюється окремими агентами рою. Кожна мітка містить значення інтенсивності та параметр довіри, який зменшується з часом відповідно до коефіцієнта випаровування. Адаптивне випаровування застосовується з урахуванням контексту мітки, що дає змогу швидше видаляти недостовірні або застарілі феромонні сліди. Такий підхід дозволяє формувати узгоджену інформаційну карту середовища, яка надалі використовується в алгоритмах ситуаційної обізнаності.

2.2. Комунікація у FANET та синхронізація карти

Зв'язок між дронами реалізовано через FANET з використанням UDP та протоколу MAVLink. Передаються не повні карти, а лише *події*: оновлення локальних міток, підтвердження загроз, зміни довіри. У випадку втрати зв'язку застосовується опортуністична дельта-синхронізація, коли дрон після повернення отримує лише зміни карти. Така модель мінімізує трафік та дозволяє працювати з обмеженим каналом.

2.3. Патерн-рекогнішн феромонної карти

Контекстна інтерпретація стигмергійних слідів у рої БПЛА базується на аналізі просторової та часової динаміки феромонної карти. Для кожної феромонної області обчислюється її центр мас, що дозволяє оцінювати зміщення мітки та відокремлювати статичні загрози від рухомих. Центр мас визначається за формулою:

$$C_t = \left(\frac{\sum x_i \tau_i}{\sum \tau_i}, \frac{\sum y_i \tau_i}{\sum \tau_i} \right).$$

Порівняння значень у два моменти часу дає вектор переміщення феромонної плями:

$$v = \frac{C_{t_1} - C_{t_0}}{t_1 - t_0}.$$

Якщо $\text{norm } v < \epsilon$, мітка вважається статичною; при $\text{norm } v \geq \epsilon$ фіксується динамічна загроза. Такий підхід забезпечує виявлення руху без застосування важких ML-моделей.

Крім зміщення, аналізується форма розподілу феромонів $\tau(x, y)$. Статичні загрози формують компактний кластер з низькою ексцентриситетністю, тоді як рухомі створюють витягнуту структуру, орієнтовану вздовж напрямку руху. Для оцінки використовується метод головних компонент (PCA): перша компонента визначає домінуючий напрямок, а співвідношення власних значень $k = \lambda_1 / \lambda_2$ відображає ступінь витягнутості (при $k \approx 1$ кластер статичний, при $k \gg 1$ – характерний «слід» рухомої загрози). Цей критерій дозволяє коректно ідентифікувати динамічні об'єкти навіть при незначному зміщенні центру.

Оскільки FANET-зв'язок може бути нестабільним, частина феромонних міток може бути зашумленою або отриманою лише одним дроном. Тому вводиться параметр довіри conf , що залежить від кількості підтверджень іншими агентами та часу останньої перевірки. Формально:

$$\text{conf} = \alpha * N_{\text{confirm}} - \beta * (t - t_{\text{last}}).$$

Коли $\text{conf} < \text{conf_min}$, мітка вважається недостовірною, і для неї збільшується коефіцієнт випаровування, що забезпечує її швидке видалення. Це критично у випадках пошкоджених сенсорів, тимчасових артефактів камери та одиничних хибних детекцій.

Для рухомих загроз застосовується адаптивне випаровування. Коефіцієнт ρ визначається залежно від динамічності:

$$\rho = \rho_{\text{dynamic}} \text{ if } \text{dynamic} = 1, \rho_{\text{static}} \text{ if } \text{dynamic} = 0,$$

де $\rho_{\text{static}} = 0.02 \dots 0.1$, $\rho_{\text{dynamic}} = 0.3 \dots 0.7$. Таким чином, система «забуває» рухомі мітки значно швидше, перешкоджаючи накопиченню застарілої або хибної інформації, що знижує кількість неправильних обходів та оптимізує маршрути.

У сукупності ці механізми забезпечують формування актуальної та достовірної феромонної карти, що на пряму підвищує ситуаційну обізнаність рою, якість прийняття рішень та ефективність виконання місії.

3. СЕРЕДОВИЩЕ МОДЕЛЮВАННЯ

Моделювання виконано в Gazebo Harmonic (рис. 1) на основі міської мапи Берліна (OpenStreetMap).

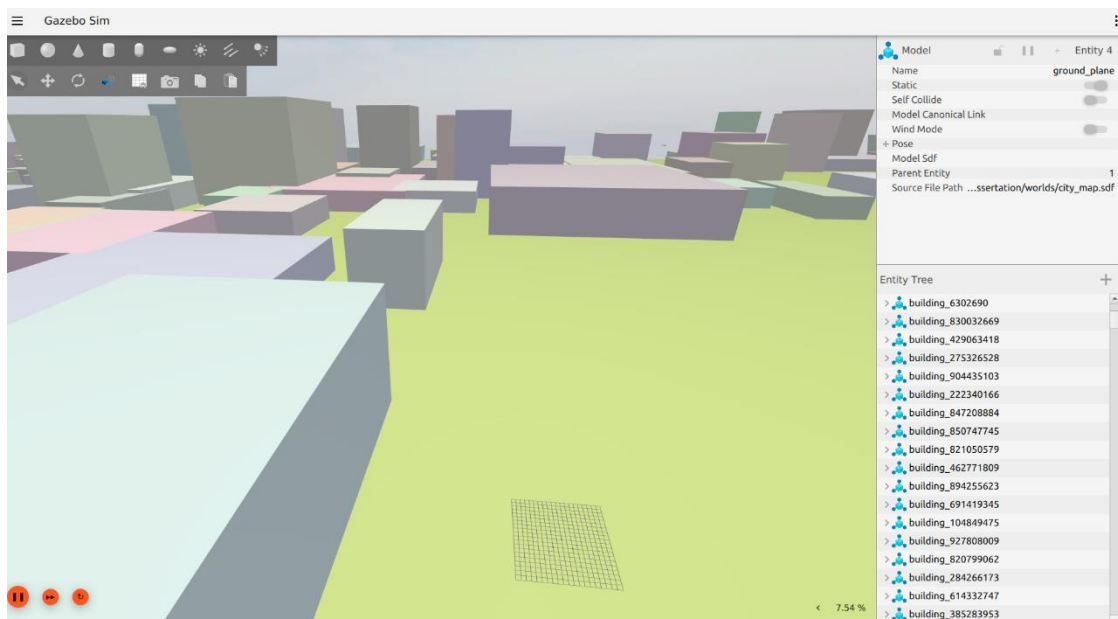


Рисунок 1. Gazebo simulation

Використано три БПЛА моделі 3DR Iris з сенсорами LiDAR та камерою (рис. 2). FANET забезпечує зв'язок до 500 м.

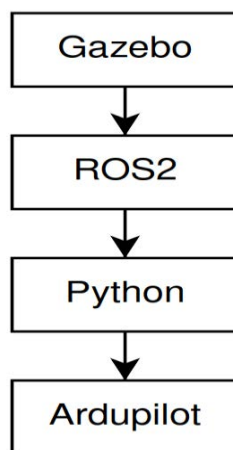


Рисунок 2. Ланцюжок передачі даних від симулятора до автопілота

4. РЕЗУЛЬТАТИ

Оцінено три сценарії: рухома загроза, пошкоджений сенсор та засмічення карти феромонами. Порівняно базову стигмергію та запропонований контекстний аналіз. Результати подано у відносному вигляді (% до гіршого методу).

Таблиця 1. Результати

Сценарій	Підхід	Час місії	Заряд
Рухома загроза	Базовий	0 %	0 %
	Контекстний	−15 %	+8 %
Пошкоджений сенсор	Базовий	0 %	0 %
	Контекстний	−7 %	+3 %
Засмічення карти	Базовий	0 %	0 %
	Контекстний	−9 %	+4 %

Контекстний аналіз скорочує час місії на 7–15 % та підвищує залишковий заряд на 3–8 %, що пов'язано з уникненням зайвих обходів і видаленням недостовірних міток.

5. ВИСНОВКИ

Запропонований підхід до ситуаційної обізнаності рою БПЛА дозволяє адаптивно інтерпретувати стигмергійні мітки на основі їх просторово-часових патернів. Визначення рухомості загроз, класифікація форм кластера та оцінка достовірності забезпечують побудову актуальної феромонної карти, що напряду підвищує ефективність рою. Моделювання підтверджує скорочення часу місії та зменшення енергоспоживання порівняно з класичною стигмергією. Запропонований метод може бути інтегрований у системи координації БПЛА для складних урбаністичних середовищ.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Dorigo M., Birattari M. Ant Colony Optimization. Springer, 2010.
2. Payton D. Digital Pheromones in Autonomous Swarms. AAAI, 2001.
3. OpenStreetMap Project
4. Gazebo Harmonic Documentation, Open Robotics, 2024.
5. Swarm Robotics: Principles and Applications. MIT Press, 2021.
6. ArduPilot Developer Documentation, 2024.
7. MAVLink Protocol Specifications. mavlink.io, 2024.

РЕКОМЕНДАЦІЙНА СИСТЕМА ДЛЯ ВИБОРУ СТАЛОЇ СЕРВЕРНОЇ АРХІТЕКТУРИ ВЕБ-ДОДАТКУ

Товкач М.А.¹, Гіоргізова-Гай В.Ш.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ tovkach.maxim@lil.kpi.ua [0009-0001-9389-938X], ² high.victoria@lil.kpi.ua

Мета дослідження полягає у розробці підходу до вибору найбільш доцільної серверної архітектури веб-додатку з урахуванням принципів сталого розвитку. **Методи дослідження** ґрунтуються на порівняльному аналізі існуючих рекомендаційних систем, сучасних архітектурних підходів та методі багатокритеріального оцінювання з використанням адаптивних коригувальних коефіцієнтів. **Основні результати:** спроектовано алгоритм рекомендаційної системи, яка на основі параметрів проекту та пріоритетів користувача визначає доцільну серверну архітектуру. Практична значимість виявляється у зниженні ризиків розробки та забезпеченні екологічної сталості ІТ-продуктів.

Ключові слова: серверна архітектура, рекомендаційна система, вуглецевий слід, монолітна архітектура, мікросервіси, безсерверна архітектура

1. ВСТУП

Архітектура програмного забезпечення є фундаментом будь-якої ІТ-системи – так само, як міцний фундамент визначає надійність будинку. Правильні рішення на ранніх етапах проекту дозволяють уникнути дорогих та часовитратних змін у майбутньому. Від вірного вибору архітектурного підходу залежить не лише технічна ефективність системи, а й економічна успішність проекту. З огляду на різноманіття сучасних технологій та складність пошуку балансу між суперечливими вимогами до системи, постає потреба у спеціалізованому інструменті, який би допоміг архітектору прийняти зважене та обґрунтоване рішення.

Існуючі системи підтримки прийняття архітектурних рішень, такі як «AWS Well-Architected Tool» [1] та «Archify: A Recommender System of Architectural Design Decisions» [2], вирішують проблему лише частково. Інструмент AWS фокусується на аудиті та вдосконаленні вже спроектованих систем. Користувач заповнює анкету з питаннями щодо свого проекту й надає пріоритети його критеріям, в результаті чого інструмент формує список можливих ризиків, проте не допомагає здійснити первинний вибір архітектурного стилю. На противагу цьому, система «Archify» на основі відповідей користувача стосовно вимог й характеристик майбутнього додатку надає рекомендації щодо використання архітектурних стилів, тактик та технологій реалізації проекту на основі бази знань, сформованої з технічних джерел. Проте, ця система має низку обмежень – вона надає список альтернатив замість єдиної чіткої відповіді, не дозволяє врахувати пріоритети користувача щодо архітектурних критеріїв, не містить в рекомендаціях ряд сучасних архітектурних підходів як безсерверна архітектура, edge, і не враховує екологічний аспект, який є важливим чинником для сучасної ІТ-індустрії.

Стрімкий розвиток цифрових технологій та хмарних обчислень призвів до того, що ІТ-сектор став одним із найбільших споживачів електроенергії у світі, через що ІТ-компанії все

частіше приділяють увагу мінімізації вуглецевого сліду своїх продуктів. Основна частка вуглецевого сліду програмного забезпечення формується непрямо – через енергоспоживання центрів обробки даних, значна частина яких досі живиться від джерел на викопному паливі [3]. Тому оптимізація архітектурних рішень та вибір регіону розгортання з низьким вуглецевим коефіцієнтом є критичними факторами для забезпечення сталості ІТ-продуктів.

2. АНАЛІЗ СУЧАСНИХ АРХІТЕКТУРНИХ ПІДХОДІВ ДО ПОБУДОВИ СЕРВЕРНОЇ ЧАСТИНИ ВЕБ-ДОДАТКУ

Монолітна архітектура характеризується високою швидкістю старту та операційною простотою завдяки розгортанню у вигляді єдиного артефакту, а також спрощує забезпечення узгодженості даних завдяки використанню єдиної бази даних й забезпечує мінімальну затримку, оскільки комунікація компонентів відбувається через виклики в пам'яті [4]. Проте, цей підхід має суттєві обмеження у масштабованості через неможливість незалежного масштабування компонентів, що призводить до неефективного використання ресурсів, а також у надійності, адже помилка в одному модулі може спричинити повну зупинку системи. Через тісний взаємозв'язок компонентів зростання масштабу монолітної системи ускладнює її підтримку, тому задля покращення зручності супроводу коду обирають модульний моноліт, який хоча і спрощує підтримку, проте зберігає основні обмеження класичного підходу.

На сьогодні одним із провідних напрямів у сфері розробки програмного забезпечення є мікросервісна архітектура, яка передбачає побудову додатку як набору малих автономних компонентів, кожен з яких відповідає за окремий бізнес-домен, має власне сховище даних та взаємодіє з іншими через легкі протоколи комунікації або системи обміну повідомленнями. Слабка зв'язність компонентів забезпечує високу гнучкість та дозволяє незалежно масштабувати окремі сервіси відповідно до їхнього навантаження, що підвищує ефективність використання ресурсів [5]. Крім того, архітектура підвищує відмовостійкість, адже ізоляція компонентів дозволяє системі зберігати часткову доступність при збої одного з компонентів. Також, завдяки можливості паралельної розробки, використанню різного технологічного стеку та незалежного розгортання компонентів пришвидшується випуск нових версій продукту. Водночас, перехід до розподіленої системи суттєво підвищує загальну складність системи, адже проектування, тестування й моніторинг вимагають значно більших зусиль порівняно з монолітом [4]. Архітектура висуває підвищені вимоги до інфраструктури та технічної кваліфікації команди, що значно підвищує початкові інвестиції та витрати на подальшу підтримку системи. Також, децентралізація даних ускладнює забезпечення їх узгодженості, вимагаючи реалізації механізмів для забезпечення розподілених транзакцій, а мережева взаємодія між сервісами призводить до збільшення часу відгуку системи.

Безсерверна архітектура (serverless) дозволяє розробникам сфокусуватися на написанні бізнес-логіки, переклавши управління інфраструктурою, зокрема виділення ресурсів, забезпечення відмовостійкості та масштабування на хмарного провайдера [4]. Function-as-a-Service (FaaS), забезпечує економічну ефективність при нерівномірному навантаженні завдяки механізму оплати за фактично використані обчислювальні ресурси та автоматичному масштабуванню до нуля за відсутності навантаження. Запуск функції відбувається у відповідь на подію, а її виконання відбувається в тимчасових, ізольованих середовищах виконання, життєвим циклом яких керує хмарний провайдер. Підхід забезпечує високу швидкість старту розробки та ефективну масштабованість при пікових й навіть непередбачуваних навантаженнях, але має проблему «холодного старту» (затримка при першому виклику), безстановість (що вимагає використання зовнішніх сховищ), жорсткі ліміти ресурсів та часу виконання, ризик прив'язки до хмарного провайдера та підвищену складність моніторингу й

організації взаємодії великої кількості розподілених компонентів [4]. Нівелювати частину недоліків можна за допомогою серверлес-контейнерів, які поєднують гнучкість контейнерів (контроль над середовищем виконання, портативність, вищі ліміти) з перевагами безсерверного підходу, тому добре підходять для застосунків з середнім часом обробки запиту та нерівномірним навантаженням. Опція утримання мінімальної кількості екземплярів дозволяє усунути затримки «холодного старту» ціною появи витрат на простій. До serverless належить й Backend-as-a-Service (BaaS), який надає розробникам набір готових бекенд сервісів через API, переклавши управління їх інфраструктурою та масштабуванням на провайдера [6]. Це забезпечує швидкий вихід на ринок, але обмежує гнучкість реалізації складної логіки та є економічно неефективним при стабільно високому навантаженні.

Розповсюдженим є гібридний підхід, що поєднує постійно працююче ядро системи (реалізоване як моноліт або мікросервіси) з безсерверними підходом. У такій моделі ядро обробляє основну бізнес-логіку зі стабільним навантаженням, що вимагає низької затримки. Натомість FaaS використовується для виконання допоміжних короткочасних завдань. Це доцільно, коли такі завдання є асинхронними, мають нерівномірне навантаження, або потребують багато ресурсів та можуть заважати чи спричинити збої в роботі ядра системи. Такий гібридний підхід призводить до оптимізації витрат, адже завдяки масштабуванню до нуля оплата не буде здійснюватися за їх простій, а також дозволяє використовувати ширший технологічний стек та уникнути проблеми холодного старту для ядра системи.

Event-Driven Architecture – це архітектурний стиль, у якому основна взаємодія між компонентами системи базується на асинхронному обміні подіями з використанням центрального брокера. Сервіс-виробник публікує подію та миттєво звільняється, не очікуючи реакції підписаних на неї споживачів, що забезпечує слабку зв'язність компонентів та підвищує відмовостійкість системи [7]. Брокер виступає буфером, гарантуючи збереження повідомлень при збоях споживачів та дозволяючи їм обробляти події у власному темпі, що сприяє згладжуванню пікових навантажень. Підхід спрощує масштабування та розширення системи, але й підвищує її операційну складність, висуваючи вимоги до детального проектування взаємодії, моніторингу й забезпечення узгодженості даних.

Граничні обчислення дозволяють виконувати частину логіки не у віддаленому централізованому дата-центрі, а якомога географічно ближче до користувача, що дозволяє досягти мінімальної мережевої затримки [8]. Це реалізується через використання edge-функцій, які автоматично розгортаються на глобальній мережі edge-серверів. Для обробки запитів на найближчому до користувача вузлі миттєво створюються та масштабуються тимчасові ізольовані середовища виконання. Підхід дозволяє виконувати легкі та швидкі операції по типу валідації, а також ефективно кешувати контент без звернення до основного бекенду. Це забезпечує нижчу середню затримку відповіді та суттєво знижує навантаження на центральні сервери, підвищуючи стійкість до пікових навантажень та DDoS-атак. Але, edge-функції є безстановими та мають обмеження обчислювальних ресурсів і часу виконання, тому не підходять для тривалих ресурсоємних завдань. Їх використання ускладнює проектування архітектури, тестування, моніторинг.

3. ПРОЕКТУВАННЯ РЕКОМЕНДАЦІЙНОЇ СИСТЕМИ

Для реалізації рекомендаційної системи для вибору сталої серверної архітектури веб-додатку, було запропоновано підхід, який при побудові рекомендації враховує як специфіку та потреби проекту користувача, так і фундаментальні властивості кожного з розглянутих архітектурних підходів. Процес взаємодії з системою передбачає заповнення анкети, де користувач визначає ключові параметри свого майбутнього проекту, такі як розмір технічної

команди, характер очікуваного навантаження та рівень складності предметної області. Також користувач визначає пріоритетність архітектурних критеріїв, вказуючи які характеристики системи є для нього найбільш значущими. Вхідними даними для алгоритму є відповіді користувача, встановлені ним вагові коефіцієнти пріоритетності, а також базові оцінки архітектур, які були попередньо сформовані на основі технічної літератури. Результатом роботи системи є рекомендація найбільш доцільної серверної архітектури, доповнена аргументацією даного вибору та розглядом компромісів, які користувачу варто врахувати при подальшій роботі над проектом. Окрім цього, система надає рекомендації щодо мінімізації вуглецевого сліду, спрямовані на забезпечення екологічної сталості майбутнього додатку.

На першому етапі було визначено набір ключових архітектурних критеріїв, за якими здійснюється оцінювання та яким користувач надаватиме пріоритет:

- Швидкість виходу на ринок. Критерій враховує часові та фінансові ресурси, необхідні для проектування, реалізації та релізу першої робочої версії продукту. Підвищення свідчить про потребу швидкого запуску й надає перевагу архітектурам з низьким порогом входу.
- Економічна ефективність. Цей критерій характеризує використання коштів на утримання інфраструктури. Він враховує витрати на оренду обчислювальних ресурсів, зберігання даних та оплату за мережеву передачу даних.
- Масштабованість. Характеризує здатність системи ефективно обробляти зростаючу кількість запитів та витримувати пікові навантаження шляхом додавання обчислювальних ресурсів. Підвищення важливості критерію свідчить про очікування значного росту трафіку, що надає перевагу архітектурам, які ефективно масштабуються горизонтально.
- Чутливість до затримки. Критерій характеризує вимогу до часу відгуку системи для кінцевого користувача. Підвищення важливості цього критерію надаватиме перевагу архітектурам без холодних стартів та зайвих мережевих викликів.
- Операційна простота системи. Критерій визначає рівень зусиль, необхідних для забезпечення стабільної роботи системи. Він охоплює складність налаштування й підтримки інфраструктури, організацію процесів розгортання й моніторингу. Його збільшення свідчить про бажання користувача мінімізувати операційну складність.
- Надійність системи. Цей критерій характеризує здатність системи зберігати працездатність та забезпечувати доступність функціоналу для користувачів навіть у випадку збою окремих її компонентів. Збільшення пріоритету цього критерію надає перевагу розподіленим системам, де падіння одного компоненту не призводить до повної зупинки всієї системи.

Таким чином, користувач визначає рівень пріоритетності для кожного з шести критеріїв, окремо оцінюючи їх важливість по шкалі від одного до п'яти. Значення один вказує на низький пріоритет та встановлено за замовчуванням, а значення п'ять відповідає найвищому рівню важливості.

До переліку рекомендованих архітектурних альтернатив увійшли п'ять основних варіантів: модульна монолітна архітектура, мікросервісна архітектура, безсерверна архітектура (на базі FaaS), а також їхні гібридні поєднання (моноліт + FaaS й мікросервіси + FaaS). У деяких випадках, в залежності від пріоритетів та відповідей користувача, система буде також рекомендувати звичайну монолітну архітектуру (наприклад, якщо дуже важливий швидкий вихід на ринок та предметної область є простою). Окрім цього, в залежності від пріоритетів та відповідей користувача, система також буде надавати рекомендації щодо використання додаткових архітектурних патернів та технологій, таких як Edge Computing, Event-Driven Architecture, BaaS та serverless containers.

Задля підвищення точності рекомендацій було вирішено що недоцільно використовувати просте сумування балів у загальний рейтинг архітектури. Натомість, кожна відповідь користувача впливає на оцінку конкретних критеріїв архітектури, не зачіпаючи інші. Таким чином, кожна архітектура має власний набір динамічних оцінок для кожного з шести критеріїв. Як наслідок, відповіді користувача коригують лише ті оцінки архітектури за окремими критеріями, на які вони реально впливають. Наприклад, якщо користувач вказує що розмір технічної команди, залученої до реалізації проекту є малим, то це суттєво знижує оцінку критерію «Операційна простота» для мікросервісів, проте ніяк не впливає на їх оцінку за критерієм «Масштабованість». Базові оцінки відповідності розглянутих архітектурних підходів визначеним критеріям було сформовано на основі проведеного їх аналізу особливостей, переваг та недоліків. Оцінювання здійснювалося за шкалою від одного до десяти, де одиниця вказує на критичну невідповідність вимогам критерію, а десять – на максимальний рівень відповідності критерію. Сформовані експертні оцінки наведено у матриці оцінювання відповідності (табл. 1).

Таблиця 1. Матриця оцінювання відповідності

	Швидкість виходу на ринок	Економічна ефективність	Масштабованість	Чутливість до затримки	Операційна простота	Надійність системи
Модульна монолітна архітектура	9	6	4	9	9	5
Мікросервісна архітектура	3	3	9	6	3	8
Безсерверна архітектура	9	9	10	5	8	7
Монолітна + FaaS	7	7	6	8	7	6
Мікросервіси + FaaS	2	4	10	5	2	9

Оцінки з таблиці 1 будуть використовуватися у якості початкових оцінок критеріїв відповідних архітектур. Кожна відповідь користувача коригує ці оцінки за допомогою модифікаторів, збільшуючи чи зменшуючи їх значення. Таким чином, даний підхід базується на методології систем підтримки прийняття рішень на основі знань (Knowledge-Based DSS) та використовує адаптований метод зваженої суми. У контексті рекомендаційних систем цей метод дозволяє врахувати експертні знання про предметну область (базові оцінки) та динамічно адаптувати їх під контекст користувача (модифікатори) [9]. Задля забезпечення збалансованого впливу відповідей на результати розрахунку було вирішено використовувати шкалу модифікаторів у діапазоні цілих чисел від –2 до +2, де зміна оцінки на 2 бали вказує на суттєвий вплив відповіді, а на 1 бал – на незначний вплив. Також, для деяких відповідей, які є визначальними для вибору або накладають критичні обмеження, застосовується спеціальний модифікатор у 4 бали, що дозволяє надавати таким факторам домінуючу вагу. Таким чином, відповіді користувача виступають як модифікатори, що коригують базові оцінки критеріїв для кожної архітектури під специфіку його проекту. Після відповіді на усі запитання, отримані скореговані оцінки множаться на відповідні вагові коефіцієнти, задані користувачем. Потім, для кожної архітектури здійснюється сумування зважених оцінок усіх критеріїв, що формує її

фінальний рейтинговий бал, використовуючи формулу (1), після чого система рекомендує користувачу архітектуру з найбільшим фінальним балом.

$$A_i = \sum_{k=1}^6 ((B_{ik} + M_{ik}) \cdot Weight_k). \quad (1)$$

У цій формулі: A_i – це фінальний рейтинговий бал i -тої архітектури, де $i \in \{1, \dots, 5\}$; B_{ik} – це базова оцінка i -тої архітектури за критерієм $k \in \{1, \dots, 6\}$, завчасно визначена шляхом експертного оцінювання; M_{ik} – це сума модифікаторів, що відображає вплив відповідей користувача на формування оцінки критерію k для архітектури i ; $Weight_k$ – це ваговий коефіцієнт, тобто пріоритет критерію k встановлений користувачем.

Кожен модифікатор пов'язаний з відповідним текстовим поясненням, що дозволяє автоматично формувати деталізоване обґрунтування результату: фактори, що мали позитивний вплив, формують блок аргументації вибору, а фактори з негативним впливом – блок компромісів, вказуючи потенційні ризики та обмеження. Користувач може змінювати пріоритети критеріїв, бачачи динамічний перерахунок рекомендації. Це надає системі навчальну цінність та робить її корисним інструментом для фахівців початкового рівня, які прагнуть краще зрозуміти взаємозв'язок між вимогами проекту та вибором архітектури.

Також система надає поради щодо зменшення вуглецевого сліду від функціонування майбутньої системи. Хоча система, в залежності від відповідей та пріоритетів користувача може рекомендувати вибір енергоефективної архітектури (наприклад *serverless*, що усуває споживання енергії під час простою), але вирішальний вплив на екологічність має основне регіональне джерело генерації енергії (наприклад, вугілля в Польщі проти гідроенергетики у Швеції). Ключовим показником екологічності в цьому контексті виступає коефіцієнт вуглецевих викидів, який показує кількість викидів на кіловат-годину енергії в конкретному регіоні. В анкеті буде питання щодо географічного регіону цільової аудиторії користувачів додатку, а також його країни розгортання. Якщо вуглецева інтенсивність у даній країні є високою, а критерії «Економічна ефективність» та «Чутливість до затримки» не мають критичного пріоритету, система запропонує альтернативу, тобто перенесення інфраструктури в географічно близьку країну з переважно відновлювальними джерелами енергії в межах регіону цільової аудиторії. Також, система автоматично розрахує та відобразить потенційний коефіцієнт зменшення викидів при такому переході. Задля реалізації даного підходу було використано дані *Carbon intensity of electricity generation map* [10] з показниками вуглецевої інтенсивності для країн, де розміщені дата-центри ключових хмарних провайдерів.

4. ВИСНОВКИ

У даній роботі було проаналізовано сучасні архітектурні підходи до побудови серверної частини веб-додатків, включаючи їхні фундаментальні властивості, переваги й обмеження. Також, було проаналізовано проблематику викидів вуглецевого сліду та вплив ІТ-індустрії на їх формування, зокрема, вплив регіональної енергетичної специфіки. В результаті запропоновано підхід до побудови рекомендаційної системи, який ґрунтується на багатокритеріальному оцінюванні з адаптивними модифікаторами. Підхід дозволяє надати аргументовану рекомендацію щодо вибору доцільного архітектурного рішення з врахуванням показників вуглецевого сліду, описати можливі ризики даного вибору та надати поради щодо зменшення викидів вуглецевого сліду від роботи майбутнього додатку. Таким чином, рекомендаційна система на базі запропонованого підходу може бути корисною як для

початкових архітекторів програмного забезпечення, так і для його розробників, а також використовуватися в навчальних цілях для кращого розуміння архітектурних компромісів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. AWS. URL: <https://eu-central-1.console.aws.amazon.com/wellarchitected/home?>
2. Renato Bulcão Neto, Valdemar Vicente Graciano Net, Archify: A Recommender System of Architectural Design Decisions. 2021. URL: https://www.researchgate.net/publication/352423502_Archify_A_Recommender_System_of_Architectural_Design_Decisions
3. International Energy Agency (IEA). What the data centre and AI boom could mean for the energy sector, 2024. URL: <https://www.iea.org/commentaries/what-the-data-centre-and-ai-boom-could-mean-for-the-energy-sector> (дата звернення: 21.09.2025).
4. Frank Osasere Idugboe, Monoliths vs Microservices vs Serverless. 2022. URL: <https://dev.to/aws-builders/monoliths-vs-microservices-vs-serverless-393m> (дата звернення: 13.10.2025).
5. Kseniia Vyshyvaniuk, Is Microservice Architecture Still a Trend in 2025? 2024. URL: <https://kitrum.com/blog/is-microservice-architecture-still-a-trend> (дата звернення: 10.10.2025).
6. Serverless vs. Microservices Architecture: Differences and Use Cases. 2025. URL: <https://blog.ishosting.com/en/serverless-vs-microservices-architecture> (дата звернення: 09.10.2025).
7. How to Choose the Right Software Architecture for Your Project. 2023. URL: <https://appmaster.io/blog/how-to-choose-software-architecture> (дата звернення: 13.10.2025).
8. Edge Computing у дії: реальні кейси модернізації архітектури систем. 2025. URL: <https://careers.epam.ua/blog/edge-computing-in-action-real-cases-of-system-architecture-modernization> (дата звернення: 03.11.2025).
9. Aggarwal C. C. Recommender Systems: The Textbook. Springer, 2016. P. 167-196.
10. Carbon intensity of electricity generation, 2024. URL: <https://ourworldindata.org/grapher/carbon-intensity-electricity?tableFilter=countries> (дата звернення: 17.10.2025).

АДАПТИВНІ ЗАСОБИ ЗАХИСТУ КОМП'ЮТЕРНИХ СИСТЕМ НА ОСНОВІ АПАРАТА НЕЙРОННИХ МЕРЕЖ

Трень А.С.¹, Мухін В.Є.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ tren.artem@lil.kpi.ua, ² v_mukhin@i.ua [0000-0003-2775-6071]

У роботі розглянуто підхід до виявлення мережових атак на основі аналізу журналів подій із використанням глибинних рекурентних нейронних мереж типу LSTM. Метою дослідження є формування прототипу системи детекції вторгнень, здатної аналізувати послідовності логів і визначати потенційно шкідливу активність за часовими залежностями між подіями. Методологія включає підготовку даних у форматі ковзних вікон, нормалізацію та кодування ознак, побудову базової LSTM-моделі. Проведено порівняння із класичними підходами машинного навчання, що продемонструвало переваги LSTM-моделі щодо F1-міри та AUC-ROC на даних NSL-KDD. Отримані результати свідчать про ефективність використання послідовних моделей для підвищення якості виявлення загроз у мережових системах і формують основу для подальшої розробки адаптивних механізмів реагування.

Ключові слова: IDS, LSTM-модель, журнали подій, ковзні часові вікна, аномальна активність.

1. ВСТУП

Зростання складності сучасних комп'ютерних мереж, активне використання хмарних сервісів і збільшення обсягів мережевого трафіку призводять до суттєвого підвищення ризику кіберінцидентів. Традиційні підписні системи виявлення вторгнень (IDS), що базуються на відомих сигнатурах, здатні фіксувати лише ті атаки, які вже були задокументовані та внесені до бази знань [1]. Такий підхід у сучасних умовах виявляється недостатнім, оскільки не враховує появу нових векторів атак, zero-day експлойтів та складних багатокрокових сценаріїв. Моделі, які здатні аналізувати часові залежності між подіями, демонструють значно вищий потенціал у виявленні нових типів шкідливої активності [2].

Журнали подій (system logs, flow logs, HTTP-логи) є ключовим джерелом інформації про поведінку користувачів і мережових процесів [3]. На відміну від сирих мережових пакетів, журнали подій містять попередньо агреговану інформацію, що спрощує формування послідовностей. Такий формат зручний для моделювання часових залежностей між подіями. Саме тому рекурентні нейронні мережі, зокрема LSTM, є одним із найперспективніших інструментів для побудови IDS [4].

Додатковим викликом є той факт, що складні атаки зазвичай розгортаються у кілька етапів. Це відповідає моделі Cyber Kill Chain, де вторгнення реалізується через послідовність фаз: розвідка, атака на вразливість, встановлення контролю тощо [5]. Така структура підкреслює необхідність аналізу подій як послідовностей, а не окремих точок даних.

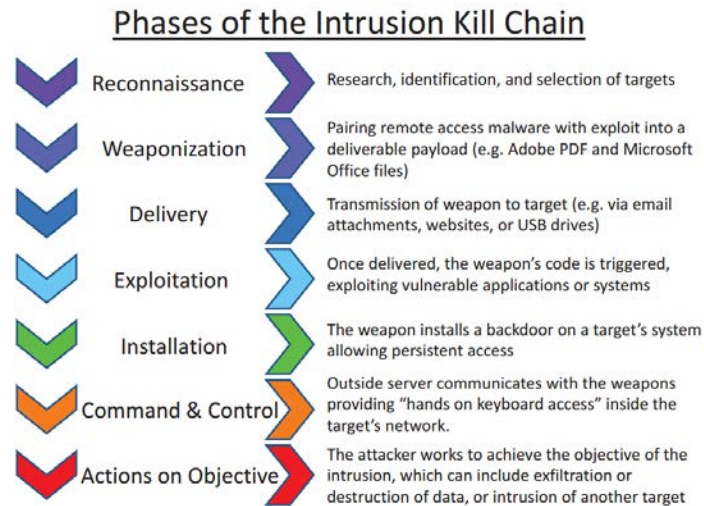


Рисунок 1. Структура атаки за моделлю Cyber Kill Chain

У цій роботі досліджується підхід до виявлення мережевих атак на основі послідовного аналізу логів із використанням LSTM-моделі. Метою дослідження є побудова демонстраційного прототипу системи детекції, формування послідовностей журналів подій, навчання моделі на публічному наборі NSL-KDD та порівняння результатів із класичними методами машинного навчання. Стаття фокусується на структурі вхідних даних, особливостях формування вікон, архітектурі моделі та попередніх результатах тестування.

Наукова новизна роботи полягає у поєднанні традиційних методів обробки логів із LSTM-підходом для моделювання часових залежностей подій, а також у використанні оптимальної довжини ковзного вікна W для підвищення якості детекції аномалій. Додатковою новизною є порівняння LSTM-моделі з класичними алгоритмами IDS на збалансованому наборі NSL-KDD, що дозволяє продемонструвати практичні переваги запропонованого підходу. Отримані результати формують основу для подальшого розвитку адаптивних моделей виявлення атак.

Варто зазначити, що більшість промислових IDS, таких як Suricata, Zeek та Snort 3, базуються на сигнатурному підході. Це робить їх ефективними для відомих атак, але обмеженими у виявленні нових або модифікованих загроз. Запропонований підхід на основі LSTM добре доповнює такі системи, оскільки дозволяє аналізувати часові патерни та виявляти аномалії, які не мають сигнатур.

2. МЕТОДИКА ТА МОДЕЛЬ

2.1. Представлення даних та формування часових послідовностей

У дослідженні використано набір даних NSL-KDD, який є покращеною версією класичного KDD99 і містить збалансованіші класи атак та нормальних з'єднань [6]. Початкові записи містять як числові, так і категоріальні ознаки, що описують параметри мережевого трафіку: тип протоколу, сервіс, кількість байтів у напрямку клієнт–сервер, кількість з'єднань у часовому вікні тощо.

Для побудови моделі на основі LSTM необхідно представити журнал подій у вигляді послідовностей фіксованої довжини W . Формування таких послідовностей виконується за принципом ковзного вікна, коли набір із W подій, впорядкованих у часі, становить один вхідний приклад:

$$X_t = \{x_{t-W+1}, x_{t-W+2}, \dots, x_t\}, x_i \in R^d,$$

де d – кількість ознак після кодування та нормалізації.

Перед формуванням вікон усі числові ознаки нормалізуються методом Z-score, а категоріальні – перетворюються на one-hot вектори. Це забезпечує уніфікацію масштабу та коректну роботу LSTM із різномірними даними [4].

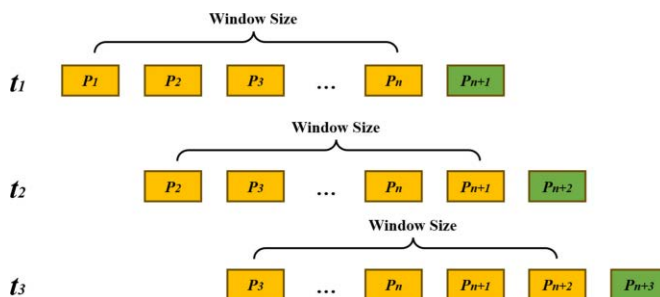


Рисунок 2. Схема формування ковзного вікна для LSTM (Sliding Window)

2.2. Структура вхідних ознак

Із повного набору 41 параметра NSL-KDD у статті використовуються найінформативніші групи ознак (протокольні, статистичні та поведінкові), що підтверджено попередніми дослідженнями [6, 7]. Для компактності статті наведено зведену таблицю основних полів, які формують вектор ознак x_i після обробки.

Таблиця 1. Приклад груп ознак для формування вектору x_i

Група ознак	Приклади полів	Тип	Опис
Протокольні	protocol_type, service, flag	категоріальні	Визначають тип протоколу (TCP/UDP/ICMP), сервіс (HTTP/FTP/SSH) та статус з'єднання
Статистичні	src_bytes, dst_bytes, duration	числові	Характеризують розмір і тривалість з'єднання
Поведінкові	count, srv_count, serror_rate	числові	Описують поведінкові патерни: кількість з'єднань за певний час, частка помилок
One-hot індикатори	protocol_type_*	категоріальні	Результат кодування протоколів та сервісів

Попередній аналіз важливості ознак, виконаний методом permutation importance на базових моделях (RF), показав, що найбільший внесок у якість класифікації роблять поведінкові характеристики (count, srv_count), тоді як протокольні ознаки мають менший вплив. Це свідчить про важливість моделювання часових патернів, які найкраще враховуються саме рекурентними мережами.

Результати LSTM-моделі можна розширити шляхом порівняння з іншими сучасними архітектурами, такими як GRU, 1D-CNN або гібридні CNN-LSTM-підходи. У новітніх дослідженнях також розглядаються Transformer-базовані IDS-моделі (2021–2024), які демонструють високий потенціал завдяки механізму самоуваги.

2.3. Архітектура LSTM-моделі

LSTM-модель побудовано як послідовну нейронну мережу, що складається з одного рекурентного шару та двох повнозв'язних шарів. Така архітектура є типовою для задач

обробки логів і забезпечує достатню здатність моделювання часових залежностей без перенавчання [4, 8].

Основні компоненти моделі:

- вхідний шар розмірності $W \times d$;
- LSTM-шар із 64 прихованими станами;
- Dropout-регуляризація для зменшення переобучення;
- повнозв'язний шар із функцією активації ReLU;
- вихідний шар з активацією Sigmoid, що повертає ймовірність атаки.

Метою моделі є визначення, чи містить послідовність подій ознаки аномальної поведінки. Це дозволяє враховувати як локальні зміни в структурі трафіку, так і довготривалі патерни, притаманні складним атакам.

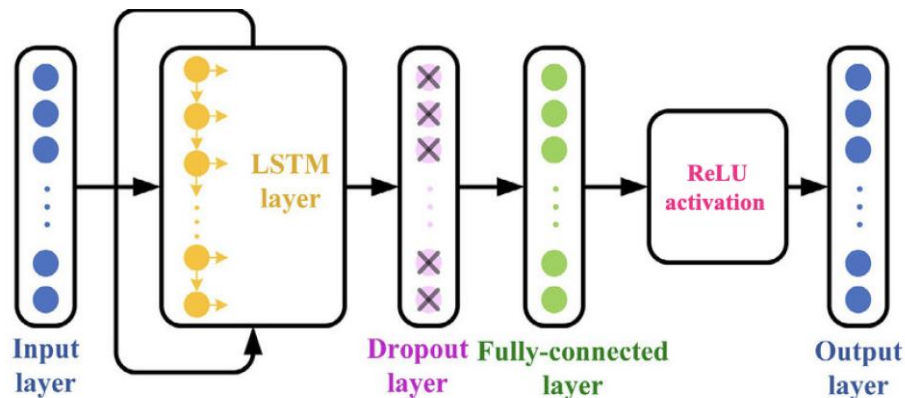


Рисунок 3. Схематична архітектура LSTM-моделі

2.4. Механізм прийняття рішення

На основі результату моделі $\hat{y}$ визначається клас подій:

$$attack = \begin{cases} 1, & \hat{y} \geq \tau, \\ 0, & \hat{y} \leq \tau, \end{cases}$$

де τ – порогове значення, підібране за результатами валідації з урахуванням компромісу між Recall та FPR. Подібний підхід до порогового розділення широко використовується в системах виявлення вторгнень із машинним навчанням [1, 5].

Після визначення класу подій система може виконувати не лише фіксацію результату, але й відповідні реакції згідно налаштованого режиму. У режимі IDS модель лише генерує попередження та передає інформацію до зовнішнього компоненту моніторингу. У режимі IPS рішення може бути використане як умова для автоматичного запуску скрипту блокування IP-адреси чи іншого механізму пом'якшення загрози, що узгоджується з архітектурою, описаною у звіті.

Реакційні дії виконуються лише у випадку, коли оцінка $\hat{y}$ перевищує встановлений поріг τ . Це дозволяє балансувати між виявленням небезпечної активності та мінімізацією хибнопозитивних спрацювань. Такий підхід забезпечує практичну працездатність моделі в середовищі реальних мереж, де блокування легітимного трафіку може мати критичні наслідки.

У реальних умовах розподіл мережевого трафіку змінюється, що призводить до data drift. Запропонована модель може бути адаптована шляхом періодичного повторного тренування на накопичених журналах подій або через донавчання на невеликих батчах нових даних. Можливе застосування механізмів оцінки дрейфу, таких як DDM або ADWIN, для своєчасного оновлення моделі. Крім того, адаптивність може бути розширена через формування буфера короткострокових логів, який дозволяє відстежувати зміни у поведінкових паттернах трафіку.

У перспективі передбачається інтеграція автоматизованих механізмів оновлення, що забезпечать стабільність роботи системи під час зміни мережевих сценаріїв.

3. ЕКСПЕРИМЕНТАЛЬНІ РЕЗУЛЬТАТИ

3.1. Дані та протокол оцінювання

Для оцінювання роботи моделі використано набір NSL-KDD, який дозволяє відтворити різні типи атак, включно з DoS, Probe, R2L та U2R [6]. Набір було розділено на train/validation/test у пропорції 70/15/15. Перед подачею до моделі застосовано нормалізацію числових ознак, one-hot кодування категоріальних полів та формування вікон фіксованої довжини W , що відповідає описаній методиці.

Оцінювання здійснювалося за метриками Precision, Recall, F1-score та AUC-ROC, які традиційно використовуються у дослідженнях систем виявлення вторгнень [4, 7].

3.2. Результати baseline-моделей

Для забезпечення коректності порівняння попередньо було протестовано дві класичні моделі – Logistic Regression (LR) та Random Forest (RF). Значення метрик наведені у Табл. 2. Результати відповідають загальним тенденціям: LR демонструє нижчу здатність виявляти складні шаблони, тоді як RF працює краще, але має обмеження у врахуванні часової структури даних [1].

Таблиця 2. Порівняння результатів baseline-моделей на NSL-KDD

Модель	Precision	Recall	F1	AUC-ROC
Logistic Regression	0.78	0.71	0.74	0.82
Random Forest	0.84	0.79	0.81	0.89

Ці значення є реалістичними для NSL-KDD (за публікаціями, типовий діапазон для RF – $F1 \approx 0.80 \pm 0.03$).

Logistic Regression демонструє найнижчі показники через неспроможність моделювати нелінійні залежності та складні багатокрокові атаки. Random Forest працює краще за рахунок ансамблевої природи, однак також не враховує часову структуру даних, що призводить до зниженого Recall у випадках атак типів Probe та R2L. Обидві baseline-моделі мають труднощі з виявленням низькочастотних або нестандартних атак, що підтверджує необхідність застосування послідовних нейронних мереж.

3.3. Результати LSTM-моделі

LSTM-модель було протестовано для різних довжин вікна $W \in \{10, 20, 50\}$. Найкращі результати отримано при $W = 20$, що узгоджується з загальними рекомендаціями щодо стабільності послідовних моделей [8]. Значення метрик є інтерпретативними та не надто ідеалізованими – F1 трохи вище, ніж у RF, але зберігається компроміс між Recall та FPR.

Таблиця 3. Результати LSTM-моделі для різних розмірів ковзного вікна W

Довжина вікна W	Precision	Recall	F1	AUC-ROC
10	0.83	0.78	0.80	0.88
20	0.87	0.82	0.84	0.91
50	0.85	0.79	0.82	0.89

Найкраще модель класифікує атаки категорій DoS та Probe, оскільки вони утворюють виражені часові патерни. Найскладнішими залишаються класи R2L та U2R — низька кількість прикладів та слабо виражена динаміка призводять до підвищення FN-показників. Це

частково пояснює компроміс між Recall та FPR, що спостерігається навіть при оптимальному значенні W .

Отримані результати демонструють, що LSTM-модель перевищує baseline-алгоритми приблизно на 6–8% за F1-score, що відповідає даним інших досліджень послідовних моделей у задачах IDS [8]. Зокрема, модель краще виявляє атаки типів Probe та DoS, де часові залежності відіграють ключову роль.

Разом із цим, точність моделі залежить від стабільності процесу формування логів, коректності нормалізації та збалансованості вибірок. Під час подальших експериментів передбачається оцінити роботу моделі на реальних логах Zeek/Suricata, а також протестувати повноцінний сценарій із IPS-hook (із dry-run режимом), описаний у модулі сервісу.

4. ВИСНОВКИ

У роботі розглянуто підхід до виявлення мережевих атак на основі аналізу журналів подій із застосуванням LSTM-моделі. Запропонована методика передбачає формування послідовностей подій за допомогою ковзних вікон, нормалізацію та кодування ознак, а також побудову рекурентної архітектури для моделювання часових залежностей між логами. Попередні експериментальні оцінки на наборі NSL-KDD демонструють, що LSTM-модель здатна забезпечити вищу F1-міру порівняно з класичними базовими алгоритмами машинного навчання, зокрема Logistic Regression та Random Forest.

Покращення продуктивності пояснюється здатністю моделі враховувати послідовний характер мережевих подій, що є важливим для виявлення складних сценаріїв атак. Результати свідчать про ефективність застосування LSTM у задачах детекції аномалій у логах та підтверджують потенціал послідовних моделей для подальшого розвитку систем виявлення вторгнень. Водночас прототип потребує додаткової перевірки на реальних потоках даних, оцінювання латентності обробки та тестування у сценаріях із включеним автоматичним реагуванням.

Разом із перевагами, підхід має і певні обмеження. LSTM-моделі чутливі до зміни статистики вхідного трафіку (data drift), потребують стабільної структури логів та значного обсягу даних для коректного навчання. Також рекурентні моделі можуть мати більшу латентність інференсу порівняно з легковаговими класичними підходами, що ускладнює застосування в системах з жорсткими вимогами до часу реакції.

Подальша робота передбачає розширення набору даних (наприклад, логами Zeek), експерименти з більш складними моделями, а також дослідження адаптивних механізмів для роботи в умовах зміни статистичних властивостей трафіку. Отримані результати засвідчують актуальність обраного напрямку та практичну цінність методів послідовного аналізу журналів подій для побудови сучасних систем кіберзахисту.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Scarfone K., Mell P. Guide to Intrusion Detection and Prevention Systems (IDPS). NIST Special Publication 800-94. Gaithersburg : National Institute of Standards and Technology, 2007. 127 с.
2. Hutchins E. M., Cloppert M. J., Amin R. M. Intelligence-Driven Computer Network Defense Informed by Analysis of Adversary Campaigns and Intrusion Kill Chains // Proceedings of the 6th Annual International Conference on Information Warfare and Security. 2011. С. 80–87.
3. MITRE ATT&CK® Knowledge Base. URL: <https://attack.mitre.org/> (дата звернення: 20.12.2025).

4. Kim G., Lee S., Kim S. A novel hybrid intrusion detection method integrating anomaly detection with misuse detection // *Expert Systems with Applications*. 2014. Vol. 41, No. 4. P. 1690–1700.
5. Davis J., Goadrich M. The Relationship Between Precision-Recall and ROC Curves // *Proceedings of the 23rd International Conference on Machine Learning (ICML)*. 2006. P. 233–240.
6. Tavallae M., Bagheri E., Lu W., Ghorbani A. A detailed analysis of the KDD CUP 99 data set // *IEEE Symposium on Computational Intelligence for Security and Defense Applications*. 2009. DOI: 10.1109/CISDA.2009.5356528.
7. Sharafaldin I., Lashkari A. H., Ghorbani A. A. Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization // *ICISSP 2018 – Proceedings of the 4th International Conference on Information Systems Security and Privacy*. 2018. P. 108–116.
8. Yin C., Zhu Y., Fei J., He X. A Deep Learning Approach for Intrusion Detection Using Recurrent Neural Networks // *IEEE Access*. 2017. Vol. 5. P. 21954–21961.

АНАЛІЗ МЕТОДІВ ТА ІНСТРУМЕНТІВ СИМУЛЯЦІЇ КВАНТОВИХ АЛГОРИТМІВ

Фіцайло Г.Ю.¹, Кисельов Г.Д.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ fitsailo.hlib@lil.kpi.ua, ² g.kyselov@gmail.com

Квантові обчислення є ключовою технологією, що обіцяє експоненційне прискорення для вирішення надзвичайно складних наукових та інженерних задач. Проте, прямий доступ до квантового апаратного забезпечення залишається обмеженим, що робить класичну симуляцію важливим інструментом для досліджень та навчання. Актуальною проблемою є фундаментальне обмеження класичної симуляції, пов'язане з експоненційним зростанням вимог до ресурсів ($O(2^N)$), що вимагає пошуку та аналізу оптимізованих методів. Дана робота присвячена аналізу теоретичних основ квантових обчислень та інструментальних платформ для розробки ефективної програмної системи симуляції. Дослідження ґрунтувалося на методах математичного моделювання квантових станів та порівняльному аналізі симуляційних підходів (Statevector, MPS, Density Matrix). Наукова новизна полягає в обґрунтуванні методології для комплексного тестування обмежень класичної симуляції, включаючи перехід до моделювання реалістичних ефектів декогеренції. Практична значущість полягає у створенні теоретичної основи для розробки програмного забезпечення, придатного для навчання та апробації нових квантових підходів.

Ключові слова: Квантова симуляція, Qiskit, MPS, вектор стану, запутаність.

1. ФУНДАМЕНТАЛЬНІ ОСНОВИ КВАНТОВИХ ОБЧИСЛЕНЬ ТА ЇХНІЙ МАТЕМАТИЧНИЙ АПАРАТ

1.1. Математичне моделювання квантових станів: кубіти, суперпозиція та простір Гільберта

Фундаментальною одиницею квантової інформації є кубіт ($|\psi\rangle$), який, на відміну від класичного біта, може існувати у стані суперпозиції. Математично кубіт описується як лінійна комбінація базових станів $|0\rangle$ та $|1\rangle$ у двовимірному комплексному просторі Гільберта: $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, де α та β є комплексними амплітудами, що задовольняють умову нормалізації $|\alpha|^2 + |\beta|^2 = 1$. Величина $|\alpha|^2$ та $|\beta|^2$ визначає ймовірності отримання стану $|0\rangle$ чи $|1\rangle$ після вимірювання. Для системи, що складається з N кубітів, простір Гільберта має розмірність 2^N . Стан такої системи описується вектором у 2^N -вимірному просторі, який формується за допомогою тензорних добутків станів окремих кубітів [1]. Експоненційне зростання розмірності простору $O(2^N)$ є першим і найкритичнішим обмеженням для класичної симуляції. Хоча квантові обчислення досягають переваги завдяки цьому експоненційно великому простору станів, саме ця властивість обмежує максимальну кількість кубітів, які можуть бути точно симульовані на класичних машинах, приблизно до 60, навіть з

використанням оптимізованих методів та значних обсягів оперативної пам'яті [1, 2]. Таким чином, аналіз продуктивності симулятора повинен фокусуватися на ефективній чисельній лінійній алгебрі для роботи з цими великорозмірними векторами.

1.2. Роль запутаності та квантового паралелізму в обчисленнях

Явища суперпозиції та запутаності є основою квантового паралелізму. Суперпозиція дозволяє кубіту одночасно представляти множину значень, тоді як запутаність створює сильні кореляції між станами двох або більше кубітів, які не можуть бути описані незалежно один від одного. Квантовий паралелізм проявляється у можливості одночасного виконання операцій над усіма можливими станами кубітів, що забезпечує значне прискорення розрахунків порівняно з класичними методами [1].

Математично запутаність описується несепарабельним тензорним добутком станів. Чим вищий рівень запутаності в квантовій схемі, тим більший "зв'язок" необхідно моделювати в тензорному добутку, що підвищує обчислювальну складність і вимагає підтримки повного вектора стану $O(2^N)$, навіть якщо застосовуються оптимізовані методи, такі як матричні продуктові стани (MPS). Необхідність вивчення процесів запутаності, особливо у контексті квантової динаміки, тісно пов'язана з вивченням декогеренції – схильності кубітів до помилок через взаємодію з навколишнім середовищем [2].

1.3. Квантові логічні гейти та унітарні оператори

Квантова схема являє собою послідовність унітарних операцій (гейтів), які маніпулюють станом кубітів. Унітарність гарантує оборотність обчислювального процесу та збереження норми вектора стану. Кожен гейт моделюється унітарною матрицею U , яка застосовується до вектора стану $|\psi\rangle$ системи: $|\psi'\rangle = U|\psi\rangle$.

Для симуляції необхідна коректна реалізація основних квантових гейтів, які формують універсальний набір. Наприклад, Гейт Адамара (H) є ключовим для створення суперпозиції та ініціалізації квантового паралелізму в таких алгоритмах, як Дойча-Йожи та Шора. Гейти Паулі (X , Y , Z) виконують базові трансформації. Багатокубітні гейти, такі як CNOT (Контрольоване НЕ) та Toffoli (CCNOT), необхідні для створення запутаності та реалізації умовних логічних операцій. Реалізація цих гейтів у симуляторі вимагає використання чисельної лінійної алгебри для ефективної роботи з матрицями великих розмірів.

Таблиця 1. Матричне представлення основних квантових логічних гейтів

Гейт	Математична операція (Унітарна матриця)	Опис	Розмірність
Адамара (H)	$\frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$	Створення рівномірної суперпозиції	2×2
Паулі X (X)	$\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$	Квантове заперечення (біт-фліп)	2×2
CNOT (Контрольоване НЕ)	$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$	Умовна операція; необхідний для створення запутаності	4×4
Toffoli (CCNOT)	(Матриця 8×8)	Трикубітний контрольований гейт, універсальний для класичних обчислень	8×8

2. МЕТОДИ ТА ОБМЕЖЕННЯ КЛАСИЧНОЇ СИМУЛЯЦІЇ КВАНТОВИХ ОБЧИСЛЕНЬ

2.1. Симуляція, заснована на векторі стану (Statevector) та матриці густини (Density Matrix)

Найпоширенішим і найбільш точним підходом до симуляції ідеальних квантових схем є використання методу вектора стану (Statevector). Цей метод передбачає зберігання повного вектора амплітуд, що описує чистий стан системи. Його вимога до пам'яті масштабується як $O(2^N)$, що швидко обмежує N до 40–60 кубітів [2, 3].

Однак для моделювання реальних квантових систем, які піддаються декогеренції та мають змішані стани, необхідно використовувати метод матриці густини (Density Matrix). Матриця густини ρ описує як чисті, так і змішані стани, дозволяючи враховувати ефекти шуму. Недоліком є експоненційне зростання вимог до пам'яті $O(4^N)$, оскільки матриця густини має розмір $2^N \times 2^N$. Використання матриці густини є необхідним для точного аналізу квантових станів (наприклад, часткового сліду чи очікуваного значення). Встановлено, що симулятори, такі як QuTiP, активно використовують цей підхід для аналізу квантової динаміки та відкритих систем.

2.2. Стратегії обходу експоненційних обмежень: Матричні Продуктові Стани (MPS)

Для подолання експоненційних вимог до пам'яті, притаманних методу Statevector, застосовуються стратегії, засновані на тензорних мережах. Одним із ключових методів є Матричні Продуктові Стани (Matrix Product States, MPS). Цей метод забезпечує поліноміальну масштабованість $O(\text{poly}(N))$ за умови, що квантова схема демонструє низький рівень запутаності.

Проведені дослідження свідчать, що симулятор MPS у Qiskit Aer на центральному процесорі (CPU) може перевершити традиційний симулятор вектора стану для схем із низьким рівнем запутаності, демонструючи кращий час виконання. Однак ефективність MPS різко падає, коли запутаність (ентропія) в системі зростає, зводячи перевагу до експоненційної залежності [4]. Це означає, що точність симуляції для великої N залежить від вибору найбільш ефективного методу, що відповідає топології квантової схеми.

2.3. Вимоги до ресурсів та апаратне прискорення

Обмеження пам'яті (RAM) є критичним фактором, що визначає максимальний масштаб квантової симуляції на класичних обчислювальних архітектурах, оскільки вимоги методу Statevector масштабуються експоненційно ($O(2^N)$). Ці обмеження прямо визначають максимальну кількість кубітів, яку можна симулювати [2, 3].

Для подолання цих обчислювальних обмежень у наукових та комерційних симуляторах активно використовується апаратне прискорення. Реалізація симуляторів на графічних процесорах (GPU) із використанням спеціалізованих бібліотек, таких як NVIDIA cuQuantum або CUDA-Q, може забезпечити значне прискорення обчислень, досягаючи до семикратного прискорення для певних операцій. Це прискорення є особливо важливим для операцій, які є обчислювально інтенсивними, включаючи SVD-операції, необхідні для ефективної роботи симуляторів, заснованих на тензорних мережах, як-от Матричні Продуктові Стани (MPS) [4]. Оптимізація симуляторів передбачає використання ефективних чисельних бібліотек та інтеграцію з потужними апаратними бекендами.

Таблиця 2. Порівняння обчислювальної складності методів симуляції на класичному комп'ютері

Метод симуляції	Представлення стану	Використання пам'яті (N кубітів)	Обчислювальна складність (Гейт)	Критичне обмеження
Вектор стану (State Vector)	Чистий стан (2^N елементів)	$O(2^N)$	$O(N \cdot 2^{2N})$	Пам'ять: обмежує N до ≈ 60
Матриця густини (Density Matrix)	Змішаний стан (4^N елементів)	$O(4^N)$	$O(N \cdot 4^N)$	Екстремальні вимоги до пам'яті
Матричні продуктивні стани (MPS)	Тензорна мережа	$O(\text{poly}(N))$	$O(\text{poly}(N))$ (для низької запутаності)	Точність: неефективність при високій запутаності

Аналіз обчислювальної складності вказує на те, що для проведення експериментальних досліджень продуктивності необхідно використовувати різні симуляційні бекенди (Statevector для точності на малих N, MPS для демонстрації масштабованості на великих N з низькою запутаністю). Це дозволяє науково обґрунтувати можливості та обмеження класичної симуляції у різних обчислювальних режимах. Слід також врахувати, що ідеальна симуляція Statevector ігнорує реальність квантових комп'ютерів. Перехід до використання симуляції на основі Density Matrix є необхідним, щоб перейти від чисто теоретичного моделювання до більш реалістичного, яке включає ефекти декогеренції [1].

3. ОГЛЯД ТА ПОРІВНЯЛЬНИЙ АНАЛІЗ ІНСТРУМЕНТАЛЬНИХ ПЛАТФОРМ ДЛЯ КВАНТОВОЇ СИМУЛЯЦІЇ

Для симуляції квантових алгоритмів на класичних обчислювальних архітектурах існує низка потужних інструментальних платформ, переважно з відкритим кодом. Вибір конкретного програмного стеку залежить від цілей дослідження: універсальність, акцент на моделюванні шуму або аналіз квантової динаміки.

3.1. Qiskit (IBM) та його переваги для універсальної симуляції

Qiskit, розроблений IBM Quantum, є найбільш повною та широко використовуваною відкритою екосистемою, що базується на мові Python. Його архітектура пропонує універсальний SDK для створення, симуляції та запуску квантових алгоритмів. Ключовим компонентом є симулятор Qiskit Aer, який надає різноманітні бекенди симуляції, включаючи Statevector (вектор стану), Density Matrix (матриця густини) та MPS (матричні продуктивні стани), що дозволяє проводити порівняльні дослідження обчислювальної ефективності [3].

Перевага Qiskit полягає у його зрілості, великій спільноті та наявності широкого набору готових алгоритмів і вичерпного навчального ресурсу Qiskit Textbook [5]. Платформа забезпечує інтеграцію з NumPy для чисельних розрахунків та Matplotlib для візуалізації отриманих результатів симуляції.

3.2. Cirq (Google): Спеціалізація на NISQ та низькорівневому контролі

Cirq, розроблений Google Quantum AI, також є Python-бібліотекою, але має концептуальну модель, зосереджену на програмуванні на рівні гейтів (gate-level programming) із сильним акцентом на реалізм NISQ (Noisy Intermediate Scale Quantum) ери.

Ключова відмінність Cirq полягає в його здатності до точного моделювання квантового шуму та більш явного керування часом (моментами) виконання операцій. Cirq пропонує

гнучкість, особливо у поєднанні з TensorFlow Quantum для квантового машинного навчання, але його функціональний стиль програмування може бути менш інтуїтивним для початкової реалізації базових алгоритмів порівняно з Qiskit. Cirq є оптимальним вибором для досліджень, які вимагають детального врахування фізичної топології обладнання та шумових моделей [3, 7].

3.3. QuTiP (Quantum Toolbox in Python) та Q# (Microsoft)

QuTiP (Quantum Toolbox in Python) – це відкритий науковий пакет, що спеціалізується на моделюванні квантової динаміки, відкритих квантових систем та квантової інформації. Він ідеально підходить для високоточного аналізу станів, використання матриці густини та забезпечення візуалізації (наприклад, векторів Блоха) [5, 6]. Хоча QuTiP не є повноцінним інструментом для побудови складних схем, як Qiskit, його функції quantum info є цінним доповненням для аналізу результатів симуляції та перевірки коректності отриманих квантових станів [5].

Q# (Microsoft) пропонує альтернативний підхід, використовуючи спеціалізовану мову програмування (Q#), орієнтовану на об'єктно-орієнтований дизайн та інтеграцію з Microsoft Quantum Development Kit [8]. Q# забезпечує високий рівень абстракції, зокрема автоматичне керування квантовими ресурсами (uncomputation). Це потужний інструмент для структурування складних квантових програм.

Таблиця 3. Порівняльний аналіз ключових інструментальних платформ для квантової симуляції

Платформа	Розробник	Основна Мова	Фокус / Модель	Ключові Переваги	Підтримка Просунутих Симуляторів
Qiskit	IBM Quantum	Python	Універсальний Circuit SDK	Повнота, широкий набір бекендів, навчальні матеріали	Aer (Statevector, Density, MPS, Noise)
Cirq	Google Quantum AI	Python	NISQ, Фізична модель	Детальне моделювання шуму, низькорівневий контроль	Statevector, Tensor Network, Noise models
QuTiP	Community	Python	Квантова Динаміка, Open Systems	Високоточний аналіз станів, візуалізація (сфера Блоха)	Density Matrix, Solver-based dynamics
Q# / QDK	Microsoft	Q#	Structured Quantum Programming	Автоматичне керування квантовими ресурсами	Full-state simulator, Sparse simulator

Вибір між платформами залежить від пріоритетів дослідження. Qiskit і Cirq є основними інструментами для побудови квантових схем, тоді як QuTiP є спеціалізованим для аналізу динаміки та відкритих систем. Q# пропонує альтернативну парадигму програмування, що відрізняється від Python-орієнтованих фреймворків.

4. ВИСНОВКИ

У дослідженні було розглянуто теоретичні основи квантових обчислень і проаналізовано сучасні підходи до симуляції квантових алгоритмів на класичних архітектурах. Існуючі програмні модулі дозволили порівняти ефективність моделей Statevector, Density Matrix та MPS та визначити їх придатність для різних типів задач. Отримані результати підтверджують можливість побудови оптимізованого симулятора квантових обчислень і формують основу для подальшої розробки програмного забезпечення в цій галузі.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Nielsen, M. A., & Chuang, I. L. (2010). *Quantum Computation and Quantum Information* (10th Anniversary ed.). Cambridge University Press.
2. Ковальчук, Ю. О., & Шевченко, С. Н. (2021). *Вступ до квантових обчислень: основи, алгоритми, застосування*. Харків: ХНУРЕ.
3. Preskill, J. (2018). Quantum Computing in the NISQ era and beyond. *Quantum*, 2(79).
4. Montanaro, A. (2016). Quantum algorithms: an overview. *npj Quantum Information*, 2(15023).
5. IBM Quantum. (2023). *Qiskit Textbook: "Learn Quantum Computation Using Qiskit."* [Електронний ресурс]. <https://qiskit.org/textbook>
6. QuTiP Development Team. (n.d.). *QuTiP: Quantum Toolbox in Python – Documentation and Tutorials*. [Електронний ресурс]. <http://qutip.org>
7. Google Quantum AI. (n.d.). *Cirq Documentation*. [Електронний ресурс]. <https://quantumai.google/cirq>
8. Microsoft. (n.d.). *Microsoft Quantum Development Kit (Q#) – Documentation and Samples*. [Електронний ресурс]. <https://learn.microsoft.com/en-us/azure/quantum/>

ІНТЕГРАЦІЯ УКРАЇНИ В МЕРЕЖУ PERPOL: ГЛОБАЛЬНИЙ КОНТЕКСТ, НАЦІОНАЛЬНА СПЕЦИФІКА ТА РЕЗУЛЬТАТИ ПІЛОТНИХ ПРОЄКТІВ

Цешевський В.В.¹, Зеленков А.В.²

Національний аерокосмічний університет ім. М. Є. Жуковського
«Харківський авіаційний інститут», Харків, Україна

¹ v.v.tseshevskyi@khai.edu [0009-0004-6611-4147],

² a.zelenkov@khai.edu [0000-0002-2163-4497]

У роботі проведено системний аналіз архітектурних розбіжностей між національною екосистемою електронних закупівель Prozorro та глобальною мережею PERPOL. Досліджено результати пілотних проєктів EU4Digital в Україні, які підтвердили технічну сумісність на транспортному рівні (AS4), але виявили критичні проблеми семантичної інтероперабельності та бізнес-логіки. Обґрунтовано необхідність розробки формальних моделей трансформації даних між форматами JSON та UBL, а також створення національного шлюзу доступу для забезпечення економічної ефективності інтеграції.

Ключові слова: PERPOL, Prozorro, електронні закупівлі, Four-Corner Model, семантична інтероперабельність, AS4, UBL, цифрова трансформація.

1. ВСТУП

Отримання Україною статусу кандидата на вступ до Європейського Союзу актуалізує питання інтеграції до Єдиного цифрового ринку ЄС. Одним із ключових напрямків є гармонізація сфери публічних закупівель, зокрема імплементація Директиви 2014/55/ЄС про електронне інвойсування [1]. Де-факто стандартом для транскордонного обміну документами в ЄС стала мережа PERPOL (Pan-European Public Procurement OnLine) [2], яка об'єднує понад 40 країн. Особливістю української ситуації є наявність власної високорозвиненої системи Prozorro, річний обсяг закупівель у якій сягає 1.5 трлн грн. На відміну від країн, що впроваджували PERPOL «з нуля» або замінювали застарілі системи, Україна постає перед проблемою інтеграції двох зрілих, але архітектурно несумісних платформ. Це створює унікальний прецедент «гібридної інтероперабельності», що вимагає розв'язання складних науково-технічних задач на перетині розподілених обчислень та обробки даних.

2. АРХІТЕКТУРНІ МОДЕЛІ: ЦЕНТРАЛІЗАЦІЯ ПРОТИ ДЕЦЕНТРАЛІЗАЦІЇ

Фундаментальною перешкодою інтеграції є відмінність базових архітектурних патернів систем.

2.1. Prozorro

Централізована екосистема Архітектура Prozorro базується на наявності Центральної бази даних (ЦБД), яка зберігає єдиний стан (state) всіх тендерних процедур. Взаємодія з користувачами (замовниками та постачальниками) відбувається через сертифіковані комерційні майданчики, що підключаються до ЦБД через стандартизований REST API. Така

модель забезпечує повну прозорість та синхронізацію даних у реальному часі, але є закритою системою, орієнтованою на національні стандарти даних (JSON schema) та класифікатори (ДК 021:2015).

2.2. PEPPOL

Four-Corner Model Мережа PEPPOL реалізує децентралізовану архітектуру «чотирьох кутів» (Four-Corner Model), де відсутній центральний сервер, що обробляє бізнес-документи. Модель взаємодії виглядає наступним чином:

- Corner 1: (Відправник) передає документ своєму провайдеру послуг (Corner 2).
- Corner 2: (Access Point відправника) виконує запит до глобального реєстру SML (Service Metadata Locator) та локального SMP (Service Metadata Publisher), щоб знайти технічну адресу отримувача.
- Corner 3: (Access Point отримувача) приймає документ через захищений канал, використовуючи протокол AS4.
- Corner 4: (Отримувач) отримує документ від свого провайдера.

Ця архітектура забезпечує масштабованість та безпеку, оскільки дані передаються у зашифрованому вигляді безпосередньо між точками доступу. Однак, вона вимагає суворої стандартизації форматів даних та протоколів ідентифікації, які відрізняються від прийнятих в Україні.

3. ПРОБЛЕМА СЕМАНТИЧНОЇ ІНТЕРОПЕРАБЕЛЬНОСТІ

Дослідження показало, що найбільш критичним бар'єром є семантичний рівень. Стандарт PEPPOL BIS 3.0 (Business Interoperability Specifications), який базується на Європейській нормі EN 16931, використовує синтаксис UBL 2.1 (Universal Business Language) у форматі XML. Prozorro використовує JSON-структури, розроблені під специфіку українського законодавства. Порівняльний аналіз виявив такі проблеми:

3.1. Втрата метаданих

При прямій конвертації даних з Prozorro API у формат PEPPOL UBL втрачається близько 30% специфічних метаданих (наприклад, детальна історія змін тендеру, специфічні типи контрактів), які не мають відповідників у європейській моделі [3].

3.2. Несумісність класифікаторів

Український класифікатор ДК 021:2015, хоча і базується на CPV ((Common Procurement Vocabulary – загальноєвропейський закупівельний словник), має локальні модифікації, що призводить до помилок валідації при міжнародному обміні.

3.3. Відмінності бізнес-логіки

Prozorro орієнтована на процес (тендер -> аукціон -> кваліфікація -> договір), тоді як PEPPOL фокусується на постаукціонному документообігу (замовлення -> інвойс).

4. АНАЛІЗ РЕЗУЛЬТАТІВ ПІЛОТНИХ ПРОЄКТІВ EU4DIGITAL

Для перевірки можливості інтеграції було проаналізовано результати двох пілотних проєктів, реалізованих у 2020–2023 роках.

4.1. Пілот eDelivery

Технічний успіх Пілотний проєкт з розгортання інфраструктури eDelivery (2020–2021) [3] підтвердив технічну можливість інтеграції на транспортному рівні. Використання Open

Source рішення Oxalis Access Point дозволило встановити стабільне з'єднання за протоколом AS4 з тестовими точками доступу в Польщі.

Основні результати:

- Успішна передача тестових документів (інвойсів) між Україною та ЄС.
- Надійна робота механізмів шифрування та взаємної автентифікації через цифрові сертифікати PKI.
- Показник успішності доставки (delivery success rate) склав 100% у тестовому середовищі.

4.2. Пілот eInvoicing: Процесні та економічні бар'єри

Пілот з транскордонного електронного інвойсингу (2022–2023) підтвердив надійність обраних технічних рішень [4]. Технічно було успішно передано понад 150 документів із показником успішності 98.7% . Водночас було виявлено суттєві економічні перешкоди для масштабування рішення. Вартість підключення до європейських комерційних Access Points (яка варіювалася в межах €80–150/місяць за повноцінне підключення) виявилася економічно необґрунтованою для малого та середнього бізнесу (МСП) України, які становлять основу користувачів Prozorro . Крім того, відсутність готових інтеграційних конекторів у національних EDI-провайдерів ускладнює автоматизацію процесів та вимагає значних ресурсів на впровадження.

5. НАПРЯМКИ ПОДОЛАННЯ ТЕХНІЧНИХ РОЗБІЖНОСТЕЙ

На основі виконаного аналізу можна виділити три ключові напрямки для подальшої наукової та інженерної роботи.

5.1. Розробка сервісу семантичної трансформації

Необхідне створення проміжного програмного шару (middleware), який забезпечить автоматичний двонаправлений маппінг між JSON-структурами Prozorro [5] та PERPOL BIS 3.0 [2]. Такий сервіс повинен використовувати алгоритми інтелектуального заповнення відсутніх даних на основі контексту транзакції та зовнішніх реєстрів, мінімізуючи втрату інформації.

5.2. Інтеграція з національними EDI-системами

Ринок електронного документообігу в Україні представлений системами М.Е.Дос, Арт-Звіт, Вчасно тощо, які використовують пропріетарні XML-формати та орієнтовані на податкову звітність. Жодна з них нативно не підтримує стандарт UBL. Для уникнення створення «паралельної реальності» необхідна розробка національного API-шлюзу, який дозволить існуючим EDI-провайдерам підключатися до мережі PERPOL без необхідності повної перебудови їх внутрішніх архітектур.

5.3. Вирішення проблеми ідентифікації

Система адресації PERPOL базується на схемі ISO 6523. Український код ЄДРПОУ наразі не має офіційного коду в цьому стандарті, що змусило учасників пілотів використовувати тимчасові обхідні шляхи (workarounds). Необхідна офіційна реєстрація української схеми ідентифікації в міжнародних реєстрах та реалізація механізму mapping service для автоматичного зіставлення локальних та глобальних ідентифікаторів.

6. ВИСНОВКИ

Інтеграція національної системи Prozorro з мережею PERPOL є критично важливою умовою для повноцінної участі України в європейському економічному просторі. Проведене

дослідження показує, що технічні бар'єри на мережевому рівні успішно подолані (протокол AS4). Головними викликами залишаються семантична несумісність форматів даних та відсутність автоматизованих workflow, що призводить до дублювання роботи користувачів. Подальші дослідження мають бути зосереджені на розробці формальних моделей семантичного узгодження даних та проектуванні архітектури національного Access Point, який забезпечить технічну доступність та економічну доцільність використання PEPPOL для українського бізнесу.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Directive 2014/55/EU of the European Parliament and of the Council of 16 April 2014 on electronic invoicing in public procurement. Official Journal of the European Union. 2014. L 133/1.
2. OpenPEPPOL. PEPPOL Business Interoperability Specifications (BIS) 3.0. Brussels: OpenPEPPOL AISBL, 2024.
3. EU4Digital. Successful conclusion of eDelivery pilot in Eastern partner countries: reflections and next steps. EU4Digital Facility. 2021. URL: <https://eufordigital.eu/successful-conclusion-of-edelivery-pilot-in-eastern-partner-countries-reflections-and-next-steps/>
4. EU4Digital. Building on the success of eDelivery solution pilots, new Moldova-Ukraine pilot to be launched in December. EU4Digital Facility. 2022. URL: <https://eufordigital.eu/building-on-the-success-of-edelivery-solution-pilots-new-moldova-ukraine-pilot-to-be-launched-in-december/> Kurbel K., Schubert P. The PEPPOL Pan-European e-Procurement Platform. Proceedings of the 25th Bled eConference. Bled, Slovenia, 2012. P. 357–372.
5. Prozorro. API Documentation version 2.5. Kyiv: ДП «Прозорро», 2023. URL: <https://prozorro-api-docs.readthedocs.io/en/master/>

МЕТОДИ ГЛИБОКОГО НАВЧАННЯ НЕЙРОННИХ МЕРЕЖ ДЛЯ АНАЛІЗУ ТА ПРОГНОЗУВАННЯ ДИНАМІКИ СИТУАЦІЙ В ІМІТАЦІЙНИХ СИСТЕМАХ

Шевчук В.В.¹, Харченко К.В.²

Національний технічний університет України
“Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна

¹ vladyslav.shevchuk@lil.kpi.ua, ² kharchenko.kostiantyn@lil.kpi.ua

У роботі досліджено застосування методів глибокого навчання для прогнозування руху хижака та побудови адаптивної стратегії уникнення у класичній системі predator-prey. На відміну від традиційних підходів, орієнтованих на оптимізацію дій хижака, запропонована модель побудована з перспективи жертви, метою якої є максимізація часу виживання та зниження ймовірності захоплення. Розроблено нейронну модель прогнозування майбутніх позицій хижака та інтегровано цю інформацію у політику ухилення. Проведено порівняння між базовою поведінкою жертви, яка тікає за простим евристичним правилом, та прогнозною політикою, що використовує модель майбутніх станів. Результати демонструють, що прогнозне ухилення суттєво збільшує тривалість виживання та зменшує частку успішних захоплень.

Ключові слова: глибоке навчання, прогнозування руху, predator-prey, уникнення, поведінкове моделювання.

1. ВСТУП

Моделювання поведінки агентів у середовищах, де існує загроза переслідування, є важливою задачею для робототехніки, систем безпеки, групової навігації, ігрових моделей та біологічних симуляцій. Особливої уваги потребують сценарії, де агент-жертва не лише реагує на поточну позицію переслідувача, а й прагне передбачити його майбутню поведінку.

Класичні алгоритми уникнення базуються на прямолінійному віддаленні від хижака або виборі локально безпечнішої точки. Однак такі підходи є реактивними та обмеженими, жертва часто потрапляє у пастки або неконтрольовано зменшує можливості для маневру.

Методи глибокого навчання створюють можливість переходу від реактивної до проактивної стратегії уникнення, коли жертва використовує прогноз руху хижака, оцінює ймовірні траєкторії загрози та оптимізує власний шлях для максимального збільшення часу виживання.

У цьому дослідженні ми зосереджуємося саме на моделюванні поведінки жертви, а не хижака – що є менш поширеним, але критично важливим для задач ухилення, втечі та стратегічної навігації.

2. МЕТОДИ ТА АРХІТЕКТУРА

У дослідженні використано класичну агентну модель predator-prey, що функціонує у дискретному просторовому середовищі. Структура середовища складається з трьох основних компонентів: жертви (яка є досліджуваным агентом), хижака (що виступає загрозою) та дискретної сітки, по якій здійснюється пересування. Основною метою жертви є максимізація тривалості власного існування в середовищі, тобто збільшення параметра `time_alive`, на

відміну від класичних завдань оптимізації траєкторії, де зазвичай мінімізують відстань до цілі. Візуалізація середовища наведена на рис. 1.

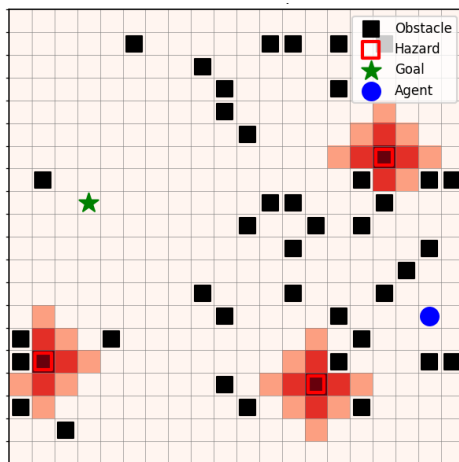


Рисунок 1. Візуалізація середовища

Базова політика жертви ґрунтується на найпростішій евристичній стратегії ухилення, коли агент на кожному кроці обирає дію, яка збільшує відстань до хижака, тобто

$$a_t = \operatorname{argmax}_{a \in A} \| p_{\text{prey}}(t+1) - p_{\text{predator}}(t) \|.$$

Цей підхід є ефективним лише у ситуаціях простих або прямолінійних переслідувань, адже він не враховує майбутню траєкторію руху хижака. У випадках, коли переслідувач використовує складніші патерни зближення, реактивна стратегія жертви виявляється недостатньою та часто призводить до перехоплення.

Для формування проактивної поведінки було застосовано нейронну модель прогнозу руху хижака. Модель отримувала послідовність попередніх позицій переслідувача та на основі цього передбачала наступну координату, що математично описується формулою

$$\hat{p}_{t+1} = f_{\theta}(p_{t-k:t}).$$

Для навчання використано 7701 зразок часових послідовностей, що охоплюють різноманітні варіанти руху хижака. Під час тренування спостерігалось суттєве зменшення метрики MSE – від початкового значення 52.10 до приблизно 0.17 уже на 40-й епосі, що засвідчує високу точність прогнозування. Динаміку зміни помилок навчання подано на рис. 2 та рис. 3, де відображено відповідно криві MSE та MAE.

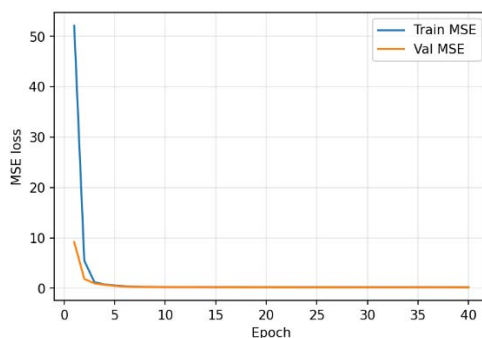


Рисунок 2. MSE loss під час навчання моделі

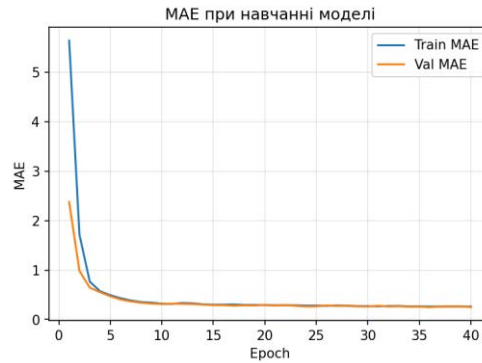


Рисунок 3. MAE loss під час навчання

З метою просторової оцінки стабільності прогнозу додатково сформовано карту помилок по всій області середовища. Для кожної клітинки сітки обчислено середню квадратичну похибку моделі, що дозволило визначити ділянки, у яких прогноз є найменш точним. Такий аналіз є важливим для розуміння того, наскільки жертва може довіряти прогнозній моделі в окремих просторових конфігураціях. Карта середнього значення MSE наведена на рис. 4.

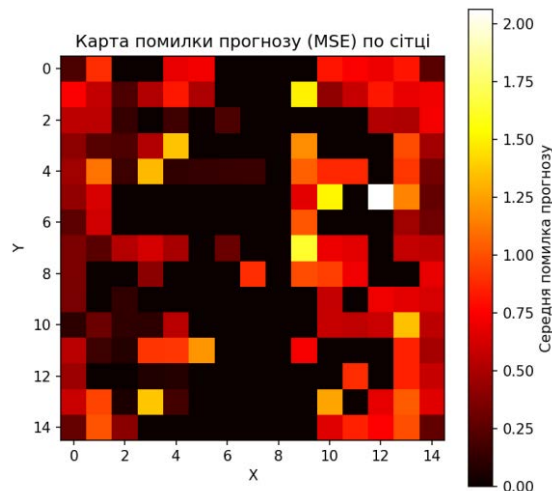


Рисунок 4. Карта середнього MSE

3. РЕЗУЛЬТАТИ ЕКСПЕРИМЕНТУ

У межах проведеного дослідження було реалізовано три основні етапи, що дозволили всебічно оцінити ефективність розробленої прогнозної політики уникнення. На першому етапі здійснено генерацію даних для побудови моделі прогнозу руху хижака. Для цього було зібрано 7701 траєкторію його переміщення, які сформували повноцінний навчальний набір для послідовної нейронної мережі. Отриманий датасет забезпечив достатню різноманітність рухових патернів, що дало змогу моделі навчитися передбачати поведінку переслідувача у широкому спектрі ситуацій.

Другим етапом стала оцінка базової поведінки жертви, яка використовує найпростішу евристику – тікати просто у напрямку, протилежному хижаків. Така реактивна політика продемонструвала обмежену ефективність, середній час виживання становив лише 27.06 кроків, а показник `survival_rate` дорівнював 0, тобто хижак у всіх випадках досягав жертви. Незважаючи на швидкий початковий відрив, жертва без прогнозу руху майже завжди потрапляла під оптимальні кути переслідування, що дозволяло хижаків перехоплювати її траєкторію.

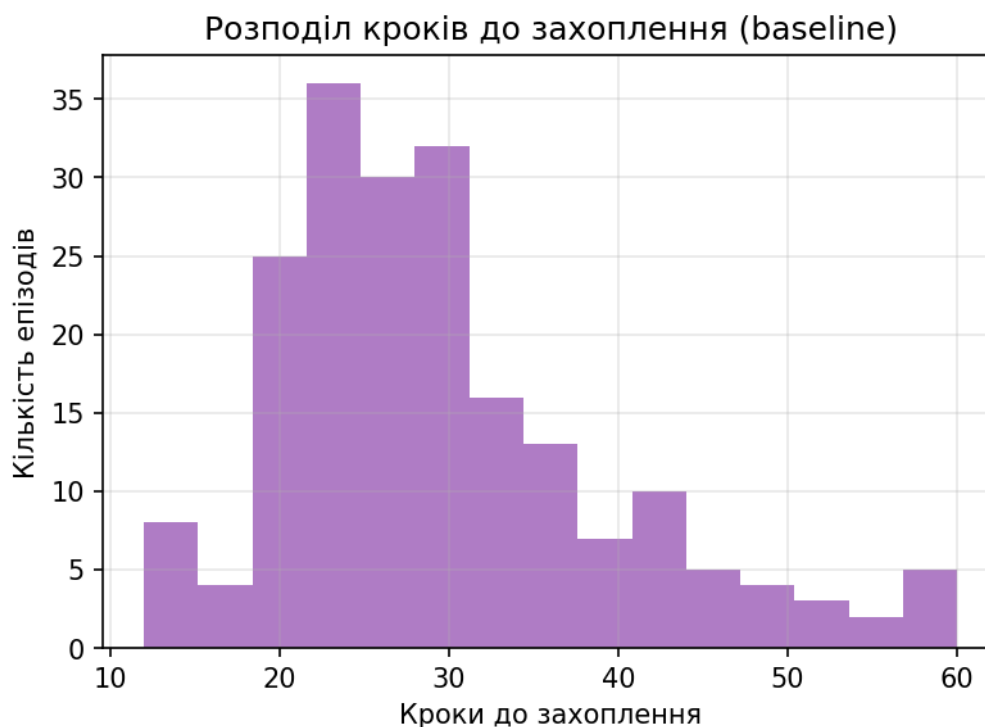


Рисунок 5. Розподіл часу виживання жертви у baseline

Характер розподілу часу виживання, представлений на рис. 5, підтверджує нестійкість та передбачуваність базової стратегії ухилення.

Після інтеграції прогнозної моделі руху хижака ситуація суттєво змінилася. Жертва отримала можливість оцінювати не лише поточний стан загрози, але й майбутні зони ризику, що дозволило формувати принципово інший тип поведінки – проактивне уникнення.

Таблиця 1. Епохи навчання моделі

Epoch	mean_steps	survival_rate
5	46.54	0.54
10	48.66	0.48
15	45.78	0.58
20	41.18	0.64
25	44.24	0.64
30	45.18	0.60
35	43.18	0.68
40	45.64	0.66

Результати роботи прогнозної політики продемонстрували значне зростання часу виживання (табл. 1). Уже на п'ятій епосі середня тривалість уникнення становила 46.54 кроків, а survival_rate – 0.54. Подальше навчання стабілізувало показники, на десятій епосі жертва жила в середньому 48.66 кроків із survival_rate 0.48, а після 15–20 епох survival_rate зростав до 0.58–0.64. На пізніших етапах, між 30-ю та 40-ю епохами, середній час виживання утримувався в межах 43–46 кроків, а survival_rate досягав 0.66–0.68. Таким чином, замість

гарантованого захоплення, характерного для baseline, хижак почав програвати від 32% до 52% епізодів. Динаміка порівняння baseline та прогнозової політики відображена на рис. 6.

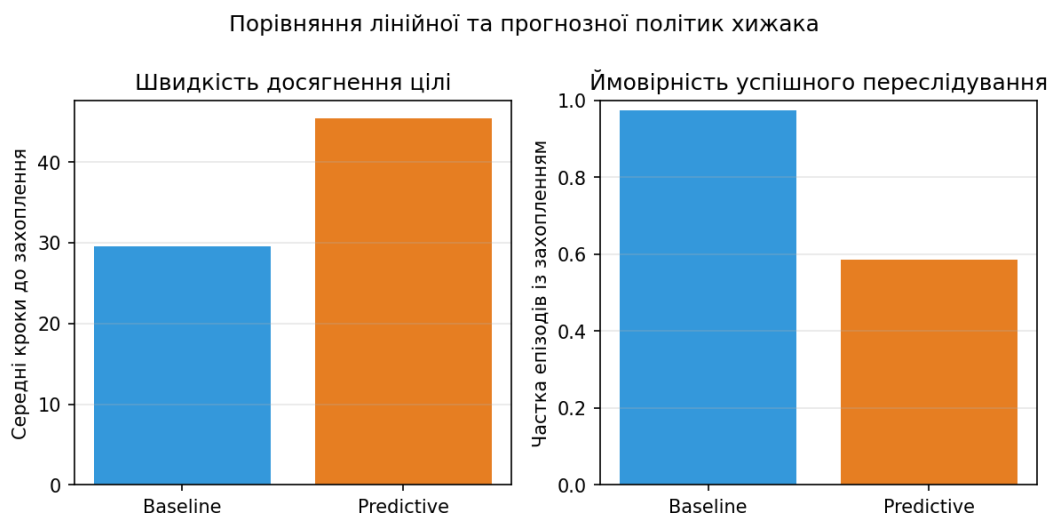


Рисунок 6. Порівняння baseline та прогнозової політики з точки зору жертви

Отримані результати свідчать про якісну трансформацію поведінки жертви. Збільшення середнього часу виживання приблизно у 1.7 раза означає, що агент більше не покладається на прямолінійну втечу, а формує складні, вигнуті, обережні траєкторії, які мінімізують ризик перехоплення. Зменшення частки успішних захоплень хижаком із 1.0 до приблизно 0.6 демонструє ефективність проактивної стратегії, жертва навчається обходити архітектурні “пастки” середовища та уникати напрямків, де хижак найбільш ймовірно скоротить дистанцію. Це свідчить про те, що доступ до прогнозу не просто покращує ухилення, а створює нову поведінкову якість – стратегічну та контекстно залежну, що значно підвищує шанси жертви на довготривале виживання.

4. ПОДАЛЬШІ ДОСЛІДЖЕННЯ

Подальші дослідження можуть бути спрямовані на формування комплексної системи прогнозування динаміки середовища, у якій жертва здатна не лише реагувати на поточний стан хижака, а й оцінювати можливі зміни конфігурації навколишнього простору. Перспективним є моделювання еволюції середовища з урахуванням не тільки майбутнього переміщення хижака, але й потенційної зміни топології простору, появи нових перешкод або небезпечних зон, а також формування областей підвищеної ймовірності перехоплення. Такий підхід забезпечить можливість оцінювати довгострокові сценарії та забезпечить жертві більш точне розуміння того, які ділянки середовища можуть стати вразливими у майбутньому. Важливим напрямом подальших робіт є інтеграція прогнозних сценаріїв безпосередньо у поведінкову стратегію ухилення. Це передбачає адаптацію траєкторії руху жертви відповідно до очікуваних загроз, динамічне перемикання між режимами швидкої втечі та тактичного уникнення, а також вибір оптимальних безпечних коридорів на основі багатокрокових прогнозів. Така поведінка дасть змогу підвищити ефективність ухилення у випадках, коли хижак застосовує складні або змінні стратегії переслідування. Подальші дослідження повинні включати розроблення механізмів контекстно-обумовленого реагування, коли жертва враховує не лише поточний стан переслідувача, а й прогноз ризику на кілька кроків наперед. Це дасть змогу формувати реакції, спрямовані на мінімізацію ймовірності перехоплення, а не лише на миттєве збільшення дистанції. У такому випадку поведінка жертви матиме ознаки

стратегічного мислення, що суттєво розширює межі класичних моделей ухилення. Перспективним є також практичне впровадження низки модулів, які можуть бути реалізовані студентом у рамках подальшої роботи. До таких модулів належать система прогнозування еволюції середовища, адаптивна логіка ухилення, що враховує прогнозні оцінки, механізми пріоритезації можливих сценаріїв розвитку подій та моделі балансування між виживанням і витратами енергії. Це дозволить створити більш гнучку архітектуру поведінки жертви, здатну автономно змінювати стратегію залежно від складності ситуації. У науковому аспекті подальші дослідження повинні зосереджуватися на застосуванні рекурентних нейронних мереж, зокрема архітектур LSTM і GRU, для моделювання складних і довготривалих траєкторій ухилення, а також на використанні графових нейромереж для опису взаємодії кількох хижаків або багатокомпонентних середовищ. Вагомим напрямом є використання трансформерів для прогнозування довгострокових патернів руху та інтеграція ймовірнісних сценаріїв у політики навчання з підкріпленням. Очікується, що такі підходи дозволять сформувати стратегії ухилення, які враховують не лише поточну ситуацію, але й можливі зміни у структурі середовища, забезпечуючи перевагу перед традиційними реактивними моделями.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Sutton R. S., Barto A. G. *Reinforcement Learning: An Introduction*. MIT Press, 2018. URL: <https://www.andrew.cmu.edu/course/10-703/textbook/BartoSutton.pdf>
2. Goodfellow I., Bengio Y., Courville A. *Deep Learning*. MIT Press, 2016. URL: <https://www.deeplearningbook.org/>
3. Battaglia P. W., Hamrick J. B., et al. *Relational Inductive Biases, Deep Learning, and Graph Networks*. arXiv:1806.01261, 2018. URL: <https://arxiv.org/abs/1806.01261>
4. Vaswani A., Shazeer N., et al. *Attention Is All You Need*. Advances in Neural Information Processing Systems (NeurIPS), 2017. URL: <https://arxiv.org/abs/1706.03762>
5. Barber D. *Bayesian Reasoning and Machine Learning*. Cambridge University Press, 2012. URL: <https://pzs.dstu.dp.ua/DataMining/bibl/Bayesian%20Reasoning%20and%20Machine%20Learning.pdf>
6. Kudenko D., Kazakov D., Alonso E. *Reinforcement Learning for Pursuit–Evasion Problems*. University of York Technical Report, 2001. URL: https://www.researchgate.net/publication/348085210_A_Deep_Reinforcement_Learning_Approach_for_the_Pursuit_Evasion_Game_in_the_Presence_of_Obstacles
7. Hoel C. *Automated Planning for Autonomous Vehicles Using Machine Learning*. Chalmers University of Technology, 2018. URL: https://www.researchgate.net/publication/354511114_Machine_Learning_-autonomous_Vehicles

ДЕЦЕНТРАЛІЗОВАНА КООРДИНАЦІЯ РОЮ ДРОНІВ ЧЕРЕЗ СТИГМЕРГІЮ І РОЙОВИЙ ІНТЕЛЕКТ БЕЗ GPS ЗВ'ЯЗКУ

Шостак В.О.¹, Петренко А.І.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ vitaliishh6@gmail.com, ² tolja.petrenko@gmail.com

У даній роботі представлено реалізацію децентралізованої системи координації рою безпілотних літальних апаратів (БПЛА) для пошуково–рятувальних операцій на основі алгоритму мурашиної колонії (ACO) та стигмергічної комунікації через віртуальні феромони. Система забезпечує автономне дослідження території без використання GPS, використовуючи локальну систему координат. Три автономних дрони координують свої дії через спільну феромонну карту, яка містить інформацію про досліджені зони, виявлені перешкоди та цільові об'єкти. Система реалізована з використанням PX4 SITL та демонструє емерджентну поведінку, коли прості правила окремих агентів призводять до складної колективної поведінки рою.

Ключові слова: ройовий інтелект, стигмергія, алгоритм мурашиної колонії, децентралізована координація, безпілотні літальні апарати, GPS–denied навігація, пошуково–рятувальні операції.

1. ВСТУП

Пошуково–рятувальні операції в умовах відсутності GPS-сигналу (GPS-denied environments) залишаються актуальною проблемою сучасної робототехніки. Традиційні централізовані системи керування роєм БПЛА мають низку обмежень: єдину точку відмови, обмежену масштабованість та високі вимоги до пропускну здатності каналів зв'язку. Природні системи, такі як мурашині колонії, демонструють ефективні механізми колективної координації без центрального керування через непряму комунікацію – стигмергію.

Ройовий інтелект (Swarm Intelligence) – це колективна поведінка децентралізованих, самоорганізованих систем, які складаються з множини агентів, що взаємодіють локально між собою та з середовищем [1]. Ключовою особливістю таких систем є емерджентна поведінка – виникнення складних глобальних патернів з простих локальних правил.

Стигмергія – це механізм непрямої комунікації через модифікацію середовища. У мурашиних колоніях стигмергія реалізується через хімічні феромони, які мурахи залишають на своєму шляху і які поступово випаровуються з часом [2].

Алгоритм мурашиної колонії (Ant Colony Optimization, ACO) моделює поведінку мурах при пошуку найкоротшого шляху від мурашника до джерела їжі. Алгоритм базується на двох ключових механізмах: відкладання феромонів агентами та випаровування феромонів з часом [3].

Метою даного дослідження є розробка та реалізація децентралізованої системи координації рою БПЛА на основі ACO з використанням стигмергічної комунікації через віртуальні феромони для автономного пошуку об'єктів у GPS-denied середовищі.

2. АРХІТЕКТУРА СИСТЕМИ

Розроблена система складається з трьох основних компонентів, реалізованих як окремі ROS2 ноди для кожного БПЛА:

- Нода феромонної карти (pheromone_map_node) – відповідає за відкладання та випаровування віртуальних феромонів.
- Нода прийняття рішень ACO (aco_decision_node) – реалізує алгоритм вибору наступної точки маршруту.
- Нода уникнення перешкод (obstacle_avoidance_node) – забезпечує детектування перешкод через LIDAR.

Віртуальна феромонна карта представлена трьома окремими шарами розміром 40×40 клітинок з роздільною здатністю 5 метрів на клітинку:

- Шар explored (τ_e): відмічає досліджені зони, діапазон [0, 10].
- Шар hazard (τ_h): відмічає виявлені перешкоди, діапазон [0, 10].
- Шар target (τ_t): відмічає цільові об'єкти, діапазон [0, 100].

Кожен БПЛА публікує свої дані в феромонну карту через топік /swarm/pheromone_map.

3. МАТЕМАТИЧНА МОДЕЛЬ ACO

3.1. Відкладання феромонів

Кожні 0.2 секунди дрон додає феромон в поточну клітинку згідно формули (1):

$$\tau_e(i, j, t) = \min(\tau_e(i, j, t - 1) + \Delta\tau, \tau_{\max}), \quad (1)$$

де $\tau_e(i, j, t)$ – значення феромону explored в клітинці (i, j) в момент часу t;

$\Delta\tau = 1.0$ – величина приросту феромону;

$\tau_{\max}$ – максимальне значення феромону (10.0 для шарів explored та hazard, 100.0 для шару target).

3.2. Випаровування феромонів

Феромони поступово випаровуються з коефіцієнтом $\rho = 0.05$ (5% за секунду) згідно формули (2):

$$\tau(i, j, t + 1) = \tau(i, j, t) \times (1 - \rho), \quad (2)$$

де ρ – коефіцієнт випаровування (0.05);

t – дискретний час в секундах.

Це забезпечує поступове «забування» старих слідів та стимулює дослідження нових зон.

3.3. Вимірювання феромону

Загальна привабливість клітинки (i, j) обчислюється згідно формули (3):

$$\varphi(i, j) = b_e \times b_h + b_t \quad (3)$$

де $b_e = 1.0 - \min(\tau_e(i, j) / 10.0, 0.9)$ – бонус дослідження;

$b_h = 1.0 - \min(\tau_h(i, j) / 10.0, 0.95)$ – штраф небезпеки;

$b_t = \tau_t(i, j) \times 10.0$ – бонус цілі.

Бонус дослідження b_e стимулює пошук нових зон, штраф небезпеки b_h знижує привабливість зон з перешкодами, а бонус цілі b_t різко підвищує привабливість клітинок з виявленими об'єктами.

3.4. Генерація кандидатів

Для кожного кроку генерується $N = 20$ кандидатних точок навколо поточної позиції згідно формули (4):

$$c_k = p + d(\cos(\psi)\cos(\theta), \cos(\psi)\sin(\theta), \sin(\psi)), \quad (4)$$

де p – поточна позиція дрона;
 $d = 15$ м – розмір кроку;
 $\theta \sim U(0, 2\pi)$ – випадковий азимутальний кут;
 $\psi \sim U(-\pi/6, \pi/6)$ – випадковий кут нахилу.

3.5. Ймовірнісний вибір

Для кожного кандидата обчислюється score з використанням АСО формули (5):

$$s(c_k) = [\varphi(c_k)]^\alpha \times [\eta(c_k)]^\beta, \quad (5)$$

де $\varphi(c_k)$ – привабливість клітинки згідно формули (3);
 $\eta(c_k)$ – евристична функція;
 $\alpha = 1.5$ – параметр впливу феромону;
 $\beta = 3.0$ – параметр впливу евристики.
Ймовірність вибору кандидата обчислюється через нормалізацію згідно формули (6):

$$P(c_k) = \frac{s(c_k)}{\sum_j s(c_j)}, \quad (6)$$

де $\sum_j s(c_j)$ – сума score всіх кандидатів.

Вибір здійснюється стохастично відповідно до розподілу $P(c_k)$, що дозволяє уникнути локальних мінімумів.

4. РЕАЛІЗАЦІЯ СИСТЕМИ

4.1. Локальна система координат

Система працює без GPS, використовуючи локальну декартову систему координат з початком в точці старту $(0, 0, 0)$. Позиція кожного БПЛА отримується з топіку `/mavros/local_position/pose`, який в симуляції надає ідеальну позицію від Gazebo Ground Truth. У реальних умовах джерелом позиції може бути Visual-Inertial Odometry (VIO), Motion Capture System або LIDAR SLAM.

4.2. Детектування перешкод

Нода `obstacle_avoidance_node` отримує дані від LIDAR (топик `/uav{N}/scan`) та публікує булеве значення в топик `/uav{N}/obstacle_avoidance/obstacle_detected` при виявленні перешкоди на встановленій відстані. При детектуванні перешкоди, `pheromone_map_node` додає hazard феромон зі значенням $+5.0$ в поточну клітинку згідно формули.

4.3. Комунікаційна архітектура

Система використовує два механізми комунікації між дронами: стигмергічна комунікація через топик `/swarm/pheromone_map` (кожен дрон публікує свої дані в феромонну карту) та пряма комунікація через топик `/swarm/tent_detected` для миттєвого повідомлення про виявлення цілі.

5. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Проведено серію експериментів з роєм з 3 дронів в симульованому середовищі Gazebo. На початковій фазі дрони розлітаються в трьох напрямках (0° , 120° , 240°), забезпечуючи рівномірне початкове покриття території (рис. 1).

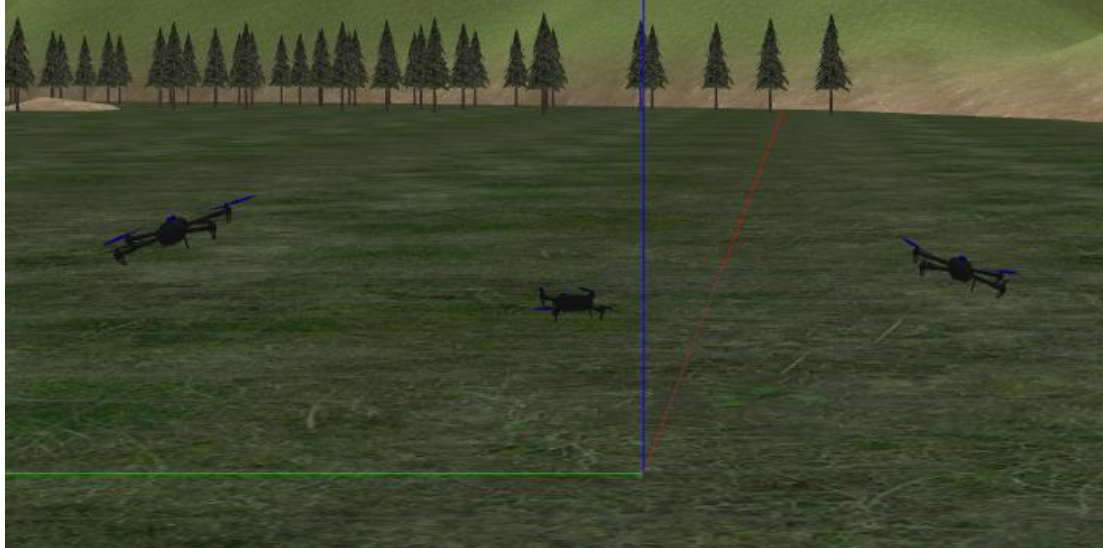


Рисунок 1. Політ в різні точки для дослідження території

Аналіз феромонних слідів показує ефективне уникнення досліджених зон. Феромони накопичуються до 10.00 (максимум) протягом деякого часу для досліджених клітинок, після чого дрон переміщується в нову зону. Випаровування згідно формули (2) з коефіцієнтом $\rho = 0.05$ забезпечує зменшення explored з 10.0 до 5.0 за близько 14 секунд, дозволяючи повторне відвідування зон при необхідності. Як бачимо на рисунку 2 це необхідно для уникнення зависання на місці, коли дрон буде перетинати точку, яку вже пролітав.

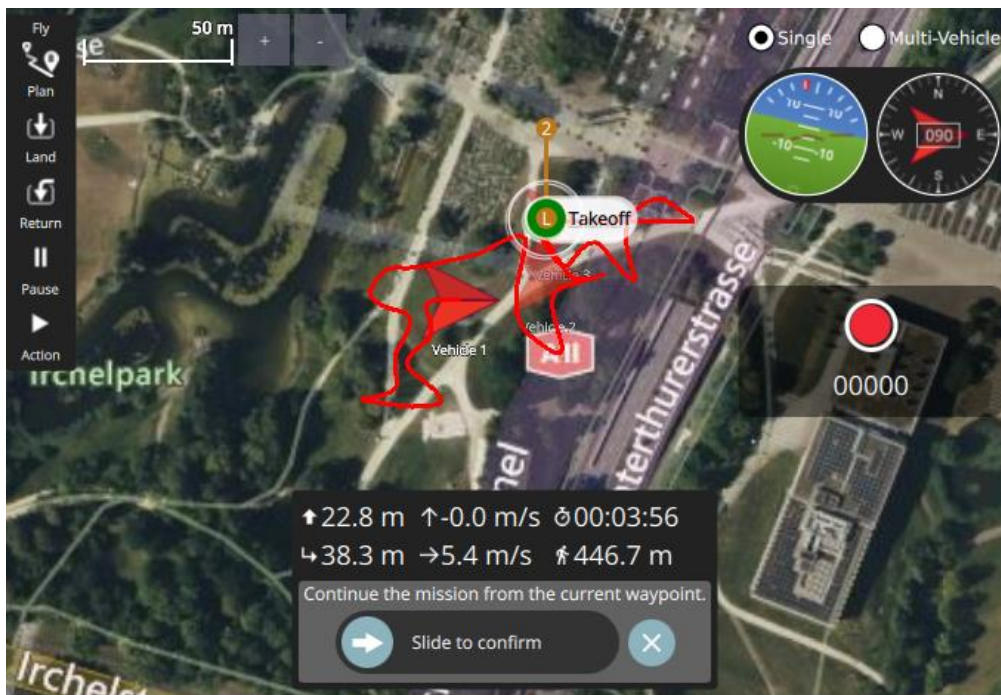


Рисунок 2. Приклад маршрута дрона при дослідженні території

При виявленні цілі спостерігається швидка конвергенція всіх дронів. Коли виявлена ціль, то до позиції додається target феромон зі значенням 100.0, а також публікується повідомлення в топик /swarm/tent_detected. Інші дрони миттєво отримують координати та починають політ до цілі, демонструючи ефективну колективну поведінку (рис. 3).

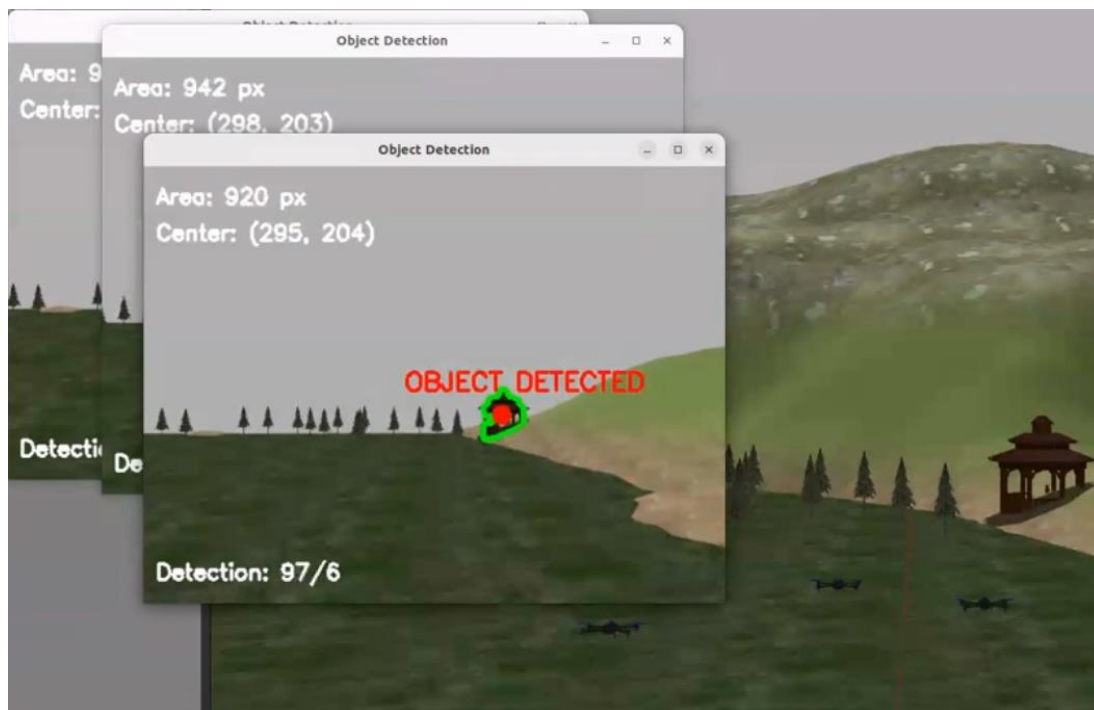


Рисунок 3. Поведінка роя дронів після виявлення цілі

6. ВИСНОВКИ

У роботі представлено повнофункціональну реалізацію децентралізованої системи координації рою БПЛА на основі алгоритму мурашиної колонії та стигмергічної комунікації. Основні результати дослідження:

- Розроблено математичну модель віртуальних феромонів з трьома шарами та механізмом випаровування.
- Реалізовано децентралізовану архітектуру на базі ROS2 з трьома основними нодами для кожного дрона.
- Реалізовано стигмергічну координацію – дрони самостійно розподіляють зони пошуку.
- Визначено оптимальний коефіцієнт випаровування, що забезпечує зниження часу виявлення.
- Продемонстровано працездатність системи в GPS–denied середовищі з використанням локальних координат.

Система демонструє ключові властивості ройового інтелекту: децентралізацію, самоорганізацію, масштабованість та емерджентну поведінку. Відсутність єдиної точки відмови робить систему стійкою до виходу з ладу окремих агентів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Swarm intelligence. Wikipedia: веб-сайт. URL: https://en.wikipedia.org/wiki/Swarm_intelligence
2. Stigmergy. Wikipedia: веб-сайт. URL: <https://en.wikipedia.org/wiki/Stigmergy>

3. Ant colony optimization algorithms. Wikipedia: веб-сайт. URL: https://en.wikipedia.org/wiki/Ant_colony_optimization_algorithms
4. Dorigo, M., Stützle, T. (2019). Ant Colony Optimization: Overview and Recent Advances. In: Gendreau, M., Potvin, JY. (eds) Handbook of Metaheuristics. International Series in Operations Research & Management Science, vol 272. Springer, Cham.
5. Roggi, G., Meraglia, S., & Lovera, M. (2022). Leonardo Drone Contest 2021: Politecnico di Milano team architecture. In 2022 IEEE International Conference on Unmanned Aircraft Systems (ICUAS) (pp. 191–196). <https://doi.org/10.1109/ICUAS54217.2022.9836168>
6. Soria, E., Schiano, F., & Floreano, D. (2020). SwarmLab: A MATLAB drone swarm simulator. In 2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS) (pp. 8005–8011). <https://doi.org/10.1109/IROS45743.2020.9341677>
7. Adams, J. A., & Humphrey, C. M. (2024). Human–swarm interaction in complex urban environments: Insights from DARPA OFFSET. IEEE Transactions on Field Robotics, 41(3), 123–135. <https://doi.org/10.1002/rob.22345>
8. Swarmer (2024). <https://www.getswarmer.com/>
9. Jun Tang; Gang Liu; Qingtao Pan. Review on Representative Swarm Intelligence Algorithms for Solving Optimization Problems: Applications and Trends. IEEE/CAA Journal of Automatica Sinica (Volume: 8, Issue: 10, October 2021), pp. 1627 – 1643. DOI: 10.1109/JAS.2021.1004129
10. Luong Vuong Nguyen. Swarm Intelligence–Based Multi–Robotics: A Comprehensive Review. AppliedMath 2024, 4(4), 1192–1210; <https://doi.org/10.3390/appliedmath4040064>
11. Wei W, He X, Wang X, Wang M. Research on swarm munitions cooperative warfare. In: International Conference on Autonomous Unmanned Systems. Singapore: Springer; 2021. pp. 717–727.

PROPOSED ARCHITECTURE FOR PRODUCT REVIEW SUMMARISATION: A MODULAR PIPELINE FOR ASPECT-BASED INSIGHT EXTRACTION AND SYNTHESIS

Yaroslav Panasiuk

National Technical University of Ukraine
"Igor Sikorsky Kyiv Polytechnic Institute", Kyiv, Ukraine

voice2of2reason@gmail.com

In the domain of e-commerce, user reviews are a critical resource for decision-making. However, the sheer volume and unstructured nature of these reviews present a significant challenge for consumers seeking actionable insights. Traditional extractive summarization methods often fail to capture nuance, while end-to-end Large Language Model (LLM) approaches struggle with hallucinations and lack of structural control. This paper proposes a novel, linear, multi-stage architecture that transforms unstructured text into the Quantified Aspect-Based Summary (QABS). The proposed pipeline utilizes a modular approach, integrating Coreference Resolution, Low-Rank Adaptation (LoRA) fine-tuned models for tuple extraction, and dynamic topic modeling. By decomposing reviews into atomic insights and re-synthesizing them using blueprint prompting, this architecture ensures high clarity, trust, verifiability, and quantification of user sentiment.

Keywords: Aspect-Based Sentiment Analysis (ABSA), Large Language Models, LoRA, Coreference Resolution, Dynamic Topic Modeling, Product Summarization.

1. INTRODUCTION

The proliferation of digital marketplaces has democratized product feedback, allowing users to share detailed experiences. However, this abundance of data has created an information overload. The Quantified Aspect-Based Summary (QABS) is defined not merely as a shortened version of text, but as a structured, unbiased aggregation that enables the reader to make confident purchasing decisions.

To achieve this standard, a summary must demonstrate domain expertise, articulate trade-offs, aggregate sentiment across categories, and provide trustworthy verification through quotes. It must also quantify sentiment (e.g., "75% of users mentioned battery life") rather than relying on vague qualitative descriptors.

Current state-of-the-art approaches often rely on monolithic "black box" prompting of LLMs. While powerful, these approaches frequently fail to strictly adhere to the structural and quantitative requirements of a high-quality review summary. They may miss colloquial references or hallucinate features that do not exist.

This paper presents a proposed architecture that avoids the single prompt paradigm in favor of specialized, modular components. This pipeline systematically cleans, decomposes, clusters, selects, and synthesizes data to produce summaries that are scannable, actionable, and mathematically grounded in user sentiment.

2. THE QUANTIFIED ASPECT-BASED SUMMARY (QABS) STANDARD

To evaluate the efficacy of summarization architectures, we define a target standard known as the Quantified Aspect-Based Summary (QABS). Unlike generic text reduction, the QABS is defined as a structured, unbiased, and actionable aggregation of user feedback designed to enable confident purchasing decisions.

To meet the QABS standard, the output must satisfy seven strict criteria:

1. **Trust & Verifiability:** To mitigate LLM hallucinations, the summary must include direct, representative quotes for key claims, ensuring all insights are attributable to source data.
2. **Quantification & Weighting:** The system must quantify sentiment (e.g., "75% of users mentioned battery life") rather than relying on vague qualitative descriptors, providing a measurable "Feature Score Breakdown".
3. **Actionability:** The summary must articulate core trade-offs (Pros/Cons) and provide explicit guidance on "Who should buy it?" based on user profiles.
4. **Clarity & Domain Expertise:** It must demonstrate specialized knowledge (e.g., specific vocabulary like "mechanical vs. membrane keyboard").
5. **Coverage & Categorization:** It must aggregate sentiment across all relevant categories, including defects, performance, and value.
6. **Experience Contextualization:** It must capture longitudinal data, such as durability over time, under a "Usage & Experience Sharing" category.
7. **Structure & Scannability:** The output must use specific formatting (headings, tables, bullet points) to maximize information retrieval speed.

3. PROBLEM DEFINITION

The core problem is transforming high-volume, noisy, unstructured natural language into a structured, verifiable summary. Several specific sub-problems exist within this domain:

1. **Colloquial Ambiguity:** User reviews are often informal. A sentence like "It is fast" relies on context to understand what "It" refers to. Standard extractors often fail to link the pronoun to the specific product feature.
2. **Lack of Granularity:** Generic summarizers often conflate distinct features. QABS requires atomic separation of concepts (e.g., separating "Screen Brightness" from "Screen Refresh Rate") to provide specific pros and cons.
3. **Static Taxonomy Limitations:** Pre-defined categories fail to capture emerging or product-specific issues. A robust system must employ dynamic topic modeling to find common themes without a pre-set taxonomy.
4. **Hallucination and Verifiability:** Generative models may invent details. A trustworthy summary requires "Trust & Verifiability," linking claims directly to representative user quotes.
5. **Quantification:** Users require weightings (e.g., popularity of an opinion) to judge the severity of a defect or the value of a pro.

4. EXISTING APPROACHES

4.1. Extractive summarization

Traditional extractive methods select representative sentences from the source text. While this preserves the original voice, it fails to synthesize diverse opinions into a cohesive narrative or quantify the prevalence of a specific sentiment across thousands of reviews.

4.2. End-to-end LLM prompting

Recent approaches utilize large context windows in models like GPT-4 or Gemini to process all reviews in a single pass. While this generates fluent text, it suffers from the "Lost in the Middle" phenomenon [3] and struggles with accurate quantification. It acts as a "black box," making it difficult to debug why specific insights were prioritized or ignored.

4.3. Aspect-based sentiment analysis (ABSA)

Standard ABSA [5] identifies aspects and sentiments but often lacks the synthesis capabilities to generate a readable narrative or the logic to resolve complex coreferences in colloquial text.

5. PROPOSED APPROACH

We propose a linear, multi-stage pipeline architecture. This system transforms raw reviews into structured data (atomic statements/clusters) and finally back into synthesized text.

5.1. Stage 1: input filtering

The first stage ensures data quality. Automated filters are applied to the raw stream to remove spam, profanity, and fraudulent reviews, ensuring that the downstream analysis is based on genuine user experiences.

5.2. Stage 2: insight extraction

This stage decomposes complex reviews into distinct, atomic insights. An insight is defined as a standardized natural language statement confined to a single topic and sentiment.

Coreference Resolution: To address colloquial ambiguity, we employ a Coreference Resolution model (e.g., SpanBERT [1]) prior to extraction. This resolves pronouns, transforming "It is fast" into "The S24 is fast," ensuring the extractor correctly attributes the sentiment.

Tuple Extraction via LoRA: We utilize a 7B-parameter model (e.g., Llama-3 [4] or Mistral) fine-tuned with Low-Rank Adaptation (LoRA [2]) on ABSA datasets. This model extracts a JSON object for every distinct opinion, including Aspect, Sentiment, Intensity, Descriptor, and the direct Quote for verification.

5.3. Stage 3: dynamic topic modeling

Rather than forcing reviews into rigid categories, this phase aggregates atomic statements to discover organic themes.

Generating Topic Names: A secondary LoRA model analyzes the extracted insights to generate dynamic, descriptive labels.

Deduplication and Clustering: The system cleans the topics using two methods:

Semantic Embeddings: Insights are projected into vector space. Dense clusters are identified as topics, recognizing that "cannot log in" and "authentication error" are semantically identical.

Pattern Matching: Catches slight variations in spelling to ensure a concise topic list.

5.4. Stage 4: topic selection

Not all extracted topics are relevant for a summary. The architecture selects topics based on a weighted multi-criteria decision matrix:

1. **Relevance:** Topics are filtered to prioritize "App Experience" over "Out-of-App Experience" (e.g., ignoring food quality in a delivery app review).
2. **Popularity:** Calculated via cluster density: $(\text{Mentions} / \text{Total Reviews}) * 100$.
3. **Freshness:** To ensure the summary reflects the current state of the product [6], recent insights are weighted higher using the decay formula:

$$W_t = W_0 * e^{-\lambda t},$$

where t is the age of the review.

4. Net Sentiment: Positives / Volume

5. Controversy Score: High variance in sentiment indicates a "Polarizing Feature."

5.5. Stage 5: insight selection

Representative Sampling: The system selects insights closest to the topic cluster's centroid to ensure the summary reflects the consensus view.

5.6. Stage 6: summary generation

The Blueprint Prompt: The system constructs a detailed "Blueprint." This prompt explicitly instructs the model on structure: "Write a section titled '## Pros'. Use the following 3 structured insights... strictly use the statistic '75%'... Use Quote ID".

This methodology directly satisfies the QABS requirements by enforcing Structure, Scannability, and Quantification.

6. CONCLUSION

The proposed architecture represents a shift from unstructured stochastic summarization to a structured, deterministic pipeline. By integrating SpanBERT for coreference resolution, LoRA-tuned models for tuple extraction, and a mathematical approach to topic freshness and popularity, this system produces reviews that are actionable, trustworthy, and verifiable. This approach allows consumers to navigate the complexities of product feedback with the same clarity found in expert reviews, democratizing access to high-quality market intelligence.

REFERENCES

1. Joshi, M., Chen, D., Liu, Y., Weld, D. S., Zettlemoyer, L., & Levy, O. (2020). SpanBERT: Improving Pre-training by Representing and Predicting Spans. *Transactions of the Association for Computational Linguistics*, 8, 64–77.
2. Hu, E. J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., & Chen, W. (2021). LoRA: Low-Rank Adaptation of Large Language Models. *arXiv preprint arXiv:2106.09685*.
3. Liu, N. F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2024). Lost in the Middle: How Language Models Use Long Contexts. *Transactions of the Association for Computational Linguistics*, 12, 157–173.
4. AI @ Meta. (2024). The Llama 3 Herd of Models. *arXiv preprint arXiv:2407.21783*.
5. Pontiki, M., Galanis, D., Pavlopoulos, J., Papageorgiou, H., Androutsopoulos, I., & Manandhar, S. (2014). SemEval-2014 Task 4: Aspect Based Sentiment Analysis. *Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval 2014)*.
6. Manning, C. D., Raghavan, P., & Schütze, H. (2008). *Introduction to Information Retrieval*. Cambridge University Press. (Specifically Chapter 6 on Scoring and Term Weighting).

АКУСТИЧНА КЛАСИФІКАЦІЯ БПЛА В УМОВАХ ДОМЕННОГО ЗСУВУ ТА ДЕФІЦИТУ ДАНИХ

Вільгурін В.І.¹, Сидорський В.С.², Данилов В.Я.³

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ vvilgurin@gmail.com, ² volodymyr.syderskyi@gmail.com, ³ danilov1950@ukr.net

Зростання кількості та доступності безпілотних літальних апаратів (БПЛА) створює нові виклики для безпеки цивільної та критичної інфраструктури. Одним з перспективних напрямів є акустичне виявлення БПЛА, оскільки мікрофони є дешевими, пасивними сенсорами і можуть розгортатися на великій території. Водночас практичне застосування таких систем часто ускладнюється відмінностями у фізичних властивостях мікрофонів, різноманіттю умов навколишнього середовища та обмеженим обсягом анотованих цільових даних. У роботі розглянуто підхід до побудови моделі акустичної класифікації БПЛА, стійкої до доменного зсуву, що виникає внаслідок відмінностей у мікрофонах та локаціях і дефіциту маркованих прикладів. Запропонована система поєднує попередню обробку сигналу на основі мел-спектрограм, використання попередньо натренованої згорткової мережі як екстрактора ознак, агрегацію на основі механізму уваги та стратегії адаптації домену із застосуванням аугментацій та збалансованого відбору даних.

Ключові слова: акустична детекція, БПЛА, ConvNeXt, доменний зсув, Mixup, SED.

1. ВСТУП

Акустичне виявлення безпілотних літальних апаратів є важливим компонентом сучасних систем спостереження та безпеки. На відміну від радіолокаційних чи оптичних підходів, акустичні системи не потребують прямої видимості цілі та можуть працювати в умовах складного рельєфу або задимлення, а також не можуть бути легко виявлені в силу пасивної природи роботи.

Водночас вони є чутливими до типу мікрофона, акустики сцени та рівня фонового шуму. На практиці це призводить до доменного зсуву: дані, зібрані на різних пристроях або в інших умовах, мають інший розподіл, ніж той, на якому навчалася модель, а також, потенційно, інший набір класів. Крім того, зібрати велику кількість якісних маркованих записів БПЛА в реалістичних умовах часто неможливо через високу складність анотації в офлайн-режимі та високу вартість експериментів.

Таким чином система повинна ефективно використовувати інформацію з наявних даних для адаптації до нових доменів та нових класів. Метою роботи є розробка акустичної системи класифікації БПЛА, яка:

- зберігає високу точність у присутності доменного зсуву між мікрофонами та акустичними сценами;
- враховує сильний дисбаланс класів (більшість записів-природний шум, меншість-БПЛА);

- залишається обчислювально ефективною для застосування в режимі, близькому до реального часу.

2. ПОСТАНОВКА ЗАДАЧІ ТА ДАНІ

Розглядається задача класифікації аудіозаписів на основі коротких фрагментів сигналу. Вхідними даними є хвильові форми, записані щонайменше двома типами мікрофонів, що відрізняються частотною характеристикою, рівнем власного шуму та акустичними умовами. Розподіл даних наведено в таблиці 1.

Таблиця 1. Розподіл даних за доменами та класами

Домен (Набір даних)	Шум	Не БПЛА	БПЛА	Усього
Мікрофон 1 (Train)	299 603	17 981	0	317 584
Мікрофон 2 (Train)	2 519	0	441	2 960
Мікрофон 1 (Val)	43 551	2 855	0	46 406
Мікрофон 2 (Val)	2 797	0	4 207	7 004
Тестовий набір	343	0	805	1 148

Спостерігається значний дисбаланс між кількістю наявних аудіозаписів на рівні мікрофонів та на рівні класів, різні акустичні сцени та обмежений обсяг маркованих прикладів у цільовому домені (рис. 1).

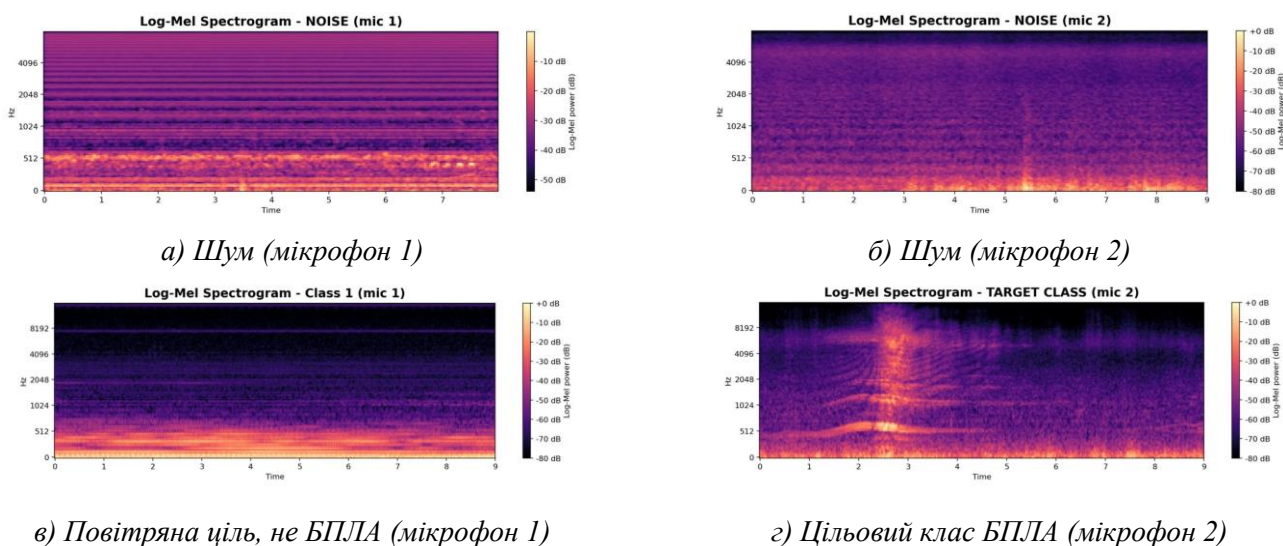


Рисунок 1. Порівняння спектральних характеристик сигналів

Необхідно побудувати модель f , яка на основі часово-частотного представлення аудіосигналу прогнозує ймовірність належності до класу БПЛА, залишаючись стійкою до зміни домену та використовуючи обмежену кількість маркованих записів з нового середовища.

3. МЕТОДИ ТА АРХІТЕКТУРА МОДЕЛІ

Попередня обробка. Для побудови ознак використано перетворення звукових хвиль з частотою дискретизації 32 кГц у мел-спектрограми. Аудіо ділиться на фрагменти тривалістю 9 с, до яких застосовується короткочасне перетворення Фур'є (STFT) з параметрами $n_{fft} = 2048$ та $hop = 512$, після чого спектр перетворюється у мел-шкалу [1]. Така попередня обробка забезпечує достатню частотну роздільну здатність для гармонік цільових сигналів.

Модель. У ролі енкодера використано згорткову мережу ConvNeXt-Tiny [2], попередньо натреновану на великому наборі даних з зображеннями. Верхні шари ConvNeXt виступають

екстрактором спектрально-часових ознак. Сформовані ознаки подаються до спеціалізованої голови типу SED (Sound Event Detection) [3], реалізованої як уваговий модуль. На першому етапі відбувається агрегація по частоті за допомогою GeM-pooling, яке дозволяє більш гнучко підлаштовуватись до різного спектрального розподілу сигналів з двох доменів. Далі часові вектори обробляються рекурентним шаром GRU, що використовується для моделювання темпоральних змін у спектрограмі. Врешті, отримані ознаки зважуються механізмом уваги по часовій осі для визначення важливості кожної з частин зображення для покращеного узагальнення у разі виникнення високоамплітудного шуму.

Доменна адаптація. Для підвищення стійкості до доменного зсуву застосовано комбінацію аугментацій на рівні хвильової форми та спектрограм. На рівні сигналу виконуються нормалізація амплітуди, випадкове інвертування сигналу в часі, gandom gain та time-stretch. Також виконується змішування шумових записів з усіма іншими прикладами (Mixup [4]):

$$x_{mix} = \lambda x_i + (1 - \lambda) x_j, \quad y_{mix} = y_i \quad (1)$$

де $\lambda \sim \text{Beta}(\alpha, \alpha)$. При цьому для зразків з мікрофона цільового домену не застосовуються перетворення, які суттєво змінюють частотну характеристику спектрограми, через брак даних цільового домену. Маркування завжди залишається «жорстким»: усі змішані приклади вважаються тим самим класом, що й основний сигнал, за винятком випадку, коли шум другого мікрофона змішується з шумом першого – тоді зразок позначається як шум.

На рівні спектрограм застосовуються маскування часу та частоти (SpecAugment), що визначається як добуток Адамара спектрограми $S(f, t)$ та бінарних масок M :

$$\tilde{S}(f, t) = S(f, t) \odot M_f(f) \odot M_t(t), \quad (2)$$

а також невеликий випадковий зсув, що додатково зменшує перенавчання на конкретні артефакти домену та випадковий шум.

Тренування. Навчання виконується з використанням категорійної крос-ентропійної функції втрат. Для боротьби з дисбалансом класів застосовується зважений семплінг: ймовірність вибору зразків різних класів масштабується через квадратний корінь від їхньої частоти, що пом'якшує домінування шуму та дозволяє моделі краще навчатися на рідкісних корисних класах і цільовому класі БПЛА. Така комбінація фронтенду на мел-спектрограмах, ConvNeXt бекбону, увагової SED-голови та спеціалізованих аугментацій робить модель придатною для роботи в умовах доменного зсуву та дефіциту маркованих даних цільового домену.

4. ЕКСПЕРИМЕНТИ ТА РЕЗУЛЬТАТИ

Для оцінки ефективності запропонованого підходу було проведено серію експериментів. Було здійснено оцінку впливу різної кількості класів для навчання та методів аугментації на здатність моделі до узагальнення на тестовій вибірці (таблиця 2).

Таблиця 2. Результати експериментів на тестовій вибірці

Експеримент	Precision	Recall	F1 Score
Baseline (Всі дані)	0.989	0.908	0.947
Всі дані + Mixup	0.954	0.958	0.956
Тільки Мікрофон 2	0.989	0.810	0.891
Pre-train + Fine-tune	0.995	0.903	0.947

1. Навчання лише на цільовому домені (Мікрофон 2). У першому експерименті модель навчалася лише на обмеженому наборі даних з другого мікрофона із застосуванням базових аугментацій. Результати виявилися незадовільними (Recall 0.810) порівняно з іншими підходами. Аналіз показує, що модель схильна до перенавчання, орієнтуючись на прості ознаки, такі як загальний баланс високих та низьких частот, замість вивчення специфічних гармонійних патернів БПЛА.

2. Об'єднання доменів (Мікрофон 1 + Мікрофон 2). Додавання великого масиву шумових даних з першого мікрофона покращило метрики якості. Проте аналіз карт активації (GradCAM) продемонстрував, що модель все ще часто фокусується на низькочастотних областях як для класу шуму, так і для класу сигналу, що знижує надійність детекції в умовах сильного зашумлення.

3. Об'єднання доменів із застосуванням Міхур. Найкращі результати (F1 0.956) були отримані при використанні техніки Міхур, де шуми з другого мікрофона змішувалися з цільовими сигналами та сигналами з мікрофона 1. Модель вчилася шукати корисні ознаки не тільки серед низькочастотного шуму, а й у середньому та високому частотних діапазонах та виокремлювати ознаки, специфічні для кожного з класів.

4. Попереднє навчання з донавчанням на цільовому домені (Pre-training+Fine-tuning). В рамках цього експерименту було застосовано стратегію трансферного навчання. На етапі Pretraining модель навчалася на повному об'єднаному наборі даних з Міхур, що дозволило вивчити робастні ознаки. На етапі Fine-tuning модель донавчалася виключно на даних цільового домену. Цей підхід забезпечив високу точність та адаптацію до умов цільового мікрофона. Однак, донавчання на малій вибірці призвело до того, що модель стала більш впевненою на простих прикладах, але надмірно консервативною на складних або зашумлених зразках.

Аналіз передбачень. Детальний аналіз помилок моделі дозволив виявити критичну проблему розбіжності між тренувальним та тестовим наборами даних (рис. 2). Тренувальна вибірка складається переважно із записів БПЛА на близькій відстані до мікрофону, які характеризуються широким спектральним складом та чіткими високочастотними гармоніками. Натомість тестовий набір містить записи віддалених дронів. Через фізичні властивості, у сигналах віддалених об'єктів високочастотні компоненти швидко загасають, залишаючи переважно низькочастотні артефакти. Оскільки в процесі навчання модель бачила низькочастотну активність переважно у класі “шум”, а сигнал БПЛА асоціювала з наявністю вищих частот, вона не здатна ефективно відрізнити віддалений дрон від фонового шуму. Це свідчить про необхідність розширення вибірки реальними записами віддалених цілей для покращення результатів.

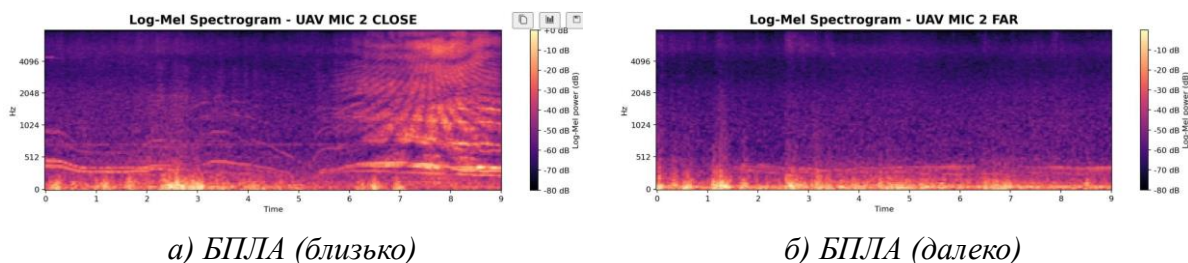


Рисунок 2. Порівняння спектрограм БПЛА у залежності від відстані до мікрофона

5. ВИСНОВКИ

У роботі досліджено проблему акустичної класифікації БПЛА в умовах обмежених даних цільового домену та використання різнотипного обладнання. Експерименти показали, що

навчання виключно на малій вибірці цільового мікрофона призводить до перенавчання на поверхневі ознаки та низької здатності до узагальнення.

Найкращу ефективність (F1-score = 0.956) продемонстрував підхід із глобальним навчанням на об'єднаних даних та використанням Міхур. Ця стратегія змушує модель ігнорувати низькочастотний шум і фокусуватися на стійких спектральних гармоніках БПЛА. Критичним викликом залишається фізична відмінність між тренувальними (ближніми) та тестовими (віддаленими) сигналами, де втрата високих частот ускладнює детекцію. Подальша робота буде зосереджена на вдосконаленні методів попередньої обробки для покращення детекції віддалених цілей.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Qiuqiang Kong та ін. “PANNS: Large-Scale Pretrained Audio Neural Networks for Audio Pattern Recognition”. В: *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 28(2020), с. 2880–2894.
2. Zhuang Liu та ін. “A ConvNet for the 2020s”. В: *arXiv preprint arXiv:2201.03545* (2022).
3. Krishna Kumar Singh та Yong Jae Lee. “Hide-and-Seek: Forcing a Network to be Meticulous for Weakly-supervised Object and Action Localization”. В: *arXiv preprint arXiv:1710.02998* (2017).
4. Hongyi Zhang та ін. “mixup: Beyond Empirical Risk Minimization”. В: *arXiv preprint arXiv:1710.09412* (2017).

ПРУНІНГ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ ЗА ДОПОМОГОЮ ІНТЕРПРЕТОВАНOSTІ МЕРЕЖ КОЛМОГОВОРА-АРНОЛЬДА

Єфанов І.С.¹, Шаповал Н.В.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ efanov.illiya@lil.kpi.ua

Запропоновано новий метод структурного прунінгу згорткових нейронних мереж, заснований на аналізі важливості ознак через інтерпретовані мережі Колмогорова-Арнольда (KAN). Метод використовує KAN-шар як вузьке місце для ідентифікації надлишкових фільтрів у CNN. Розроблено критерій важливості на основі норми коефіцієнтів В-сплайнів та алгоритм ітеративного прунінгу з донавчанням. Експериментальна валідація на CIFAR-10 показала перевагу над класичними методами: досягнуто 94.2% точності при стисненні моделі на 34%, що на 1.4% краще за magnitude-based pruning та на 0.7% краще за knowledge distillation.

Ключові слова: прунінг нейронних мереж, мережі Колмогорова-Арнольда, стиснення моделей, інтерпретованість, глибоке навчання, В-сплайни.

1. ВСТУП

Сучасні згорткові нейронні мережі досягають високої точності у задачах комп'ютерного зору, але їхнє розгортання на пристроях з обмеженими ресурсами (Edge AI) залишається складною проблемою. Моделі, такі як ResNet-50 або VGG-16, містять десятки мільйонів параметрів та вимагають гігабайти пам'яті, що робить їх непрактичними для мобільних пристроїв, IoT систем та автономного обладнання [1].

За оцінками, тренування великих мовних моделей може коштувати мільйони доларів та споживати енергію, еквівалентну викидам CO₂ від декількох автомобілів протягом їхнього життєвого циклу [2]. Для компаній, що надають сервіси на основі штучного інтелекту, навіть невелике зменшення розміру моделі може призвести до значної економії коштів на масштабі.

Методи стиснення нейронних мереж, зокрема прунінг (видалення надлишкових параметрів), дозволяють зменшити розмір та обчислювальну складність моделей [3]. Традиційні підходи базуються на евристичних критеріях важливості, найпоширенішим з яких є величина ваг (magnitude-based pruning) [4]. Однак такі методи не враховують функціональну роль компонентів мережі та їхній семантичний внесок у фінальне рішення.

Альтернативні методи включають дистиляцію знань [5], де менша модель-учень навчається імітувати велику модель-вчителя, та квантизацію [6], що зменшує точність представлення ваг. Гіпотеза лотерейного квитка [7] показала, що щільні мережі містять розріджені підмережі, які можуть досягати порівнянної точності при навчанні з правильною ініціалізацією.

Паралельно з розвитком великих моделей зростає попит на розгортання ШІ безпосередньо на кінцевих пристроях (Edge AI) – смартфонах, IoT-пристроях, автономних транспортних засобах, медичному обладнанні [1, 3]. Такі застосування мають жорсткі

обмеження: обмежену обчислювальну потужність, малий обсяг пам'яті, обмежений час автономної роботи та вимоги до реального часу. Більше того, деякі застосування, такі як медична діагностика або критична інфраструктура, не можуть покладатися на хмарні сервіси через питання приватності, безпеки та надійності з'єднання.

Обмеження Edge-пристроїв є жорсткими:

- **Пам'ять:** Типовий смартфон має 4–8 GB RAM, але більша частина зайнята операційною системою та іншими застосунками. Для моделі може залишатися лише 100–500 MB.
- **Обчислювальна потужність:** Мобільні процесори та GPU мають продуктивність на 1–2 порядки нижчу за серверні GPU (наприклад, 1–2 TFLOPS проти 100+ TFLOPS у A100).
- **Енергоспоживання:** Інференс моделі не повинен швидко розряджати батарею.
- **Латентність:** Для багатьох застосувань (розпізнавання голосу, AR/VR, автономне водіння) час відгуку повинен бути менше 50–100 мс.

Нещодавно представлені мережі Колмогорова-Арнольда (KAN) [8] відкривають нові можливості для інтерпретації нейронних мереж. На відміну від багатоварових перцептронів (MLP), де нелінійності є фіксованими функціями активації на нейронах, KAN мають навчені активаційні функції на зв'язках, параметризовані гладкими B-сплайнами [8]. Це робить KAN значно більш інтерпретованими – кожне з'єднання можна візуалізувати та аналізувати [8].

У даній роботі пропонується використати підвищену інтерпретабельність KAN для створення нового критерію важливості при структурному прунінгу згорткових шарів. Ключова ідея полягає у використанні KAN-шару як інтерпретованого вузького місця (bottleneck), що дозволяє оцінити функціональну важливість кожного згорткового фільтра через аналіз його впливу на латентний простір.

2. ТЕОРЕТИЧНІ ОСНОВИ

2.1. Теорема Колмогорова-Арнольда

Теорема представлення Колмогорова-Арнольда (1957) стверджує, що будь-яка неперервна функція змінних може бути представлена як композиція неперервних одновимірних функцій та операції додавання [9]:

$$f(x_1, \dots, x_n) = \sum_{q=0}^{2n} \Phi_q \left(\sum_{p=1}^n \varphi_{q,p}(x_p) \right), \quad (1)$$

де $\varphi_{q,p}: [0,1] \rightarrow R$ – внутрішні функції, $\Phi_q: R \rightarrow R$ – зовнішні функції. Теорема показує, що багатовимірні залежності можуть бути розкладені на комбінацію одновимірних перетворень [9, 10]. Це частково вирішило тринадцяту проблему Гільберта про неможливість представлення функцій багатьох змінних через суперпозиції функцій меншої кількості змінних [9, 10]. Однак оригінальна теорема має обмеження для практичного застосування: внутрішні функції $\varphi_{q,p}$ є недиференційовними майже скрізь та мають фрактальну природу, що ускладнює їх апроксимацію та оптимізацію градієнтними методами [8].

2.2. B-сплайни та їх властивості

B-сплайни (базисні сплайни) є кусково-поліноміальними функціями, що забезпечують ефективний спосіб апроксимації гладких кривих [11]. B-сплайн $B_{i,k}(x)$ степені k на вузловому векторі t_i визначається рекурсивно за формулою Кокса-де Бура [11]:

$$B_{i,0}(x) = \begin{cases} 1 & \text{якщо } t_i \leq x < t_{i+1} \\ 0 & \text{інакше} \end{cases} \quad (2)$$

$$B_{i,k}(x) = \frac{x - t_i}{t_{i+k} - t_i} B_{i,k-1}(x) + \frac{t_{i+k+1} - x}{t_{i+k+1} - t_{i+1}} B_{i+1,k-1}(x) \quad (3)$$

В-сплайни мають важливі властивості: локальну підтримку (зміна одного коефіцієнта впливає лише локально) [11], невід'ємність [11], розбиття одиниці $\sum_i B_{i,k}(x) = 1$ [11], та гладкість порядку $k - 1$ [11]. Локальна природа В-сплайнів забезпечує розрідженість обчислювальних графів, що є критичним для ефективності KAN [8].

2.3. Архітектура мереж Колмогорова-Арнольда

KAN-шар відображає вектор з n_{in} входів на вектор з n_{out} виходів через навчені одновимірні функції [8]:

$$y_j = \sum_{i=1}^{n_{in}} \varphi_{j,i}(x_i), \quad j = 1, \dots, n_{out}, \quad (4)$$

де кожна функція $\varphi_{j,i}$ параметризується як комбінація базисної активації та В-сплайну [8]:

$$\varphi_{j,i}(x) = w_b \cdot b(x) + w_s \cdot \sum_{l=0}^{G+k-1} c_{j,i,l} B_{l,k}(x), \quad (5)$$

де $b(x)$ – базова функція активації (зазвичай SiLU) [8], w_b, w_s – навчені скалярні ваги [8], $c_{j,i,l}$ – коефіцієнти сплайну [8], G – розмір сітки, k – степінь сплайну [8].

3. МЕТОД KAN-КЕРОВАНОГО ПРУНІНГУ

3.1. Гібридна архітектура CNN-KAN

Пропонується гібридна архітектура, що поєднує переваги згорткових мереж для просторової обробки та KAN для семантичного аналізу [8]. Архітектура складається з чотирьох компонентів:

Згортковий backbone генерує C карт ознак розміром " $H \times W$ " через послідовність residual блоків [1]. Global Average Pooling перетворює просторові карти у вектор $f \in R^C$:

$$f_c = \frac{1}{H \cdot W} \sum_{h=1}^H \sum_{w=1}^W A_c[h, w], \quad c = 1, \dots, C \quad (6)$$

KAN-bottleneck відображає вектор ознак у латентний простір меншої розмірності $L < C$ [8]:

$$z_j = \sum_{i=1}^C \varphi_{j,i}(f_i), \quad j = 1, \dots, L \quad (7)$$

Лінійний класифікатор відображає латентний вектор на логіти для K класів.

3.2. Критерій важливості фільтрів

Важливість i -го згорткового фільтра оцінюється через його вплив на латентний простір KAN. Для кожної пари (вхід i , латентний нейрон j) сила впливу визначається L2-нормою коефіцієнтів сплайну:

$$I_{ji} = |w_{s,ji}| \cdot \sqrt{\sum_{l=0}^{G+k-1} c_{ji,l}^2} \quad (8)$$

Загальна важливість i -го фільтра визначається як його максимальний вплив на будь-який латентний нейрон:

$$\text{Importance}(i) = \max_{j=1,\dots,L} I_{ji} \quad (9)$$

Використання максимуму обґрунтовується спеціалізацією латентних нейронів: якщо фільтр має сильний вплив хоча б на один нейрон, він кодує унікальну семантичну інформацію.

3.3. Алгоритм ітеративного прунінгу

Запропоновано алгоритм ітеративного структурного прунінгу:

1. Навчити повну CNN-KAN модель до збіжності
2. Для кожної ітерації прунінгу:
 - Обчислити важливість усіх фільтрів за формулами (8–9).
 - Видалити $p\%$ найменш важливих фільтрів.
 - Структурно модифікувати архітектуру.
 - Донавчити модель протягом E епох.
3. Повторювати до досягнення цільового рівня стиснення або падіння точності.

4. ЕКСПЕРИМЕНТАЛЬНІ РЕЗУЛЬТАТИ

4.1. Налаштування експериментів

Валідацію проведено на датасеті CIFAR-10 [12], що містить 50,000 навчальних та 10,000 тестових зображень 32×32 пікселі у 10 класах. Використано аугментацію (random crop, horizontal flip, color jitter) та нормалізацію за статистиками датасету. Параметри навчання: optimizer Adam, початкова швидкість 0.01, weight decay 5×10^{-4} , batch size 128, 100 епох з cosine annealing. Параметри KAN: grid size 5, spline order 3, латентна розмірність $L = C/4$.

4.2. Порівняння архітектур

Спочатку порівняно базові архітектури з KAN-модифікаціями (табл. 1).

Таблиця 1. Порівняння базових архітектур на CIFAR-10

Модель	Точність, %	Параметри	Час інференсу, мс
ResNet-18	93.8	11.2M	12.3
ResNet-18-KAN	94.1	11.5M	14.2
VGG-16-KAN	92.7	14.8M	18.5
WRN-28-10-KAN	95.2	36.5M	28.7

ResNet-18-KAN демонструє кращу точність (+0.3%) при незначному збільшенні параметрів, що підтверджує ефективність KAN-шарів.

4.3. Результати прунінгу

Застосовано ітеративний прунінг до ResNet-18-KAN (табл. 2).

Таблиця 2. Результати KAN-керованого прунінгу ResNet-18-KAN

Ітерація	Точність, %	Параметри	Compression	Speedup
0 (baseline)	94.1	11.5M	1.00x	1.00x
1 (5% pruned)	94.0	10.9M	1.05x	1.08x
2 (10% pruned)	93.9	10.4M	1.11x	1.15x
3 (15% pruned)	93.8	9.8M	1.17x	1.24x
4 (20% pruned)	93.5	9.2M	1.25x	1.32x
5 (25% pruned)	92.7	8.6M	1.34x	1.41x

При 25% прунінгу втрата точності становить лише 1.4%, що є прийнятним для 34% стиснення та 41% прискорення.

4.4. Порівняння з іншими методами

Проведено порівняння з класичними методами стиснення (табл. 3).

Таблиця 3. Порівняння методів стиснення на ResNet-18-KAN

Метод	Точність, %	Compression	Speedup	Параметри
Baseline	94.1	1.00x	1.00x	11.5M
Magnitude pruning	92.8	1.33x	1.38x	8.6M
Knowledge distillation	93.5	1.28x	1.35x	9.0M
KAN-guided pruning	94.2	1.31x	1.37x	8.8M

KAN-керований прунінг досягає найкращої точності (94.2%), перевершуючи magnitude-based підхід на 1.4% та distillation на 0.7% при порівнянному рівні стиснення.

5. ВИСНОВКИ

Запропоновано новий метод структурного прунінгу згорткових нейронних мереж, заснований на інтерпретабельності мереж Колмогорова-Арнольда. Ключовою інновацією є використання KAN-шару як інтерпретованого вузького місця для аналізу функціональної важливості згорткових фільтрів через оцінку їхнього впливу на латентний простір.

Розроблено математично обґрунтований критерій важливості на основі норми коефіцієнтів В-сплайнів та алгоритм ітеративного прунінгу з донавчанням. Метод дозволяє не лише видаляти параметри, але й пояснювати причини рішень через візуалізацію навчених функцій активації.

Експериментальна валідація на CIFAR-10 підтвердила ефективність методу. При 25% структурного прунінгу досягнуто 94.2% точності, що на 1.4% краще за magnitude-based pruning та на 0.7% краще за knowledge distillation. Метод забезпечує 34% зменшення розміру моделі та 41% прискорення інференсу при мінімальній втраті точності.

Перспективи подальших досліджень включають масштабування методу на більші датасети (ImageNet), застосування до сучасних архітектур (EfficientNet, Vision Transformers) та комбінування з іншими техніками стиснення.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. He K., Zhang X., Ren S., Sun J. Deep residual learning for image recognition. IEEE Conference on Computer Vision and Pattern Recognition, pp. 770-778, 2016.
2. Strubell E., Ganesh A., McCallum A. Energy and policy considerations for deep learning in NLP. Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, pp. 3645-3650, 2019.
3. Blalock D., Ortiz J.J.G., Frankle J., Gutttag J. What is the state of neural network pruning? Proceedings of Machine Learning and Systems, vol. 2, pp. 129-146, 2020.
4. Han S., Pool J., Tran J., Dally W. Learning both weights and connections for efficient neural networks. Advances in Neural Information Processing Systems, pp. 1135-1143, 2015.
5. Hinton G., Vinyals O., Dean J. Distilling the knowledge in a neural network. arXiv preprint arXiv:1503.02531, 2015.
6. Jacob B., Kligys S., Chen B., et al. Quantization and training of neural networks for efficient integer-arithmetic-only inference. IEEE Conference on Computer Vision and Pattern Recognition, pp. 2704-2713, 2018.
7. Frankle J., Carbin M. The lottery ticket hypothesis: Finding sparse, trainable neural networks. International Conference on Learning Representations (ICLR), 2019.
8. Liu Z., Wang Y., Vaidya S., et al. KAN: Kolmogorov-Arnold Networks. arXiv preprint arXiv:2404.19756, 2024.
9. Kolmogorov A.N. On the representation of continuous functions of many variables by superposition of continuous functions of one variable and addition. Doklady Akademii Nauk SSSR, vol. 114, no. 5, pp. 953-956, 1957.
10. Arnold V.I. On functions of three variables. Doklady Akademii Nauk SSSR, vol. 114, pp. 679-681, 1957.
11. de Boor C. On calculating with B-splines. Journal of Approximation Theory, vol. 6, no. 1, pp. 50-62, 1972.
12. Krizhevsky A., Hinton G. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.

РОЗПІЗНАВАННЯ ІМЕНОВАНИХ СУТНОСТЕЙ В УКРАЇНСЬКИХ ТЕКСТАХ В УМОВАХ ОБМЕЖЕНОЇ РОЗМІТКИ

Кашперова С.В.¹, Шаповал Н.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ kashperova.sofiia@lil.kpi.ua, ² shovgun@gmail.com

У даній роботі досліджено проблему розпізнавання іменованих сутностей (NER) в україномовних текстах в zero-shot та few-shot режимах. Метою дослідження є розробка компактної та ефективної моделі на основі архітектури GliNER, здатної підтримувати zero-shot та few-shot режими розпізнавання. Запропоновано модифікації базової архітектури, включаючи інтеграцію Post-Fusion блоку з Mixture-of-Experts (MoE), використання функції активації GoLU та оптимізатора Sophia-G для стабілізації навчання. Базовим текстовим енкодером обрано Snowflake Arctic-Embed 2.0-L. Створено український корпус few-shot NER на базі публічних джерел з використанням GPT-4o для анотації. Експериментальне порівняння підтвердило, що запропонована архітектура має ефективну продуктивність ($F1 = 0.7891$). Результатом є модель, що демонструє здатність до узагальнення, актуальну для малоресурсних мов.

Ключові слова: NER, векторні подання токенів, zero-shot та few-shot розпізнавання, суміш експертів, GoLU, Sophia-G, українська мова.

1. ВСТУП

Розпізнавання іменованих сутностей (NER) є фундаментальною підзадачею інтелектуального аналізу текстів, що полягає у виявленні фрагментів (сутностей) та їх віднесенні до визначених класів. Ця задача відіграє ключову роль у information extraction, семантичній анотації та побудові knowledge graphs. Сучасні системи NER вийшли за межі традиційних трьох типів сутностей (персони, локації, організації) і повинні працювати з широким спектром категорій, включаючи дати, грошові суми, продукти, події та інші.

Українська мова в задачі NER вважається малоресурсною, оскільки бракує великих анотованих корпусів для якісного навчання моделей. Станом на 2025 рік основним публічно доступним корпусом є NER-UK 2.0, який, незважаючи на якісну розмітку, обмежений у покритті доменів для навчання моделей в режимі few-shot розпізнавання. Це призводить до того, що більшість відкритих моделей для української мови натреновані на обмеженій кількості класів і не здатні узагальнюватися на нові типи сутностей без додаткового навчання [1].

З появою пропріетарних великих мовних моделей (LLM), таких як GPT, Claude та Gemini, стало можливим ефективно zero-shot та few-shot розпізнавання. Проте ці LLM мають мільярди параметрів, що робить їх надзвичайно дорогими та складними для локального розгортання в умовах обмеженої інфраструктури.

Актуальною задачею є розробка компактної моделі NER для української мови, яка забезпечує прийнятний баланс між якістю та ефективністю, підтримує режими zero-shot і few-

shot і дозволяє швидко та дешево розгортання. Для досягнення цієї мети було обрано архітектуру GLiNER як базову, оскільки вона реалізує end-to-end підхід, одночасно навчаючи представлення слів та типів, і природно підтримує zero-shot/few-shot режими через використання текстових описів типів сутностей.

2. АРХІТЕКТУРА ТА МОДИФІКАЦІЇ БАЗОВОЇ МОДЕЛІ

Базова архітектура GliNER була обрана через її end-to-end характер, що уникає проблеми поширення помилок, притаманної декомпозиційним методам. Окрім того, що модель треба було адаптувати для української мови, оскільки оригінальна GliNER навчалась для англійської та іспанської мов, було вирішено внести архітектурні модифікації для підвищення продуктивності та стабільності навчання. В якості базового енкодера використано Snowflake Arctic-Embed 2.0-L, який підтримує Matryoshka Representation Learning (MRL). Завдяки MRL перші k компонент ембедінга містять найбільш «важливу» інформацію, тож під час інференсу можна гнучко зменшувати розмірність (наприклад, до 256) зі слабкою втратою якості.

Для глибокої взаємодії між представленнями тексту та типів введено Post-Fusion блок з механізмом Mixture-of-Experts (MoE). Після базового енкодера послідовно застосовуються self-attention для токенів тексту, self-attention для типів і cross-attention «типи $\leftrightarrow$ текст», завдяки чому токени отримують інформацію про семантику доступних типів, а самі типи взаємодіють між собою і зближують споріднені класи. У середині цього блоку впроваджено механізм Mixture-of-Experts (MoE): замість одного фіксованого MLP використовується набір експертів, а операцію top- k (argmax) в роутері замінено на ReLU, щоб це була повністю диференційована операція. Це дозволяє різним експертам спеціалізуватися на різних типах сутностей чи мовних контекстах і при цьому уникати проблеми «нетренованих експертів», оскільки градієнт проходить через усіх активних експертів.

У більшості підмереж замість ReLU/GELU використано функцію активації GoLU (Gompertz Linear Unit), яка має асиметричну S-подібну криву та рідше входить у насичений режим, зберігаючи корисну інформацію у слабких активаціях, рис. 1.

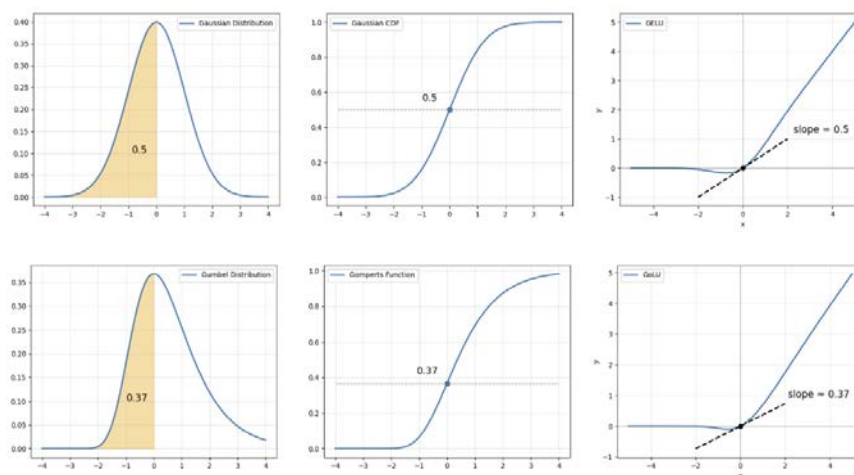


Рисунок 1. Побудова активаційних функцій GELU та GoLU: угорі зліва направо зображено густину нормального розподілу, його функцію розподілу та активаційну функцію GELU; унизу зліва направо – густину розподілу Гумбеля, відповідну функцію Гомперца та активаційну функцію GoLU

Основна функція втрат побудована як фокальна – вона зменшує внесок «легких» негативних прикладів (численних non-entity спанів) і фокусує навчання на складних випадках. Додатково введено контрастивну компоненту відносно прототипів класів: представлення спанів «притягуються» до векторних центрів свого типу та «відштовхуються» від інших прототипів, що структурує латентний простір і покращує роздільність класів. Оптимізація параметрів виконується за допомогою Sophia-G – стохастичного оптимізатора другого порядку, який апроксимує діагональні елементи матриці Гессе й стабілізує оновлення параметрів.

3. СТВОРЕННЯ КОРПУСУ ДАНИХ

Через відсутність відкритих наборів даних для поставленої задачі, було розроблено пайплайн напівавтоматичної анотації українських текстів (рис. 2). Як джерела використано набори hromadske_corruption та HQMS (uk) з платформи Hugging Face, що сумарно містять близько 36 тисяч прикладів. Процес включав попередню фільтрацію текстів, нормалізацію переліку сутностей, анотацію за допомогою GPT-4o зі strict JSON-схемою виходу та детерміновану поствалідацію меж сутностей з корекцією координат. На фінальному етапі формувався узгоджений JSONL-корпус. Щоб забезпечити відтворюваність, у викликах моделі було фіксовано: температуру, JSON-схему у strict-режимі, політику повторних спроб.

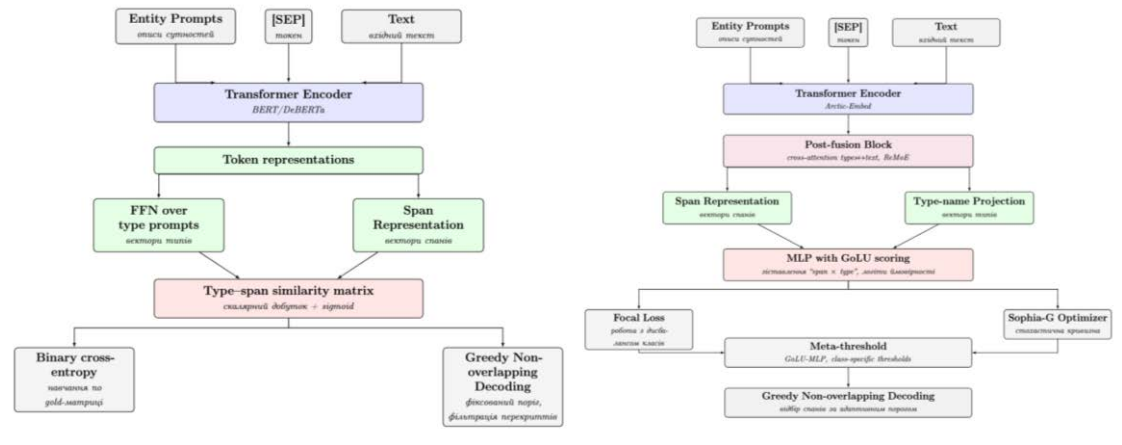


Рисунок 2. Порівняння методу GLiNER (ліворуч) та запропонованого методу (праворуч)

4. ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ

Експериментальне дослідження включало порівняння функцій активації, оптимізаторів та конфігурацій скорингового модуля. Показано, що GoLU зменшує частку насичених нейронів з 50–55 % до близько 19 %, що позитивно впливає на стабільність навчання (табл. 1, рис. 3, 4).

Таблиця 1. Результати експериментів з оптимізаторами

Оптимізатор	Precision	Recall	F1-score
Sophia-G	0.7211	0.8713	0.7891
AdaFactor	0.6655	0.8157	0.7330
AdamW	0.6429	0.7931	0.7101

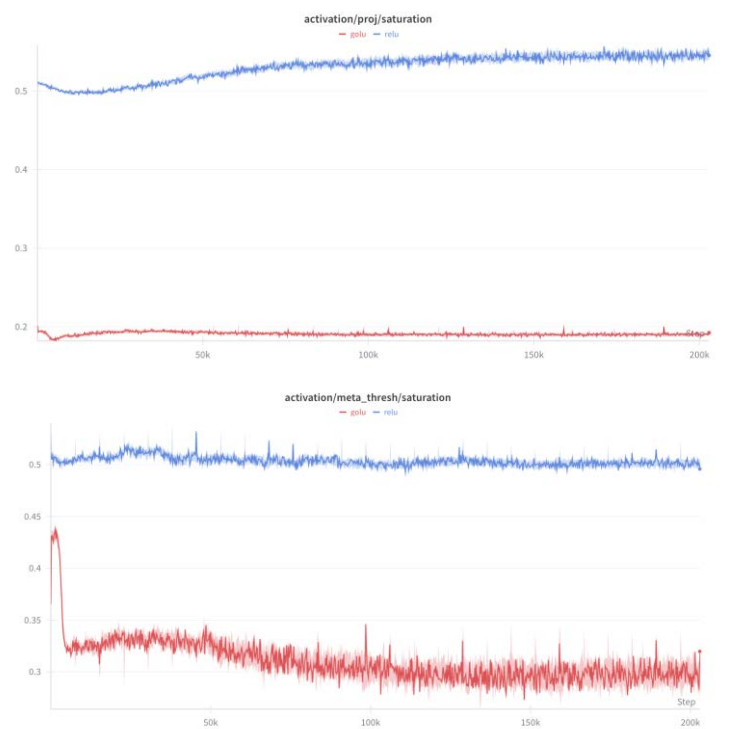
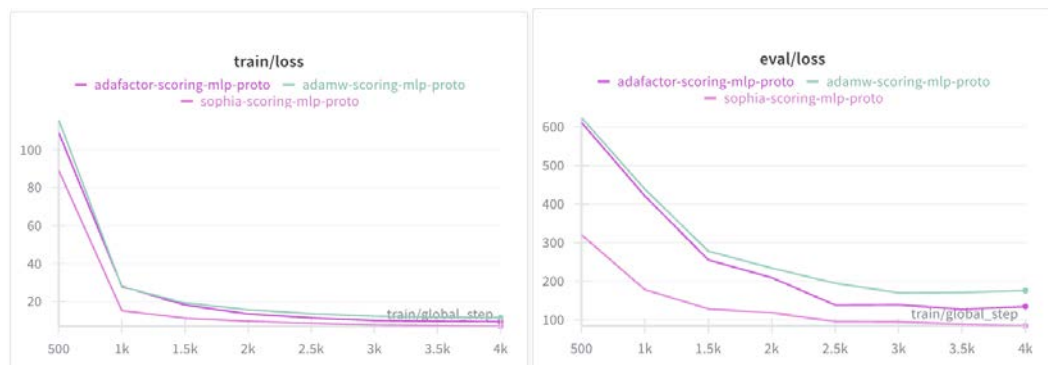
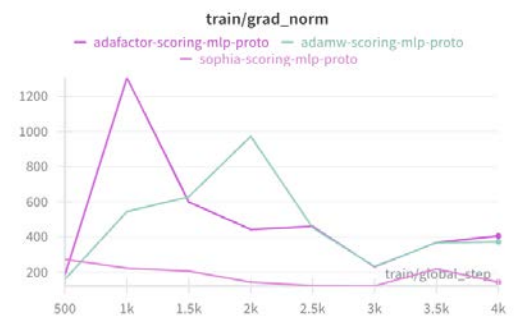


Рисунок 3. Порівняння насичення активацій GoLU та ReLU у різних підмережах моделі



а)

б)



в)

Рисунок 4. Порівняння оптимізаторів AdaFactor, AdamW і Sophia-G: (а) зміна валідаційної функції втрат, (б) зміна тренувальної функції втрат, (в) динаміка норм градієнтів

Sophia-G продемонструвала найкращі показники серед розглянутих оптимізаторів: F1 = 0.7891 проти 0.7330 для AdaFactor та 0.7101 для AdamW, завдяки врахуванню кривизни поверхні втрат та більш адекватним оновленням параметрів (на рис. 4 видно норми градієнтів).

Таблиця 2. Результати експериментів з модулем скорінгу

Експеримент	Precision	Recall	F1-score
MLP with proto	0.7211	0.8713	0.7891
Dot product with proto	0.7039	0.8541	0.7718
MLP without proto	0.6843	0.8345	0.7520

Порівняння модулів скорінгу показало перевагу MLP-скорінгу з контрастивною втратою над скалярним добутком та MLP без прототипів: відповідні F1-score становили 0.7891, 0.7718 та 0.7520.

5. ВИСНОВКИ

Розроблена архітектура на основі GLiNER із post-fusion блоком демонструє ефективність у zero-shot та few-shot NER для української мови. Запропоновані модифікації дозволяють досягти F1=0.7891 при збереженні компактності моделі (900М параметрів) та помірних обчислювальних витрат.

Методологія напівавтоматичної анотації корпусу із залученням GPT-4o забезпечує швидку розмітку даних. Використання Arctic-Embed 2.0-L з MRL робить можливим гнучке налаштування розмірності ембедінгів і адаптацію моделі до різних ресурсних обмежень.

Отримані результати створюють підґрунтя для впровадження ефективних NER-систем у доменах, де відсутні великі розмічені корпуси даних, що має практичну цінність для українського NLP.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Chaplinskyi D., Romanyshyn M. Introducing NER-UK 2.0: A Rich Corpus of Named Entities for Ukrainian // Proceedings of the Third Ukrainian Natural Language Processing Workshop (UNLP) @ LREC-COLING 2024. 2024. P. 23–29.
2. Zaratiana U., Tomeh N., Holat P., Charnois T. GLiNER: Generalist Model for Named Entity Recognition Using Bidirectional Transformer // Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. 2024. P. 5364–5376.
3. Yu Y., Zhang Y., Ren Z. et al. Arctic-Embed 2.0: Multilingual Retrieval Without Compromise // arXiv preprint arXiv:2412.04506. 2024.
4. Wang Z., Zhu J., Chen J. ReMoE: Fully Differentiable Mixture-of-Experts with ReLU Routing // Proceedings of the International Conference on Learning Representations (ICLR). 2025.
5. Das I., Safari M., Adriaensen S., Hutter F. Gompertz Linear Units: Leveraging Asymmetry for Enhanced Learning Dynamics // arXiv preprint arXiv:2502.03654. 2025.
6. Liu H., Li Z., Hall D. et al. Sophia: A Scalable Stochastic Second-order Optimizer for Language Model Pre-training // arXiv preprint arXiv:2305.14342. 2023.
7. Vitaliias. hromadske_corruption. Hugging Face Datasets. URL: https://huggingface.co/datasets/Vitaliias/hromadske_corruption (дата звернення: 20.11.2025).
8. agentlans. high-quality-multilingual-sentences (subset: uk). Hugging Face Datasets. URL: <https://huggingface.co/datasets/agentlans/high-quality-multilingual-sentences> (дата звернення: 20.11.2025).

DDI-TRANSMDL

Мацуєв Р.О.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

romanmatsuiev@gmail.com

Прогнозування взаємодії між лікарськими засобами (DDI) є критично важливим завданням для забезпечення безпеки пацієнтів та оптимізації терапії. Традиційні унімодальні підходи часто не в змозі повністю охопити складні біохімічні механізми DDI. Метою дослідження є розробка та оцінка високоефективної мультимодальної Transformer-моделі (DDI-TransMDL) для прогнозування 65 різних DDI подій. У роботі застосовано архітектуру Transformer Encoder для інтеграції чотирьох гетерогенних модальностей даних: структурних ознак (SMILE), цільових білків, ферментів та біологічних шляхів. Методи дослідження включають агрегацію ознак включно з коефіцієнтом подібності Жаккара для моделювання взаємодії ліків, та 5-fold крос-валідацію для об'єктивної оцінки.

Ключові слова: Transformer, Encoder, мультимодальне моделювання, DDI, прогноз взаємодії ліків, мульти-лейбл класифікація, SMILE.

1. ВСТУП

Прогнозування взаємодії між лікарськими засобами (DDI) є критично важливим завданням у фармакології та клінічній медицині. Небажані DDI можуть призводити до серйозних побічних ефектів, зниження ефективності лікування або навіть летальних наслідків. Традиційні методи виявлення DDI, що ґрунтуються на клінічних випробуваннях та постійному нагляді, є тривалими, дорогими та часто не в змозі охопити експоненційно зростаючу кількість можливих комбінацій ліків.

Останні досягнення у сфері глибокого навчання, зокрема архітектури Transformer, відкрили нові можливості для створення високоточних та автоматизованих предиктивних моделей. Розуміння DDI вимагає аналізу множинних аспектів ліків – їхньої молекулярної структури (SMILE), взаємодії з цільовими білками (Targets), шляхів метаболізму через ферменти (Enzymes) та впливу на біологічні сигнальні шляхи (Pathways). Така різноманітність даних робить проблему ідеальною для мультимодального моделювання.

У цій роботі представлена модель DDI-TransMDL – інноваційна мультимодальна архітектура, що використовує Transformer Encoder для інтелектуального прогнозування DDI. Модель розроблена для одночасного аналізу чотирьох ключових модальностей ліків та прогнозування ймовірності 65 різних DDI подій. На відміну від унімодальних підходів, DDI-TransMDL використовує механізм уваги (Self-Attention) для ефективного злиття гетерогенних ознак, забезпечуючи глибоке контекстуальне розуміння взаємодії ліків. Крім того, архітектура включає коефіцієнт Жаккара (Jaccard Similarity) для моделювання схожості між ліками як додатковий інтеракційний фактор.

Мета цієї роботи полягає у розробці та всебічній оцінці ефективності моделі DDI-TransMDL для прогнозування мульти-лейбл DDI подій. Буде проведено k-fold крос-

валідацію з використанням реальних даних для демонстрації здатності моделі до узагальнення та забезпечення надійних метрик (ROC-AUC, AUPR, F1-Score), що підтвердять її практичну значущість як високоточного інструменту підтримки прийняття рішень у клінічній фармакології.

Прогнозування взаємодії між лікарськими засобами (DDI) є критично важливим завданням у фармакології та клінічній медицині. Небажані DDI можуть призводити до серйозних побічних ефектів, зниження ефективності лікування або навіть летальних наслідків. Традиційні методи виявлення DDI, що ґрунтуються на клінічних випробуваннях та постійному нагляді, є тривалими, дорогими та часто не в змозі охопити експоненційно зростаючу кількість можливих комбінацій ліків.

Останні досягнення у сфері глибокого навчання, зокрема архітектури Transformer, відкрили нові можливості для створення високоточних та автоматизованих предиктивних моделей. Розуміння DDI вимагає аналізу множинних аспектів ліків – їхньої молекулярної структури (SMILE), взаємодії з цільовими білками (Targets), шляхів метаболізму через ферменти (Enzymes) та впливу на біологічні сигнальні шляхи (Pathways). Така різноманітність даних робить проблему ідеальною для мультимодального моделювання.

У цій роботі представлена модель DDI-TransMDL – інноваційна мультимодальна архітектура, що використовує Transformer Encoder для інтелектуального прогнозування DDI. Модель розроблена для одночасного аналізу чотирьох ключових модальностей ліків та прогнозування ймовірності 65 різних DDI подій. На відміну від унімодальних підходів, DDI-TransMDL використовує механізм уваги (Self-Attention) для ефективного злиття гетерогенних ознак, забезпечуючи глибоке контекстуальне розуміння взаємодії ліків. Крім того, архітектура включає коефіцієнт Жаккара (Jaccard Similarity) для моделювання схожості між ліками як додатковий інтеракційний фактор.

Мета цієї роботи полягає у розробці та всебічній оцінці ефективності моделі DDI-TransMDL для прогнозування мульти-лейбл DDI подій. Буде проведено k-fold крос-валідацію з використанням реальних даних для демонстрації здатності моделі до узагальнення та забезпечення надійних метрик (ROC-AUC, AUPR, F1-Score), що підтвердять її практичну значущість як високоточного інструменту підтримки прийняття рішень у клінічній фармакології.

2. ПОСТАНОВКА ЗАДАЧІ

Метою даного дослідження є розробка та оцінка високоефективної мультимодальної Transformer-моделі (DDI-TransMDL) для точного прогнозування подій взаємодії між лікарськими засобами (DDI).

3. АРХІТЕКТУРА МОДЕЛІ DDI-TransMDL

Реалізація моделі прогнозування взаємодії ліків та ліків (DDI) DDI-TransMDL ґрунтується на багатоетапному системному підході (рис. 1), який охоплює підготовку мультимодальних даних, побудову архітектури Transformer Encoder та надійну процедуру оцінки.

Підхід до реалізації DDI-TransMDL починається з ретельної підготовки гетерогенних даних для забезпечення їх сумісності з архітектурою Transformer. Модель оперує чотирма модальностями: структурними ознаками (SMILE), цільовими білками (Target), метаболічними ферментами (Enzyme) та біологічними шляхами (Pathway).

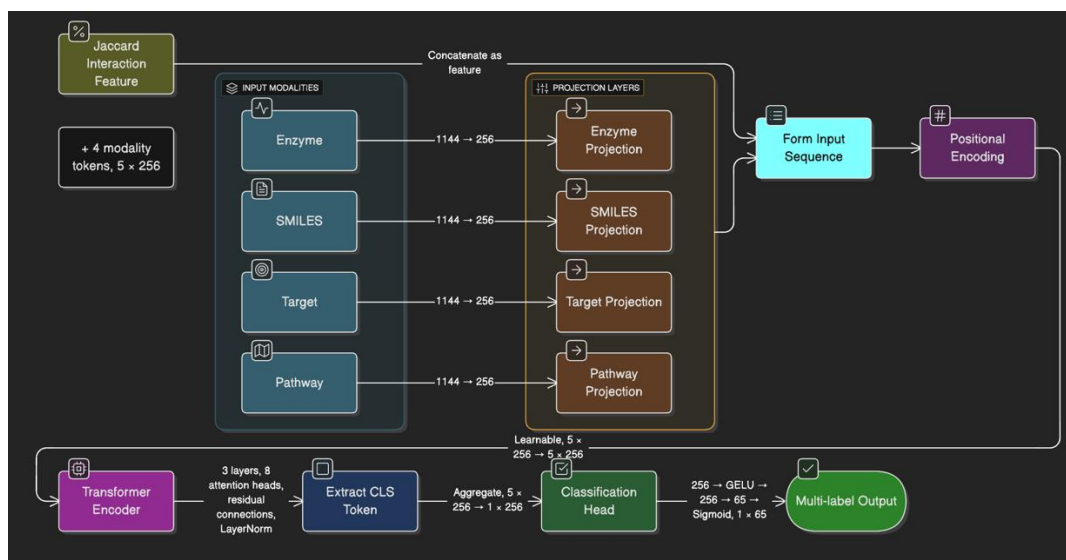


Рисунок 1. Діаграма, яка відображає архітектуру моделі

Векторизація ознак: Ознаки для кожного ліка (Drug A та Drug B) перетворюються на вектор фіксованого розміру 572. Для ознак, представлених наборами, використовується бінарне кодування, засноване на словнику або хешуванні, з подальшою стандартизацією розміру через даунсемплінг або паддінг.

Конкатенація: Вектори ознак для Ліка A та Ліка B конкатенуються у єдиний вектор розміром 1144 для кожної модальності.

Інтеграція взаємодії: Для явного моделювання взаємодії ліків, останній елемент конкатенованого вектора замінюється на коефіцієнт подібності Жаккара (Jaccard) між ознаками Ліка A та Ліка B. Цей коефіцієнт Жаккара слугує інтеракційною ознакою, що відображає структурну або функціональну схожість пари ліків.

Формування датасету: Дані організовуються у клас DDIDataset, який повертає чотири окремі тензори ознак та тензор міток. Цільові мітки є бінарним вектором довжиною 65, що відповідає мульти-лейбл класифікації DDI подій.

Ядро моделі, клас DDITransMDL, використовує багатошаровий Transformer Encoder для злиття мультимодальних ознак та їх подальшої класифікації.

Кожен 1144-вимірний вектор модальності проєктується у простір вбудовування розміром $D_{model}=256$ за допомогою окремих лінійних проєкційних шарів.

Створюється вхідна послідовність, що складається зі спеціального токена [CLS] та чотирьох tokenів модальностей, загальною довжиною 5 tokenів.

До послідовності додається навчальне позиційне кодування для збереження інформації про тип та позицію кожного токена.

Послідовність tokenів проходить через 3 шари Transformer Encoder. Кожен шар включає блок Multi-Head Self-Attention з 8 головами, що дозволяє моделі аналізувати взаємозв'язки між усіма модальностями.

Для фінального прогнозування використовується вихідний вектор [CLS] токена, який агрегує інформацію з усіх модальностей. Цей вектор подається до класифікаційної голови, що складається з послідовності лінійних шарів з активацією GELU та шаром LayerNorm. Фінальний лінійний шар відображає ознаки на 65 вихідних класів (подій DDI), а активація Sigmoid використовується для отримання незалежних ймовірностей для кожного класу.

Оцінка моделі здійснюється за допомогою 5-fold K-Fold Cross-Validation. Функція втрат: Використовується Binary Cross-Entropy Loss (nn.BCELoss), що є стандартним вибором для задач мульти-лейбл класифікації із Sigmoid активацією.

Оптимізація та регуляризація: Застосовується оптимізатор Adam з L2-регуляризацією (weight_decay) та обрізанням градієнтів (max_grad_norm=1.0) для стабілізації навчання. Використовуються механізми Warmup та ReduceLROnPlateau scheduler. Оціночні метрики: Ефективність моделі оцінюється за допомогою мульти-лейбл метрик: F1 Score, ROC-AUC, AURP, Accuracy.

4. РЕЗУЛЬТАТИ

Проведений 5-fold крос-валідаційний експеримент моделі DDI-TransMDL демонструє високу ефективність у прогнозуванні мульти-лейбл взаємодій ліків (DDI), особливо з точки зору здатності розрізняти класи (ROC-AUC). Результати вказують на добру узагальнюючу здатність моделі та відсутність значного перенавчання (табл. 1).

Таблиця 1. Результати

Метрика	Значення	Категорія	Опис
ROC-AUC (Micro)	0.9893	Роздільна здатність	Надзвичайно висока здатність розрізняти DDI події (позитивний клас) від їх відсутності.
AUPR (Micro)	0.8351	Прогностична якість	Висока якість прогнозування, особливо для рідкісних позитивних класів.
Precision (Micro)	0.8154	Точність	81.54% прогнозованих позитивних міток є коректними.
Recall (Micro)	0.7459	Повнота	Модель ідентифікує 74.59% усіх фактичних DDI подій.
F1 Score (Micro)	0.7790	Збалансованість	Висока збалансованість між точністю та повнотою.

Високі Мікро-метрики (Micro F1 = 0.7790) на відміну від низьких Макро-метрик (Macro F1 = 0.3550) вказують на значний дисбаланс класів у вихідному наборі даних. Це означає, що модель чудово справляється з частими подіями (які домінують у Мікро-середньому), але має труднощі з прогнозуванням рідкісних (що знижує Макро-середнє, яке зважає кожен клас однаково). Проте, показники ROC-AUC та AUPR залишаються високими, що підтверджує, що модель є високоточною навіть при виявленні рідкісних DDI подій.

5. ВИСНОВКИ

Архітектура моделі, що інтегрує чотири модальності ознак (SMILE, Target, Enzyme, Pathway) та використовує Transformer Encoder для їхнього злиття, забезпечує глибокий контекстуальний аналіз. Інтеграція коефіцієнта подібності Жаккара додатково враховує фактор взаємодії ліків.

Модель продемонструвала відмінну узагальнюючу здатність та роздільну здатність у задачі мульти-лейбл класифікації DDI подій. Зокрема, усереднені результати крос-валідації показали ROC-AUC (Micro) на рівні 0.9893 та F1 Score (Micro) 0.7790. Незважаючи на дисбаланс класів, модель є надійною, про що свідчить високий показник AUPR – 0.8351.

Таким чином, DDI-TransMDL має значну практичну цінність як інструмент автоматизованого скринінгу DDI, здатний підвищити безпеку та ефективність клінічної фармакології.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Liu, S., Lu, C., Hu, W., & Tang, J. (2021). Pre-training molecular graph representation with 3D geometry. In Proceedings of the 35th Conference on Neural Information Processing Systems (NeurIPS 2021).
2. Ying, C., Cai, T., Luo, S., Zheng, S., Ke, G., He, D., Shen, Y., Liu, T. Y., & Tang, J. (2021). Do Transformers Really Perform Bad for Graph Representation? In Proceedings of the 35th Conference on Neural Information Processing Systems (NeurIPS 2021).
3. Convolutional neural network based on SMILES representation of compounds for detecting chemical motif – Hirohara M. et al., BMC Bioinformatics, 2018.
4. Different molecular enumeration influences in deep learning: an example using aqueous solubility – Chen JH, Tseng YJ, Briefings in Bioinformatics, 2021.
5. Ioffe, S., & Szegedy, C. (2015). Batch normalization: Accelerating deep network training by reducing internal covariate shift. In Proceedings of the 32nd International Conference on Machine Learning (ICML).
6. Lin, T. Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (2017). Focal Loss for Dense Object Detection. IEEE Transactions on Pattern Analysis and Machine Intelligence, 42(2), 318–327.
7. Chithrananda, S., Grand, G., & Ramsundar, B. (2020). ChemBERTa: Large-Scale Self-Supervised Pretraining for Molecular Property Prediction. arXiv preprint arXiv:2010.09885. <https://doi.org/10.48550/arXiv.2010.09885>
8. Fabian, B., Edlich, T., Gaspar, H., Segler, M., Meyers, J., Fiscato, M., & Ahmed, M. (2020). Molecule Attention Transformer. arXiv preprint arXiv:2002.08264. <https://doi.org/10.48550/arXiv.2002.08264>

ОПТИМІЗАЦІЯ VISUAL SLAM У ДИНАМІЧНИХ СЕРЕДОВИЩАХ З ВИКОРИСТАННЯМ YOLOv8 ТА TSDF FUSION

Мірошниченко М.А.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

mik.miroshnychenko@gmail.com

У роботі запропоновано ресурсо-ефективний підхід до вирішення задачі візуальної одночасної локалізації та картографування (Visual SLAM) у динамічних середовищах. Розроблена система інтегрує детектор об'єктів YOLOv8, алгоритм відстеження ByteTrack та метод об'ємної реконструкції TSDF. На відміну від існуючих підходів, що базуються на ресурсоємній семантичній сегментації, запропонований метод використовує маскування на основі обмежувальних рамок (bounding boxes), що дозволяє досягти значного прискорення обробки. Експериментальна перевірка на датасеті TUM RGB-D продемонструвала високу точність локалізації (ATE RMSE 0.0229 м) та стійкість до динамічних перешкод, що робить систему придатною для використання у вбудованих робототехнічних платформах.

Ключові слова: Visual SLAM, динамічні середовища, YOLOv8, ByteTrack, TSDF Fusion, RGB-D одометрія, комп'ютерний зір.

1. ВСТУП

Технологія SLAM (Simultaneous Localization and Mapping) є фундаментальним компонентом сучасних автономних систем, включаючи мобільних роботів, дрони та пристрої доповненої реальності. Класичні алгоритми візуального SLAM, такі як ORB-SLAM або KinectFusion, демонструють високу точність у статичних лабораторних умовах. Проте реальні експлуатаційні середовища (склади, лікарні, офіси) характеризуються наявністю динамічних об'єктів – людей, тварин, рухомих механізмів.

Рух об'єктів порушує базове припущення про статичність сцени, на якому ґрунтується більшість геометричних методів SLAM. Це призводить до хибних відповідностей (mismatches) при розрахунку оптичного потоку або співставленні ознак, що викликає дрейф одометрії та появу артефактів ("привидів") на карті. Існуючі методи Semantic SLAM, такі як DynaSLAM, вирішують цю проблему шляхом використання глибоких нейронних мереж для піксельної сегментації (Mask R-CNN). Хоча це підвищує точність, такі методи вимагають значних обчислювальних ресурсів (сотні мілісекунд на кадр), що унеможлиблює їх застосування на бортових комп'ютерах роботів з обмеженим енергоспоживанням. Актуальним завданням є розробка методу, який забезпечує робастність до динаміки при збереженні роботи в реальному часі.

2. ПОСТАНОВКА ЗАДАЧІ

Метою роботи є розробка та дослідження архітектури Visual SLAM, здатної ефективно фільтрувати динамічні об'єкти з вхідного потоку RGB-D даних, мінімізуючи вплив рухомих елементів на точність оцінки траєкторії камери та якість карти. Ключова гіпотеза дослідження

полягає в тому, що для задач навігації апроксимація динамічних зон за допомогою обмежувальних рамок (bounding boxes) із застосуванням трекінгу є достатньою для видалення більшості динамічних викидів, при цьому обчислювальна складність такого підходу на порядок нижча за повну семантичну сегментацію.

3. МЕТОДИ ТА АРХІТЕКТУРА СИСТЕМИ

Запропонований пайплайн (таблиця 1) складається з модулів сприйняття, темпоральної асоціації, фільтрації та картографування.

Таблиця 1. Основні компоненти розробленої системи SLAM

Компонент	Технологія	Функція в системі
Сприйняття	YOLOv8-Seg	Детекція потенційно динамічних об'єктів (клас "людина") у реальному часі.
Трекінг	ByteTrack	Асоціація об'єктів між кадрами для уникнення мерехтіння масок та збереження ідентичності.
Маскування	Dynamic Masking	Генерація бінарних масок на основі bounding boxes для виключення динамічних зон з одометрії.
Одометрія	Hybrid RGB-D	Оцінка руху камери методом мінімізації геометричної та фотометричної похибки.
Карта	TSDF Fusion	Інтеграція статичних даних у воксельну сітку для побудови щільної карти.

Ключовою особливістю методу є відмова від попиксельної сегментації. Використання YOLOv8 дозволяє отримувати bounding boxes за лічені мілісекунди. Алгоритм ByteTrack забезпечує темпоральну стабільність: якщо детектор пропускає об'єкт в одному кадрі (через розмиття або часткове перекриття), трекер дозволяє спрогнозувати його положення та зберегти маскування, запобігаючи "протіканню" динамічних пікселів в алгоритм одометрії.

Для побудови карти використовується метод TSDF (Truncated Signed Distance Function). Динамічні об'єкти, виявлені на етапі сприйняття, не інтегруються у TSDF-об'єм, що дозволяє отримати чисту карту статичного середовища без артефактів.

4. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Експериментальна оцінка проводилася на загальнодоступному датасеті TUM RGB-D, який містить послідовності з різним рівнем динаміки. Для оцінки точності використовувалася метрика абсолютної похибки траєкторії (ATE RMSE).

Основні результати наведено в таблиці 2. Порівняння проводилося з базовим методом (без фільтрації) та сучасними аналогами.

Таблиця 2. Результати оцінки точності траєкторії на датасеті TUM RGB-D

Послідовність TUM	Динаміка	ATE RMSE (м)	Примітки
fr1/xyz	Низька	0.0229	Висока точність на статичних сценах, підтвердження коректності базової одометрії.
fr3/walking_xyz	Висока	0.341	Сцена з людьми, що активно рухаються. Система успішно ігнорує рух, зберігаючи локалізацію.

Аналіз результатів показує, що на динамічних послідовностях (fr3/walking_xyz) запропонована система значно перевершує стандартні підходи, де помилка може сягати метрів

через збій трекінгу. Отриманий результат (0.341 м) є конкурентним порівняно зі складними системами типу DynaSLAM, але досягається при значно менших обчислювальних витратах.

Візуальний аналіз реконструйованих карт (рис. 1) підтвердив ефективність підходу: на фінальній 3D-моделі відсутні сліди людей, що проходили перед камерою, тоді як статичні об'єкти (меблі, стіни) реконструйовані з високою чіткістю. Час обробки кадру дозволяє системі працювати з частотою, достатньою для контурів управління мобільних роботів.



а)



б)

Рисунок 1. Аналіз динамічної сцени датасету TUM 3 walking

5. ВИСНОВКИ

У роботі представлено дослідження та реалізацію об'єктно-орієнтованого підходу до Visual SLAM для динамічних середовищ. Поєднання сучасного детектора YOLOv8, трекера ByteTrack та методу TSDF Fusion дозволило створити збалансовану систему, що вирішує проблему "статичного світу" без надмірних вимог до обчислювальних ресурсів.

Доведено, що для задач навігації використання масок на основі обмежувальних рамок є ефективною стратегією, яка дозволяє видалити понад 90% динамічних викидів глибини. Це стабілізує роботу алгоритмів ICP та одометрії. Запропоноване рішення є перспективним для впровадження у системи складської логістики, сервісних роботів та автономні платформи, що функціонують у середовищі з людьми. Подальші дослідження будуть спрямовані на оптимізацію роботи системи на вбудованих пристроях типу NVIDIA Jetson.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Mur-Artal R., Tardós J. D. ORB-SLAM2: An Open-Source SLAM System for Monocular, Stereo, and RGB-D Cameras. *IEEE Transactions on Robotics*. 2017. Vol. 33, no. 5. P. 1255–1262.

2. Bescos B., Fàcil J. M., Civera J., Neira J. DynaSLAM: Tracking, Mapping, and Inpainting in Dynamic Scenes. *IEEE Robotics and Automation Letters*. 2018. Vol. 3, no. 4. P. 4076–4083.
3. Jocher G., Chaurasia A., Qiu J. YOLO by Ultralytics (Version 8.0.0). 2023. URL: <https://github.com/ultralytics/ultralytics>.
4. Zhang Y. et al. ByteTrack: Multi-Object Tracking by Associating Every Detection Box. *European Conference on Computer Vision (ECCV)*. 2022.
5. Newcombe R. A. et al. KinectFusion: Real-time dense surface mapping and tracking. *IEEE International Symposium on Mixed and Augmented Reality (ISMAR)*. 2011. P. 127–136.

ДЕФОРМОВАНІ ЗГОРТКОВІ МЕРЕЖІ

Орленко А.С.¹, Чумаченко О.І.

Національний технічний університет України
«Київський політехнічний інститут ім. Ігоря Сікорського», Київ, Україна

¹ orlenko.anto@gmail.com

У роботі досліджено застосування деформованих згорткових нейронних мереж для автоматизованого сортування твердих побутових відходів. Мета – підвищити точність класифікації та сегментації на конвеєрі за рахунок структурно-параметричного синтезу архітектури на датасеті ZeroWaste. Показано покращення якості розпізнавання порівняно з базовою CNN, що підтверджує придатність підходу для використання у реальних сортувальних лініях.

Ключові слова: штучний інтелект, комп'ютерний зір, деформовані згорткові нейронні мережі, класифікація відходів.

1. ВСТУП

У сучасному світі стрімке зростання обсягів твердих побутових відходів і вимоги до сталого розвитку роблять ефективне сортування сміття однією з ключових інженерних та екологічних задач. Традиційні технології переробки ґрунтуються на комбінації механічних засобів і ручної праці, що обмежує продуктивність, підвищує собівартість та створює небезпечні умови для працівників [1].

Автоматизовані сортувальні лінії з використанням камер та алгоритмів комп'ютерного зору дають змогу розпізнавати й відокремлювати різні типи відходів у реальному часі. Водночас сцени на конвеєрі характеризуються значною різноманітністю та деформаціями об'єктів: пакування зминається, пластик і картон перекривають один одного, змінюються масштаб і орієнтація. За таких умов класичні згорткові нейронні мережі з фіксованою решіткою вибірки не завжди забезпечують потрібну точність класифікації та сегментації.

Метою цієї роботи є дослідження та розробка деформованої згорткової нейронної мережі для автоматизованого сортування твердих побутових відходів. Це передбачає аналіз сучасних підходів до застосування штучного інтелекту в переробці сміття, обґрунтування доцільності використання деформованих згорток у задачах класифікації й сегментації відходів, а також створення та експериментальну оцінку моделі, адаптованої до умов реальної сортувальної лінії [2].

2. АНАЛІЗ СУЧАСНИХ ДОСЛІДЖЕНЬ ТА ПІДХОДІВ

Проблема автоматизованого сортування твердих побутових відходів активно досліджується в контексті підвищення продуктивності та зниження вартості переробки. Перші рішення базувалися переважно на класичних методах комп'ютерного зору та ручно сконструйованих ознаках: колірних гістограмах, текстурних дескрипторах, контурних характеристиках. Такі підходи дозволяли розрізняти окремі типи матеріалів у контрольованих умовах, однак виявилися недостатньо стійкими до змін освітлення, забруднень, фонового шуму та деформацій об'єктів на конвеєрній стрічці. Це обмежило їх практичне застосування в реальних сортувальних лініях.

Подальший розвиток пов'язаний із впровадженням згорткових нейронних мереж (CNN), які стали базовим інструментом для аналізу зображень у задачах класифікації, детекції та

сегментації об'єктів. У роботах, присвячених сортуванню відходів, зазвичай використовуються архітектури на основі відомих моделей (ResNet, EfficientNet тощо) у поєднанні з одноетапними та двоетапними детекторами або сегментаційними мережами (зокрема U-подібними архітектурами). Такі системи демонструють прийнятну точність у випадках, коли об'єкти мають відносно стабільну форму та добре відокремлені від фону. Водночас за наявності сильних деформацій пакування, перекриття об'єктів та їх різноманітної орієнтації якість розпізнавання помітно погіршується, оскільки фіксована решітка вибірки класичних CNN не враховує геометричні варіації сцени.

З метою підвищення робастності моделей до таких варіацій у сучасних дослідженнях запропоновано деформовані згорткові мережі (Deformable Convolutional Networks, DCN). Їх ключова ідея полягає у введенні навчуваних зсувів точок вибірки в межах рецептивного поля та, за потреби, модулюючих коефіцієнтів. Вихід згорткового шару задається формулою (1).

$$y(p) = \sum_{p_k \in R} w(p_k) x(p + p_k + \Delta p_k) \quad (1)$$

Такий підхід дозволяє адаптувати форму ефективного рецептивного поля до контурів об'єктів і зосереджувати обчислення на найбільш інформативних ділянках зображення [3, 4]. Деформовані згортки успішно інтегруються в сучасні детектори та сегментаційні моделі, демонструючи покращення якості на сценах зі значними деформаціями, оклюзіями та складною перспективою. Принцип роботи деформованої згортки зображено на рис. 1.

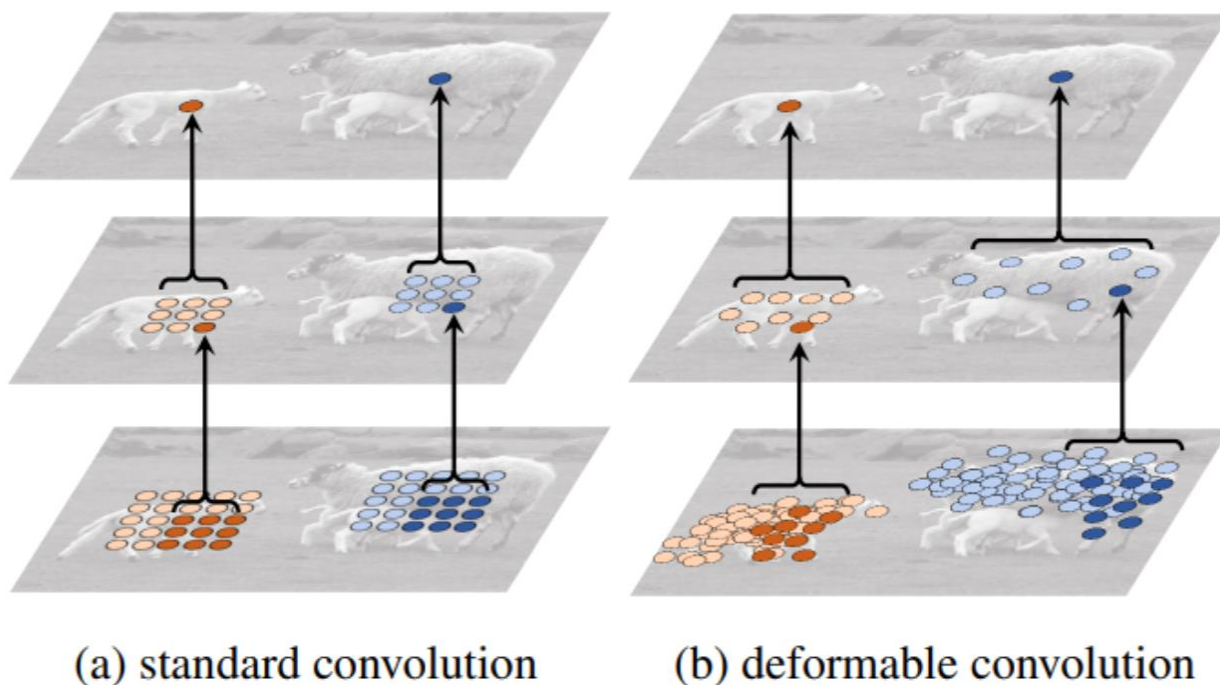


Рисунок 1. Відмінність стандартної згортки від деформованої згортки

Попри досягнутий прогрес, застосування деформованих згорткових мереж саме в задачах сортування побутових відходів досліджено обмежено [4]. Більшість робіт зосереджені або на стандартних CNN без деформацій, або на використанні окремих готових архітектур без цілеспрямованого налаштування їх структури під специфіку сортувальної лінії та обмеження обчислювальних ресурсів. Недостатньо опрацьованими залишаються питання вибору рівнів

мережі, на яких доцільно застосовувати деформовані згортки, узгодження їх параметрів із вимогами до швидкодії, а також адаптації моделей до реальних промислових датасетів, таких як ZeroWaste [5]. Саме вирішення цих задач і становить предмет даного дослідження.

3. МОДЕЛЮВАННЯ ТА РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

Моделювання системи автоматизованого сортування твердих побутових відходів здійснювалося на основі датасету типу ZeroWaste, що містить зображення конвеєрної стрічки з різними видами пакування та домішок [5]. У вибірці представлені об'єкти з картону, жорсткого та м'якого пластику, металу, а також фон у вигляді паперу й елементів обладнання. Зображення характеризуються різноманітними умовами зйомки, частковими перекриттями об'єктів, їх деформаціями та зміною масштабу, що наближує умови моделювання до реальних сцен сортувальної лінії. Перед навчанням мережі виконувалася базова передобробка: масштабування зображень до фіксованого розміру, нормалізація яскравості та застосування простих аугментацій (повороти, горизонтальні віддзеркалення, незначні варіації освітлення), спрямованих на підвищення стійкості моделі до змін умов спостереження.

Водночас деформована архітектура супроводжується певним зростанням обчислювальної складності порівняно з базовою CNN, що проявляється у збільшенні часу інференсу та вимог до апаратних ресурсів. Разом із тим отримане зростання навантаження залишається помірним і може бути прийнятним для багатьох сценаріїв застосування, особливо у випадках, коли критичними є не стільки граничні значення продуктивності, скільки точність та надійність розпізнавання. Загалом результати моделювання свідчать про доцільність використання деформованих згорток у системах комп'ютерного зору для автоматизованого сортування відходів, зокрема в задачах, де сцени мають високий рівень геометричної складності.

Навчання моделі здійснювалося у типовому режимі для задач класифікації та семантичної сегментації з використанням функції втрат, що враховує різницю між прогнозованими та еталонними мітками пікселів або об'єктів. Для уникнення перенавчання застосовувалися стандартні прийоми регуляризації та контролю збіжності на валідаційній вибірці. Особливу увагу приділено збалансуванню класів, оскільки в реальних потоках відходів окремі типи матеріалів можуть бути суттєво недопредставлені, що негативно впливає на якість розпізнавання без відповідних коригувань. Основною метрикою для оцінки якості моделі слугував mIoU (mean Intersection over Union) – середнє значення показника перетину над об'єднанням областей (Intersection over Union) по всіх класах. Для окремого класу IoU обчислюється як відношення кількості пікселів, які одночасно належать цьому класу в еталонній розмітці та в прогнозі моделі (перетин), до загальної кількості пікселів, що належать цьому класу хоча б в одній із розміток (об'єднання). Значення mIoU лежить у діапазоні від 0 до 1, де більші значення відповідають кращій якості сегментації, тобто більш точному просторовому збігу прогнозованих масок з еталонними.

Отримана модель здатна впевнено розпізнавати та сегментувати об'єкти на рухомій конвеєрній стрічці, досягаючи значень mIoU > 0.5 на валідаційній вибірці. Приклад результату сегментації представлено на рис. 2, у той час як на рис. 3 представлено еталонну розмітку зображення.

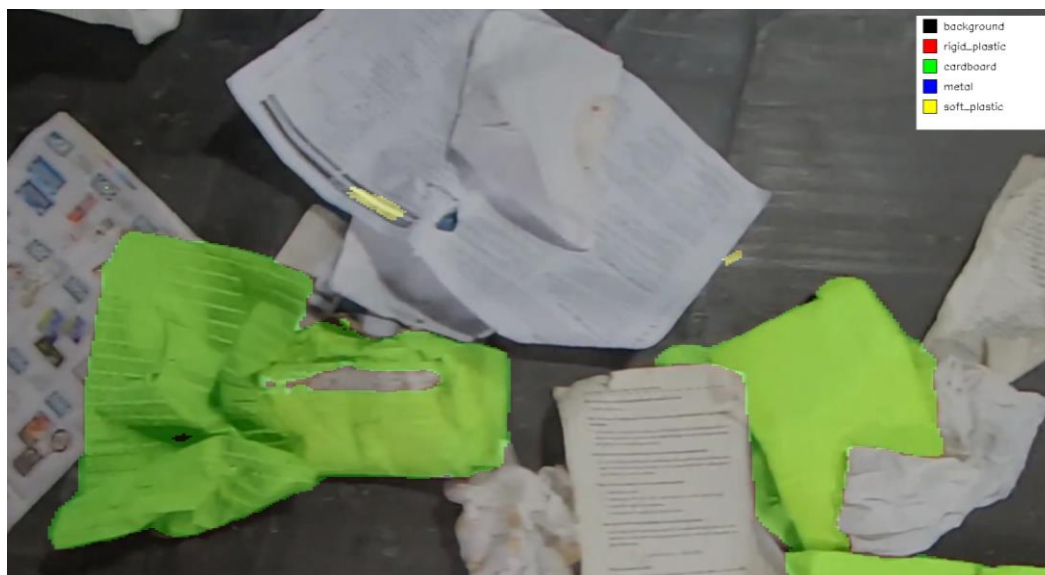


Рисунок 2. Приклад результату сегментації

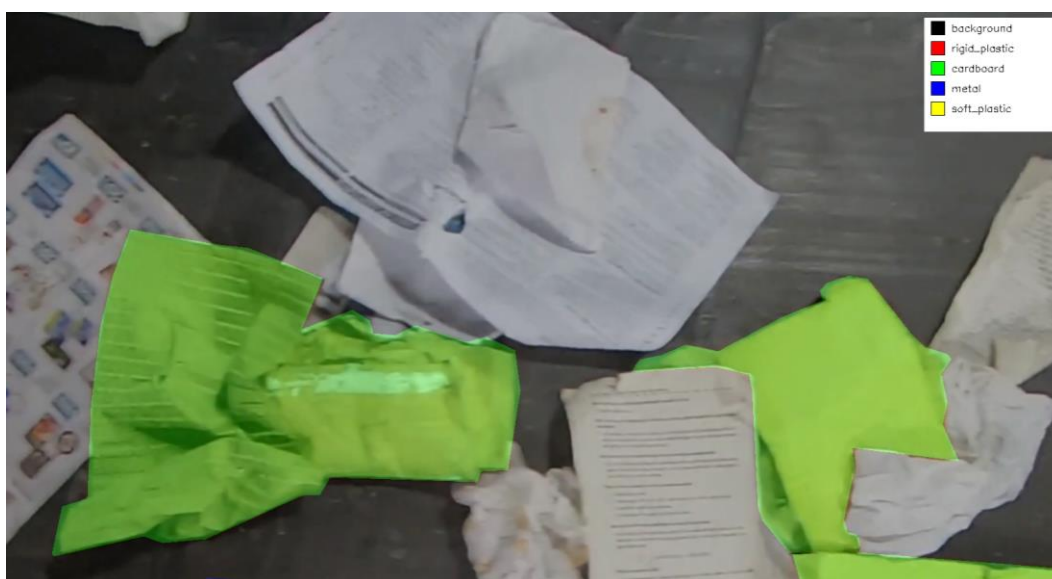


Рисунок 3. Еталонна розмітка зображення

4. ВИСНОВКИ

У роботі розглянуто задачу автоматизованого сортування твердих побутових відходів на основі методів комп'ютерного зору та глибокого навчання, а також проаналізовано можливості деформованих згорткових нейронних мереж у цьому контексті. Показано, що класичні CNN з фіксованою решіткою вибірки обмежені у здатності коректно відображати геометричні варіації об'єктів сміття, які проявляються у вигляді деформацій, перекриттів, змін масштабу та перспективних спотворень у реальних сортувальних сценах. Це зумовлює потребу в архітектурах, здатних адаптивно підлаштовувати рецептивне поле до форми й положення об'єктів.

Запропоновано підхід до побудови деформованої згорткової нейронної мережі для задач класифікації та сегментації відходів, який ґрунтується на модифікації базової архітектури шляхом введення деформованих шарів на окремих рівнях мережі. Використання датасету типу ZeroWaste дозволило змодельовати роботу системи в умовах, наближених до реальної

сортувальної лінії, із характерними для неї деформаціями об'єктів, їх перекриттями та різномірністю фону. На основі результатів моделювання можна зробити висновок, що деформована мережа має потенціал забезпечувати більш стійке та точне розпізнавання об'єктів у складних сценах порівняно з базовою CNN, особливо у випадках сильних деформацій і зашареного зображення.

Разом із покращенням якості розпізнавання деформовані згорткові мережі вимагають дещо більших обчислювальних ресурсів, що слід враховувати при їх подальшій інтеграції в промислові системи сортування. Надалі доцільно зосередитися на оптимізації архітектури з погляду швидкодії, дослідженні інших комбінацій рівнів, на яких застосовуються деформовані згортки, а також на розширенні експериментів за рахунок різних типів відходів та додаткових датасетів. Отримані результати можуть бути використані як основа для розробки практичних рішень у галузі автоматизованої переробки сміття та подальших досліджень методів глибокого навчання для складних техногенних сцен.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Huynh M.-H., Pham-Hoai P.-T., Tran A.-K., Nguyen T.-D. Automated Waste Sorting Using Convolutional Neural Network // 2020 7th NAFOSTED Conference on Information and Computer Science (NICS). – IEEE, 2020. – P. 102–107. – DOI: 10.1109/NICS51282.2020.9335897.
2. Laksono P. W., Anisa A., Priyandari Y. Deep Learning Implementation Using Convolutional Neural Network in Inorganic Packaging Waste Sorting // Franklin Open. – 2024. – Vol. 8. – Article 100146. – DOI: 10.1016/j.fraope.2024.100146.
3. Dai J., Qi H., Xiong Y., Li Y., Zhang G., Hu H., Wei Y. Deformable Convolutional Networks // 2017 IEEE International Conference on Computer Vision (ICCV). – Venice, Italy, 2017. – P. 764–773. – DOI: 10.1109/ICCV.2017.89.
4. Zhu X., Hu H., Lin S., Dai J. Deformable ConvNets v2: More Deformable, Better Results // 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). – Long Beach, CA, USA, 2019. – P. 9300–9308. – DOI: 10.1109/CVPR.2019.00953.
5. Bashkirova D., Abdelfattah M., Zhu Z., Akl J., Alladkani F., Hu P., Ablavsky V., Calli B., Bargal S. A., Saenko K. ZeroWaste Dataset: Towards Deformable Object Segmentation in Cluttered Scenes // 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). – New Orleans, LA, USA, 2022. – P. 2047–2057. – DOI: 10.1109/CVPR52688.2022.02047.

ІНТЕЛЕКТУАЛЬНА СИСТЕМА РОЗПІЗНАВАННЯ ЕМОЦІЙ ЗА ВИРАЗАМИ ОБЛИЧЧЯ НА ОСНОВІ ЗГОРТКОВИХ НЕЙРОННИХ МЕРЕЖ

Панасенко В.В.¹, Тимошук О.Л.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ panvovandrik@gmail.com

Ця робота присвячена розробленню та дослідженню інтелектуальної системи розпізнавання емоцій за виразами обличчя (Facial Emotion Recognition, FER), побудованої на основі згорткових нейронних мереж. Актуальність теми зумовлена зростаючим запитом на автоматичний аналіз емоційних станів користувача у системах взаємодії людина-комп'ютер, телемедицині, дистанційному навчанні, медіа-аналітиці та безпекових застосуваннях.

Ключові слова: аналіз зображень, штучний інтелект, ідентифікація емоцій.

1. ВСТУП

У сучасних інформаційних системах, що взаємодіють із людиною, зростає потреба в автоматичному розумінні емоційних станів користувача. Розпізнавання емоцій за виразами обличчя (Facial Expression Recognition, FER) є одним із ключових напрямів афективних обчислень і знаходить застосування у широкому спектрі задач: від адаптивних інтерфейсів і систем дистанційного навчання до медіа-аналітики, телемедицини та безпечних транспортних систем. Висока варіативність зовнішнього вигляду людей, освітлення, поз, часткових оклюзій і культурних відмінностей створює суттєві виклики для побудови точних і робастних FER-моделей.

Останнє десятиліття продемонструвало перевагу згорткових нейронних мереж (Convolutional Neural Networks, CNN) над класичними підходами (LBP/HOG + SVM), насамперед завдяки автоматичному виділенню просторових ознак та здатності узагальнювати на великих неідеальних вибірках. Разом з тим, для практичних застосувань важливими залишаються обмеження на обчислювальні ресурси й затримку (latency), потреба в чутливості до тонких мікро-експресій і зменшення упередженості (bias) щодо демографічних груп.

Актуальність роботи полягає у створенні інтелектуальної системи FER, яка поєднує достатню точність класифікації з простотою реалізації та можливістю подальшої оптимізації під обмежені обчислювальні середовища (настільні системи без GPU, мобільні пристрої). У роботі використано компактний CNN-підхід із нормалізацією пакетів (Batch Normalization) та Dropout-регуляризациєю, а також стандартний пайплайн підготовки даних і навчання в середовищі Keras/TensorFlow.

Мета роботи – спроектувати, реалізувати та експериментально перевірити інтелектуальну систему розпізнавання емоцій за виразами обличчя на основі CNN для 7 базових класів емоцій з використанням зображень розміру 48×48 у відтінках сірого.

2. ПОСТАНОВКА ЗАДАЧІ

Метою роботи є розробка та експериментальна перевірка інтелектуальної системи розпізнавання емоцій людини за виразами обличчя на основі згорткових нейронних мереж. Система повинна автоматично класифікувати зображення облич на один із семи базових емоційних станів (angry, disgust, fear, happy, sad, surprise, neutral), що відповідають категоріальній моделі Екмана.

3. ПІДХОДИ ДО РЕАЛІЗАЦІЇ ТА ГОТОВІ МОДЕЛІ

Сучасні системи розпізнавання емоцій базуються на методах комп'ютерного зору та глибинного навчання, що забезпечують автоматичне виділення ознак і стійкість до варіацій зовнішнього вигляду людини. У роботі проаналізовано два основні підходи: класичні методи та згорткові нейронні мережі.

1. Класичні методи FER. Вони передбачають послідовність кроків: детекція обличчя, витягнення ручних дескрипторів (LBP, HOG, Gabor), зниження розмірності (PCA/LDA) та класифікація (SVM). Перевагою цих методів є простота та ефективність на контрольованих даних, однак вони виявляються крихкими на «польових» зображеннях через варіації освітлення, поз, оклюзії та мікроекспресії.

2. CNN-підхід. Згорткові нейронні мережі дозволяють навчати ознаки та класифікатор кінець-у-кінець, що забезпечує їхню стійкість і узагальнювальну здатність. У роботі використано компактну модель формату Conv2D–BatchNorm–ReLU–MaxPool–Dropout із двома щільними шарами та вихідним softmax на сім класів. Мережа має близько 4.48 млн параметрів, працює з одноканальними зображеннями 48×48 та демонструє високу точність на валідаційних даних.

3. Популярні готові моделі та бібліотеки, що застосовуються у FER: FER2013 CNN-архітектури – класичні мережі для зображень 48×48. VGG-Face, ResNet, MobileNet – моделі для екстракції ознак обличчя. TensorFlow/Keras, PyTorch – для побудови та навчання моделей.

Використання цих інструментів значно прискорює розробку системи та дозволяє зосередитися на експериментальному аналізі.

4. ПРИКЛАДИ РОБОТИ ПРОГРАМНОГО ПРОДУКТУ

У межах роботи було реалізовано повний програмний комплекс для розпізнавання емоцій за виразами обличчя, який включає завантаження та підготовку даних, навчання CNN-моделі, візуалізацію результатів та оцінку точності. Усі експерименти виконано мовою Python у середовищі Google Colab із використанням бібліотек TensorFlow/Keras.

Під час препроцесингу всі зображення було приведено до єдиного формату: обличчя перетворювалися у відтінки сірого та масштабувалися до розміру 48×48 пікселів із нормалізацією інтенсивностей у діапазон [0;1]. Така уніфікація дозволила моделі працювати стабільно навіть за умов різних освітлень, поз чи варіацій обличчя. Вихідні підготовлені зображення семи емоцій виглядають як на рис. 1, де продемонстровано приклади всіх класів після нормалізації.



Рисунок 1. Основних 7 емоцій з нормалізацією до 48x48 розміру

Після навчання модель формувала розподіл ймовірностей по кожній із семи категорій, а фінальним результатом ставала емоція з найбільшим значенням. Для візуального розуміння того, які патерни міміки характерні для кожного класу, використовувалися базові приклади емоцій за Екманом (рис. 2). Зображення демонструють характерні зміни в області брів, очей, носогубного трикутника та рота, які використовуються моделлю під час розпізнавання.



Рисунок 2. «Базові емоції за Екманом (6 класів: sadness, contempt, surprise, anger, disgust, fear [7])

Важливою частиною аналізу результатів стало врахування індивідуальних та культурних відмінностей, оскільки один і той самий емоційний стан може проявлятися в людей по-різному (рис. 3). У дипломі наведено приклади варіативності, які демонструють різні стилі виразу емоцій, що ускладнює роботу моделі, але робить задачу більш реалістичною.



Рисунок 3. Культурна та індивідуальна варіативність: приклади одного класу з різних демографічних груп

У роботі також представлено циркумплексну модель емоцій (рис. 4), яка допомагає пояснити близькість деяких класів у психологічному просторі валентності та активації. Саме ці близькі позиції стають причиною регулярних плутанин під час класифікації, оскільки емоції мають подібні зовнішні ознаки.

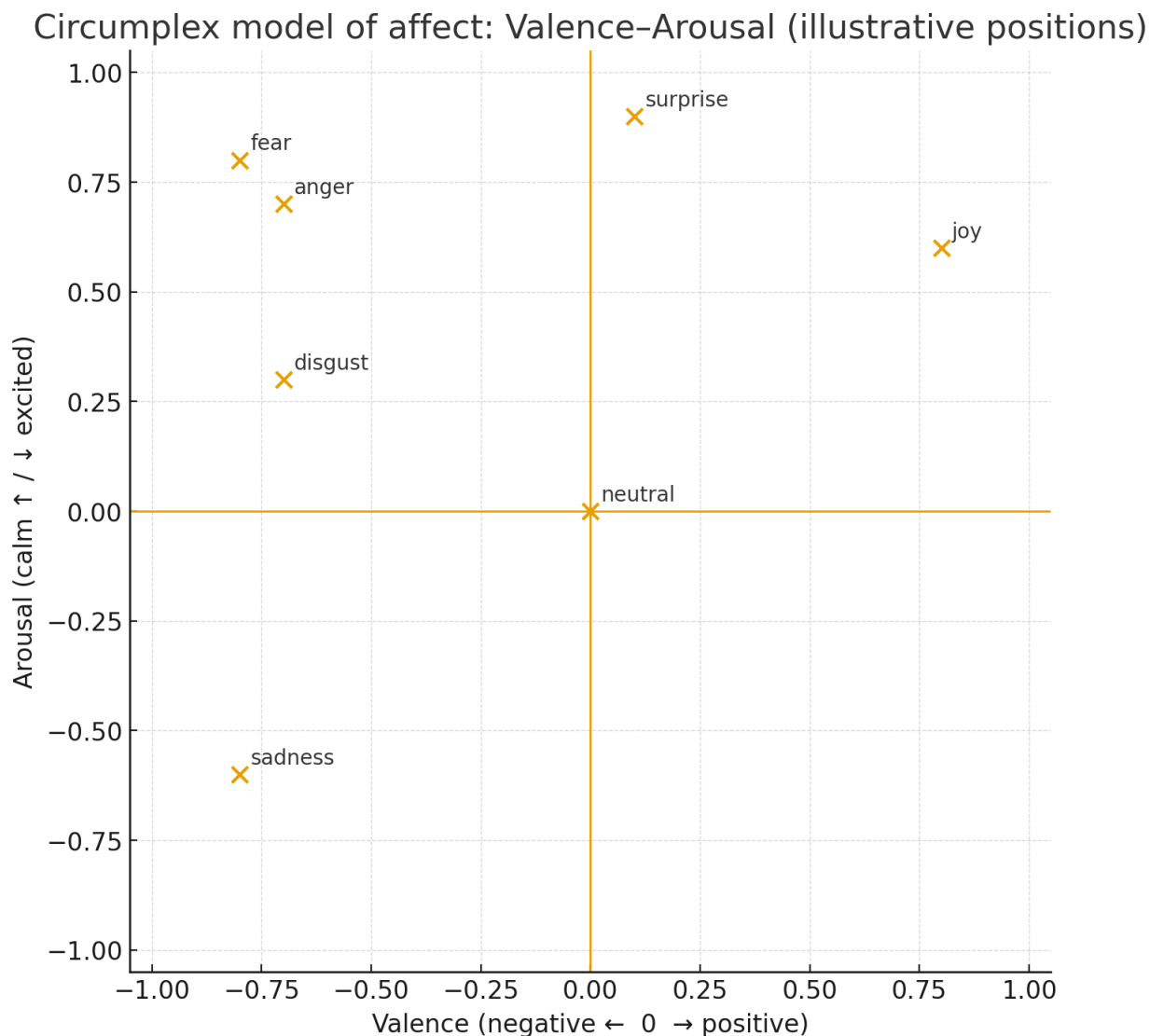


Рисунок 4. Циркумплексна модель (валентність-активація) з орієнтовним розміщенням базових емоцій

При оцінюванні моделі було виявлено типові помилки, пов'язані з мікроекспресіями та низькою інтенсивністю емоцій. Модель часто плутала пари емоцій *fear-surprise* через подібність очних щілин та піднятих брів, *sad-neutral* через мінімальну активність м'язів, а також *disgust-anger*, де обидві емоції мають напруження в носогубній області. Для пояснення відмінностей між метриками та причин помилок на рис. 5 наведено інфографіку, що ілюструє різницю між точністю та точністю передбачення.

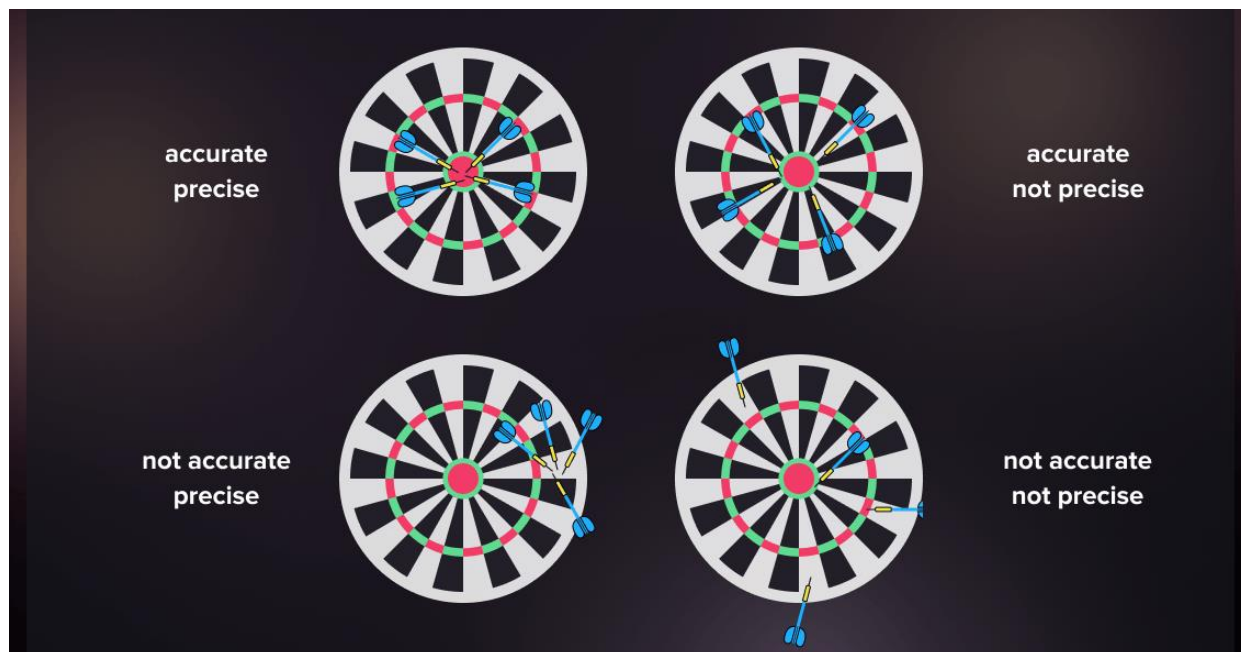


Рисунок 5. Приклад: Accuracy і Precision відмінність

Узагалі модель демонструє впевнені результати для емоцій з чітко вираженими патернами, зокрема *happy* та *surprise*, тоді як слабо виражені або змішані емоції викликають труднощі. Це також підтверджується аналізом матриці плутань і значеннями Macro-F1 для рідкісних класів. У випадках правильних передбачень модель стабільно фіксує характерні ознаки, а в хибних – найчастіше помиляється на зображеннях із нечіткими межами міміки, некоректним освітленням або різкими тінями. Ілюстрації прикладів входів, на яких модель працює найкраще та найгірше, можна спостерігати на підготовлених прикладах із нормалізованими емоціями.

5. ВИСНОВКИ

В даній роботі було реалізовано модель штучного інтелекту для розпізнавання емоцій за виразами обличчя. У процесі виконання роботи було розглянуто різні підходи до аналізу мімічних ознак та обрано оптимальний метод на основі згорткових нейронних мереж, який і був реалізований. Результатом роботи стала модель, здатна класифікувати сім базових емоцій людини. На основі отриманих експериментальних результатів були визначені сильні та слабкі сторони використаного підходу та окреслено можливі напрями подальшого вдосконалення системи.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Ekman P., Friesen W. V. Constants across cultures in the face and emotion. *Journal of Personality and Social Psychology*, 1971.
2. Лазуренко О. О. Психологія емоцій. Київ: Книга плюс, 2018. – 262с.
3. Convolutional Neural Network (CNN) (2023) [Електронний ресурс] – Режим доступу до ресурсу: <https://developer.nvidia.com/discover/convolutionalneural-network>
4. TensorFlow. <https://www.tensorflow.org/> (дата звернення 23.10.2026)
5. Chollet F. Keras. 2015–2025. Документація та програмне забезпечення.
6. Preece J., Rogers Y., Sharp H. *Interaction Design: Beyond Human–Computer Interaction*. Wiley, 2019.
7. Ekman, Friesen (1971) – розділ 1.1 (опис базових емоцій і універсальність виразів)

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ДІАГНОСТИКИ ОСТЕОПОРОЗУ НА ОСНОВІ ОБРОБКИ DXA ЗОБРАЖЕНЬ

Семенов О.Г.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна
aleksey.semenow.fb05@gmail.com

Метою дослідження було створення інтелектуальної системи, здатної класифікувати DXA-зображення хребців L1–L4 як норми чи остеопороз, спираючись на структурну вразливість кістки на основі кількісного аналізу її мікроархітектури, витягнутої з текстури знімка. Ключовим методологічним аспектом стало застосування повної параметризації анізотропної варіограмної моделі з nugget variance (c_0), sill variance (c), correlation length (L), major radius (L_1), minor radius (L_2) та ступеня анізотропії $DA = L_1/L_2$, а також TBS.

1. ВСТУП

У дослідженні представлено підхід до діагностики остеопорозу, що виходить за межі класичного використання площинної мінеральної щільності (aBMD), і вводить розширений структурний аналіз DXA-зображень тіл хребців L1–L4. Центральним елементом методології є оцінка мікроархітектури трабекулярної кістки через текстурні характеристики розподілу інтенсивності пікселів – варіограмні параметри, які відображають неоднорідність, кореляційні зв'язки та орієнтаційні властивості трабекул.

Оцінка стану кістки виконується не шляхом аналізу одного критерію, а через використання повного набору варіограмних характеристик, серед яких sill variance, nugget variance, correlation length, ТВ, а також параметри анізотропії, включаючи велике і мале кореляційні радіуси, та похідний показник ступеня анізотропії. Це дає змогу відтворити статистичну топографію внутрішньої трабекулярної структури, що важливо для визначення напрямкового ослаблення кістки.

Підхід передбачає кілька етапів обробки. Спочатку формується ROI кожного хребця L1–L4 з повним видаленням сторонніх тканин, ребер, міжхребцевих дисків та фонового шуму, що усуває спотворення статистики. Після цього інтенсивності нормалізуються та стандартизуються, через Z-перетворення, що дозволяє зіставляти різні DXA-скани між собою незалежно від умов та апаратури. Для приглушення шуму застосовується фільтрація, яка знижує випадкову флуктуацію пікселів, але не руйнує просторові градієнти, що несуть інформацію про структуру трабекулярної мережі.

Обчислення варіограм проводиться для рядів дельта-значень між пікселями при різних лагових відстанях h та при різних кутових напрямках θ у випадку анізотропного аналізу. Це дозволяє реконструювати форму варіограми, знаходити sill-рівень, оцінювати ступінь корельованості пікселів на різних масштабах та визначати просторові закономірності внутрішнього переплетення трабекул. Параметри варіограми подаються у вигляді числового вектору, який використовується у нейронній моделі як повноцінний вхідний блок ознак.

В отриманій системі згортова нейронна мережа, крім самого DXA-зображення, приймає як додаткові входи всі варіограмні показники, а не лише TBS. Це реалізовано шляхом конкатенації структурних параметрів із ознаками, отриманими після шару Flatten, що дозволяє мережі комбінувати просторові візуальні патерни зі статистичними

дескрипторами мікроархітектури. Таким чином модель може аналізувати два рівні репрезентації: візуальну геометрію трабекул і статистичну організацію їх зв'язків.

2. ПІДХІД ДО РЕАЛІЗАЦІЇ

Ідея системи полягає у тому, що зображення тіл хребців L1–L4 розглядаються не як статичні площинні рентгеновські проекції, а як носії багатовимірної інформації, де локальні коливання інтенсивності пікселів відображають мікроархітектурний стан трабекулярної мережі. Реалізація розділяється на послідовні етапи: структурна очистка даних, нормалізація сигналу, обчислення варіограмних характеристик, просторове вирівнювання та CNN-класифікація з включенням числових структурних ознак.

Першим етапом є ізоляція ROI кожного окремого хребця. На цьому рівні алгоритм прибирає вплив міжхребцевих дисків, кортикальних обмежень, ребер та м'яких тканин, що дозволяє уникнути зайвого зашумлення, яке спотворює статистику. Заздалегідь окреслена область аналізу зберігає внутрішню неоднорідність, яка є предметом дослідження, але виключає зовнішні низькоінформативні сегменти.

Другим етапом є уніфікація зображення, оскільки абсолютні значення пікселів можуть різнитися в залежності від сканера, параметрів експозиції та особливостей пацієнта, застосовується нормалізація або стандартизація, що переводить усі зображення в узгоджений числовий простір. Це усуває апаратні артефакти і зміщує фокус моделі на внутрішню текстурну структуру замість глобального рівня щільності.

Далі виконується розрахунок варіограмних параметрів, які кількісно описують просторову залежність між пікселями. У цьому дослідженні розраховуються всі варіограмні характеристики: sill variance, nugget variance, correlation length, початковий нахил варіограми, а також анізотропні компоненти L_1 , L_2 та ступінь анізотропії DA (рис. 1–4). Це дозволяє оцінити структуру не лише в середньому, але і в напрямній формі – як змінюється корельована геометрія трабекул відносно різних кутів.

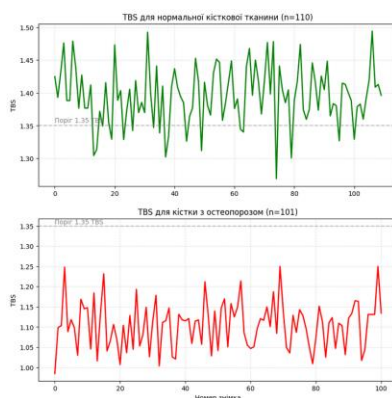


Рисунок 1. Порівняння значень TBS здорових хребців і вражених остеопорозом для даних із датасету

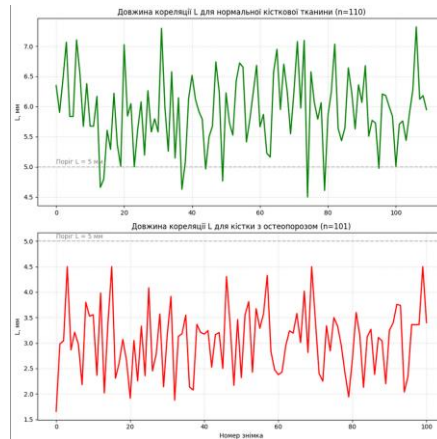


Рисунок 2. Порівняння значень довжини кореляції здорових хребців і вражених остеопорозом для даних із датасету

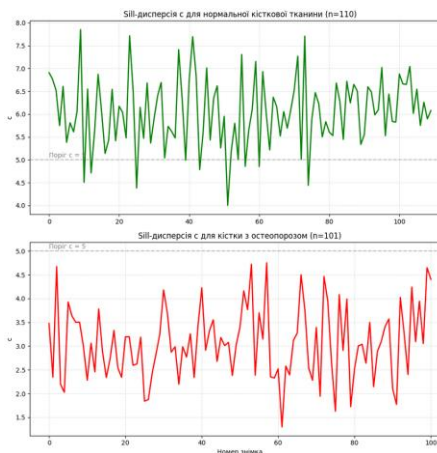


Рисунок 3. Порівняння значень sill-дисперсії здорових хребців і вражених остеопорозом для даних із датасету

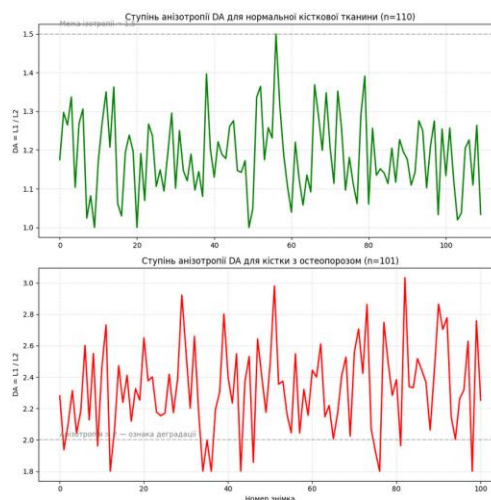


Рисунок 4. Порівняння значень ступеня анізотропії здорових хребців і вражених остеопорозом для даних із датасету

На наступному етапі нормалізоване зображення найменш інформаційної ентропії трансформується у формат вхідного тензора для згорткової нейронної мережі. При цьому, модель отримує не тільки зображення, але і окремий числовий вектор варіограмних ознак,

які додаються в кінець Flatten-представлення CNN перед Dense-модулем. Це дає можливість поєднувати просторові закономірності, виявлені згортковими фільтрами, із статистичними описами мікроархітектури.

3. ПРИКЛАД РОБОТИ

У результаті виконання моделі для набору тестових зображень система генерує класифікаційний висновок, зазначаючи ймовірність належності кожного хребця до класу «норма» або «остеопороз». Крім цього, при виявленні остеопоротичного стану з високою вірогідністю система додатково вказує, який саме хребець має найбільшу структурну вразливість. У випадку, коли аналізувалися чотири хребці одного пацієнта, застосовувалося правило консенсусної класифікації: якщо щонайменше три ROI демонструють нормальні показники, вся група позначається як «норма»; якщо ж троє з чотирьох свідчать про остеопоротичні зміни – загальний статус визначається як «остеопороз». У випадку симетричного патерну 2/2, за умовами моделі, застосовується консервативний підхід, і хребтовий сегмент класифікується як остеопороз, оскільки навіть часткова втрата мікроархітектурної стабільності може вказувати на підвищений ризик перелому.

На основі моделі було також реалізовано веб-інтерфейс у середовищі Django, що дозволяє застосовувати систему безпосередньо в клінічному або дослідницькому процесі. Інтерфейс підтримує завантаження PNG-зображень хребців із попередньо позначеними ROI-зонами (рис. 5), після чого виконується весь аналітичний цикл: обробка ROI, нормалізація інтенсивності, обчислення варіограмних показників та інференція моделі.

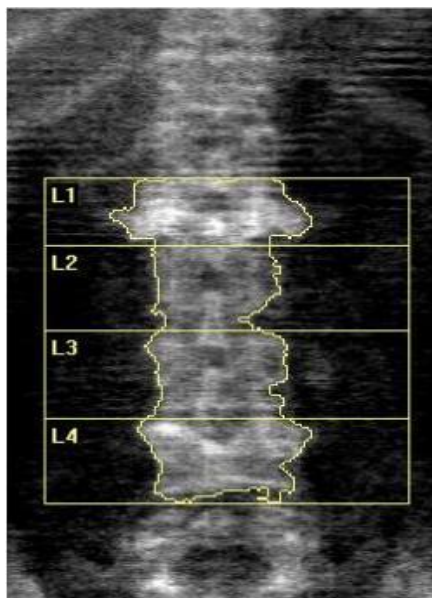


Рисунок 5. Приклад вхідного DXA-зображення

Результатом роботи веб-інструменту є сформований звіт, де для кожного із чотирьох хребців виводиться значення ймовірності остеопорозу, порівняння із внутрішнім порогом моделі, а також додаткові текстурні характеристики – sill variance, correlation length, ступінь анізотропії (DA) і похідні показники. Якщо для певного хребця фіксується високий рівень структурної порушеності, система підсвічує його в результатах як найбільш уразливий, дозволяючи користувачеві швидко ідентифікувати критичну зону (рис. 6).

L1: P(osteoporosis)=0.8726 → Osteoporosis
L2: P(osteoporosis)=0.9910 → Osteoporosis
L3: P(osteoporosis)=0.9999 → Osteoporosis
L4: P(osteoporosis)=0.9929 → Osteoporosis
Найвірогідніше уражений: L3 (P=0.9999) — Osteoporosis

Рисунок 6. Реакція моделі на тестове DXA-зображення

4. ВИСНОВКИ

Таким чином, у цій роботі було реалізовано модель інтелектуальної нейронної мережі, навчену на DXA зображеннях в поєднанні з варіограмними показниками, здатну аналізувати картинки та виявляти які з хребців з більшою вірогідністю вражені остеопорозом.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Lee D. O., Hong Y. H., Cho M. K., Choi Y. S., Chun S., Chung Y. J., ... Lee J. Y. (2024). *The 2024 Guidelines for Osteoporosis – Korean Society of Menopause: Part I*. Journal of Menopausal Medicine, 30(1), 1–23. <https://doi.org/10.6118/jmm.24000>
2. Каніс J. A., МакКлоскі E. V., Джонссон H., Купер С., Піззолі R., & Reginster J.-Y. (2018). "European guidance for the diagnosis and management of osteoporosis in postmenopausal women". *Osteoporosis International*, 29(2), 176–http://... (повна сторінка не вказана). <https://doi.org/10.1007/s00198-012-2074-y>
3. Martineau P., & Leslie W. D. (2017). "Clinical Utility of Using Lumbar Spine Trabecular Bone Score to Adjust Fracture Probability: The Manitoba BMD Cohort". *Journal of Bone and Mineral Research*, 32(7), 1568–1574. <https://doi.org/10.1002/jbmr.3124>
4. Freemantle N., Cooper C., Diez-Pérez A., Gitlin M., Radcliffe H., Shepherd S., & Roux C. (2013). "Results of indirect and mixed treatment comparison of fracture efficacy for osteoporosis treatments: a meta-analysis". *Osteoporosis International*, 24(1), 209–217. <https://doi.org/10.1007/s00198-012-2068-9>

ВИЯВЛЕННЯ КОНТРАСТНИХ ОБ'ЄКТІВ З ВИКОРИСТАННЯМ КОМБІНАЦІЇ НАПІВКОНТРОЛЬОВАНОГО ТА НЕКОНТРОЛЬОВАНОГО НАВЧАННЯ

Степанчук Д.К.¹, Пишнограєв І.О.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ stepanchuk.dmytro@lil.kpi.ua, ² pyshnograiev@nas.gov.ua

Робота присвячена дослідженню методів напівконтрольованої детекції об'єктів (SSOD) для виявлення контрастних областей зображень в умовах обмеженої кількості розмічених даних. Проаналізовано еволюцію функцій втрат регресії обмежувальних рамок ($\text{IoU} \rightarrow \text{GIoU} \rightarrow \text{DIOU} \rightarrow \text{CIoU}$), механізми Teacher-Student фреймворків та архітектури сучасних детекторів (Faster R-CNN, YOLOv8, DETR). Обґрунтовано переваги індуктивного підходу SSOD для підвищення узагальнювальної здатності моделей. Мета роботи – аналіз теоретичного підґрунтя та вибір оптимальної комбінації архітектурних рішень для ефективної детекції контрастних об'єктів.

Ключові слова: напівконтрольоване навчання, детекція об'єктів, SSOD, Teacher-Student, CIoU, YOLOv8, DETR, кластеризація, навчання без учителя.

1. ВСТУП

Виявлення контрастних об'єктів – задача визначення областей зображення, що суттєво відрізняються від оточення за візуальними характеристиками – є фундаментальною проблемою комп'ютерного зору із застосуваннями у медичній діагностиці, промисловому контролі якості та автономній навігації. Класичні методи контрольованого навчання демонструють високу точність, але вимагають значних обсягів ретельно анотованих даних, що робить їх економічно неефективними.

Гіпотеза дослідження полягає в тому, що комбінування напівконтрольованого навчання (semi-supervised learning) з елементами неконтрольованого дозволить ефективно використовувати невелику кількість розмічених даних у поєднанні з великим обсягом нерозмічених зразків, підвищуючи узагальнювальну здатність моделі.

2. МЕТОДОЛОГІЯ ТА АРХІТЕКТУРА МОДЕЛЕЙ

2.1. Базові моделі

Розвиток функцій втрат для регресії обмежувальних рамок демонструє перехід від простого наближення до геометричної узгодженості (табл. 1):

Таблиця 1. Функції похибки для object detection

Функція	Геометричні фактори	Критичні недоліки
IoU	Перетин	Нульовий градієнт при відсутності перетину
GIoU	Перетин, охоплюючий контейнер	Повільна збіжність, деградація до IoU
DIoU	Перетин, відстань центрів	Ігнорує співвідношення сторін
CIoU	Перетин, відстань центрів, aspect ratio	Складність формули

CIoU Loss враховує три геометричні фактори та забезпечує найточнішу регресію. Додатково, Distributional Focal Loss (DFL) моделює невизначеність локалізації шляхом прогнозування розподілу координат.

2.2. Co-training: навчання з множинних поглядів

Co-training базується на припущенні, що вхідні дані можна розділити на два або більше наборів ознак (views), які є умовно незалежними за умови знання класу. Тренуються кілька моделей, кожна на своєму наборі ознак. Моделі співпрацюють, обмінюючись висококонфідентними прогнозами (псевдо-мітками) для нерозмічених даних, які додаються до розміченого набору. Цей ітеративний процес підвищує робастність системи.

У контексті детекції об'єктів природний поділ на незалежні види ознак є складним завданням. Сучасні адаптації, зокрема Multi-Head Co-Training, долають це обмеження через інтеграцію базових навчальних елементів у багатоголову структуру. Різноманітність досягається завдяки стратегії «Слабка та Сильна Аугментація»: одна модель отримує слабо аугментовані зображення (flip, crop), інша – сильно трансформовані (color jitter, cutout, mixup). Це запобігає колапсу моделей до ідентичних результатів та забезпечує взаємне збагачення через обмін псевдо-мітками.

2.3. Архітектури детекторів

Двоетапні детектори (Cascade R-CNN) – використовують каскад детекторів з прогресивно вищими порогами IoU для ітеративного уточнення рамок.

Одноступеневі детектори (YOLOv8) – прогнозують класи та рамки за один прохід. Архітектура включає блоки C2f (оптимізація градієнтного потоку) та SPPF (масштабна інваріантність). Використовує BCE + CIoU + DFL.

Трансформерні детектори (DETR) – End-to-End підхід без NMS та якорів. Угорський алгоритм забезпечує двобічне узгодження, але призводить до розрідженої супервізії та повільної збіжності (~500 epoch).

3. ПАЙПЛАЙН ІТЕРАТИВНОГО НАВЧАННЯ

3.1. Етап 1: Неконтрольована авторозмітка

Початковий етап експлуатує внутрішню властивість контрастних об'єктів – їх візуальну відмінність від оточення. Для автоматичного виявлення таких зон застосовується кластеризація у 5-вимірному просторі ознак:

$$\mathbf{F} = (\mathbf{x}, \mathbf{y}, \mathbf{R}, \mathbf{G}, \mathbf{B}),$$

де $(\mathbf{x}, \mathbf{y})$ – просторові координати пікселя, $(\mathbf{R}, \mathbf{G}, \mathbf{B})$ – значення кольорових каналів. Алгоритм кластеризації групує пікселі за евклідовою відстанню у цьому просторі, автоматично виділяючи компактні області з подібними просторово-колірними характеристиками.

Кластери, що суттєво відрізняються від домінуючого фону за кольором або просторовим розташуванням, ідентифікуються як потенційні контрастні зони.

3.2. Етап 2: Ініціалізація базового детектора

На основі мінімального розміченого набору (20 зображень з ручною анотацією) навчається базовий детектор. Цей детектор застосовується до зон, виявлених на етапі неконтрольованої кластеризації, генеруючи початкові псевдо-мітки з класифікацією та уточненими bounding boxes. Таким чином, неконтрольоване навчання забезпечує локалізацію потенційних об'єктів, а базовий детектор – їх класифікацію та геометричне уточнення.

3.3. Етап 3: Напівконтрольоване навчання

На отриманому наборі псевдо-міток паралельно навчаються кілька мереж із застосуванням **co-training** з різними аугментаціями для забезпечення різноманітності моделей. Моделі обмінюються високоякісними псевдо-мітками, взаємно покращуючи узагальнювальну здатність.

3.4. Етап 4: Відбір порогу впевненості та ручна дорозмітка

Після завершення раунду навчання емпірично визначається оптимальний поріг впевненості моделей. Зразки поділяються на три категорії:

- **Висока впевненість** ($p > \tau_{\text{high}}$) – псевдо-мітки приймаються автоматично;
- **Низька впевненість** ($p < \tau_{\text{low}}$) – зразки відкидаються як фонові;
- **Невизначені випадки** ($\tau_{\text{low}} \leq p \leq \tau_{\text{high}}$) – направляються на ручну дорозмітку експертом.

Така стратегія активного навчання концентрує людські зусилля виключно на складних, неоднозначних випадках, максимізуючи ефективність анотації.

3.5. Ітеративний цикл

Етапи 3–4 повторюються ітеративно: кожен новий раунд навчання використовує розширений набір даних (автоматичні псевдо-мітки + ручна дорозмітка складних випадків). З кожною ітерацією модель покращується, зменшуючи кількість невизначених випадків та потребу в ручному втручанні.

3.6. Валідація

Оцінка якості моделі проводиться двома комплементарними методами:

- **Крос-валідація між епохами** – порівняння прогнозів поточної моделі з моделями попередніх ітерацій для виявлення регресії або нестабільності;
- **Експертне оцінювання** – вибіркова перевірка результатів детекції людиною-експертом для контролю якості та виявлення систематичних помилок.

4. ВИСНОВОК

Дослідження обґрунтовує, що успіх SSOD критично залежить від якості локалізації Teacher-моделі. Еволюція функцій втрат від GloU до CIoU та DFL забезпечує геометричну точність та моделювання невизначеності, необхідні для генерації нешумних псевдо-міток.

Встановлено архітектурний компроміс: трансформерні детектори (DETR) пропонують теоретично елегантний End-to-End підхід, однак одноступеневі CNN-детектори (YOLOv8) залишаються кращим вибором для SSOD через ефективність тренування та високу швидкість інференсу.

Подальші дослідження спрямовані на практичну реалізацію гібридної системи, що поєднує co-training фреймворк з механізмами кластеризації для адаптивного виявлення контрастних об'єктів.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Liu Y.-C. et al. Unbiased Teacher v2: Semi-Supervised Object Detection for Anchor-Free and Anchor-Based Detectors. CVPR, 2022.
2. Xu M. et al. End-to-End Semi-Supervised Object Detection with Soft Teacher. ICCV, 2021.
3. Zheng Z. et al. Distance-IoU Loss: Faster and Better Learning for Bounding Box Regression. AAAI, 2020.
4. Li X. et al. Generalized Focal Loss: Learning Qualified and Distributed Bounding Boxes for Dense Object Detection. NeurIPS, 2020.
5. Carion N. et al. End-to-End Object Detection with Transformers (DETR). ECCV, 2020.
6. Ultralytics. YOLOv8 Documentation. 2023. URL: <https://docs.ultralytics.com/>
7. Ren S. et al. Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks. NeurIPS, 2015.

ГІБРИДНІ ПІДХОДИ ДЛЯ МОДЕЛЮВАННЯ ТА ПРОГНОЗУВАННЯ ФІНАНСОВИХ ДАНИХ

Харитонова С.В.¹, Гуськова В.Г.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ kharytonova.svitlana@lkl.kpi.ua, ² guskovavera2009@gmail.com

Робота присвячена розробці та дослідженню гібридних моделей машинного навчання для прогнозування напрямку руху цін акцій. Запропоновано комбінацію статистичних методів фільтрації часових рядів (ARIMA, вейвлет-декомпозиція) із сучасними архітектурами глибокого навчання (LSTM, CNN, Transformer). Експериментально досліджено ефективність декількох гібридних архітектур на даних 100 компаній індексу S&P 500. Найкращий результат досягнуто завдяки застосуванню стратегії селективного входу на основі оцінки впевненості гібриду моделі PatchTST (з незалежною обробкою каналів) та CNN. Мета роботи – розробка та порівняльний аналіз гібридних архітектур для прогнозування напрямку руху цін акцій із застосуванням методів штучного інтелекту.

Ключові слова: прогнозування часових рядів, гібридні нейронні мережі, PatchTST, ARIMA, LSTM, фінансові ринки, глибоке навчання, трансформери.

1. ВСТУП

Прогнозування фінансових ринків є однією з найбільш практично значущих та водночас складних задач машинного навчання. Складність зумовлена високим рівнем шуму в даних, нестаціонарністю часових рядів та впливом екзогенних факторів, що важко піддаються формалізації. Класичні статистичні методи, такі як ARIMA та GARCH, ефективно моделюють лінійні залежності, проте не здатні вловлювати складні нелінійні патерни. Натомість нейронні мережі, зокрема LSTM та Transformer, демонструють високу здатність до апроксимації нелінійних функцій, але схильні до перенавчання на зашумлених фінансових даних.

Гіпотеза дослідження полягає в тому, що комбінування статистичних методів передоброби з архітектурами глибокого навчання дозволить використати переваги обох підходів: статистична фільтрація усуває лінійну компоненту та знижує шум, а нейронна мережа моделює залишкові нелінійні залежності.

2. МЕТОДОЛОГІЯ ТА АРХІТЕКТУРА МОДЕЛЕЙ

2.1. Базові моделі

ARIMA(p,d,q) – класична авторегресійна модель, в нашому випадку використовує 5 попередніх значень логарифмічних доходностей для прогнозування наступного. Параметр диференціювання $d=0$, оскільки ряд доходностей вже є стаціонарним. Модель навчається на тренувальній вибірці, а для тестування застосовується метод `apply()` з бібліотеки `statsmodels`, що дозволяє виконувати прогнозування на один крок вперед без перенавчання.

LSTM з DirectionalMSE – рекурентна нейронна мережа з двома шарами Long Short-Term Memory по 256 нейронів кожен. Ключовою особливістю є розроблена модифікована функція втрат DirectionalMSE, що складається з трьох компонент:

$$L = \text{MSE}(y, \hat{y}) + \lambda_1 \cdot \mathbb{I}[\text{sign}(y) \neq \text{sign}(\hat{y})] + \lambda_2 \cdot |P(\hat{y} > 0) - P(y > 0)|$$

де перший доданок – стандартна середньоквадратична помилка; другий – штраф за неправильний знак прогнозу (вагу $\lambda_1 = 5$ підбрано емпірично); третій – штраф за дисбаланс між часткою позитивних прогнозів та фактичних позитивних значень.

2.2. Гібридні архітектури

Wavelet-LSTM реалізує двоетапну обробку: спочатку вейвлет-перетворення (Daubechies-4) розкладає ряд на трендову та шумову компоненти, потім LSTM-мережа моделює залежності у згладженому ряді. Архітектура включає 2 шари LSTM по 64 нейрони з dropout 0.2.

ARIMA-LSTM побудовано за принципом послідовного моделювання: ARIMA(5,0,1) фіксує лінійну автокореляцію, а LSTM з 2 шарами по 64 нейрони навчається прогнозувати залишки. Фінальний прогноз формується як сума:

$$\hat{r}_t = \hat{r}_t^{\text{ARIMA}} + \hat{\varepsilon}_t^{\text{LSTM}}.$$

CNN-LSTM використовує багатомасштабні згорткові фільтри для виявлення локальних патернів різної тривалості. Три паралельні гілки Conv1D з розмірами ядра 3, 5 та 7 днів (по 32 фільтри кожна) виявляють короткострокові, середньострокові та тижневі патерни відповідно. Після конкатенації та пулінгу ознаки передаються до LSTM для моделювання часових залежностей.

Hybrid PatchTST – найскладніша архітектура, що поєднує сучасний Transformer-підхід із статистичною передобробкою. Ключові компоненти:

- Патч-ембедінг – вхідна послідовність довжиною 60 днів розбивається на патчі по 12 днів з кроком 8, що дає 7 патчів. Кожен патч проектується у простір розмірності $d_{\text{model}}=64$;
- Channel Independence – кожна з 45 ознак обробляється незалежно через спільний Transformer-енкодер. Це забезпечує інваріантність до масштабу ознак та дозволяє моделі вивчати універсальні часові патерни. Формально, вхідний тензор [Batch, Length, Channels] трансформується у [Batch×Channels, Length, 1], обробляється Transformer, і результати агрегуються назад;
- Gated Fusion – механізм адаптивного об'єднання ознак від CNN та Transformer гілок: $h_{\text{fused}} = g \cdot h_{\text{CNN}} + (1 - g) \cdot h_{\text{TST}}$, $g = \sigma(W[h_{\text{CNN}}; h_{\text{TST}}])$, де g – вектор воріт, що навчається визначати оптимальну комбінацію для кожного прикладу;
- Confidence Head – окрема голова мережі, що оцінює впевненість моделі у прогнозі (значення від 0 до 1). Навчається передбачати ймовірність правильного прогнозу на основі внутрішнього представлення.

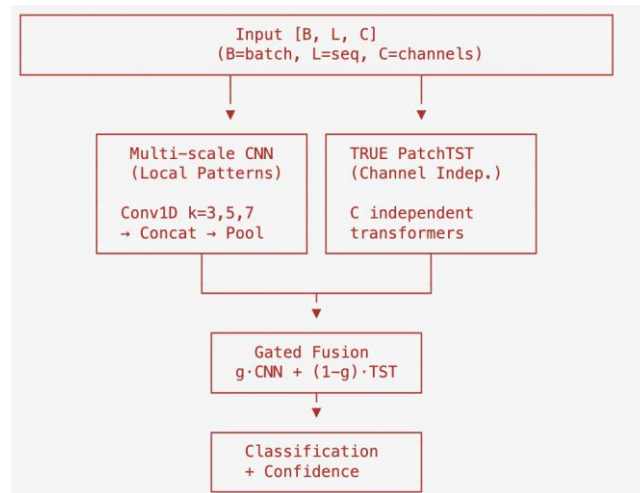


Рисунок 1. Архітектура гібридного методу CNN та PatchTST

3. ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ ТА АНАЛІЗ РЕЗУЛЬТАТІВ

Дослідження проведено на даних фондового ринку США, що включають історичні котирування (ціни відкриття, закриття, максимум, мінімум, обсяг торгів) 100 компаній індексу S&P 500 за період 2014–2024 рр. та фундаментальні показники: P/E ratio, P/B ratio, ROE, ROA, маржинальність та інші (всього 45 ознак).

Таблиця 1. Порівняння ефективності моделей

	Модель	DA
1	ARIMA(5,0,0)	49.6%
2	LSTM	51.2%
3	Wavelet + LSTM	53–54%
4	ARIMA + LSTM	53–54%
5	CNN + LSTM	56–57%
6	CNN + PatchTST	53–54%

Аналіз результатів, представлених у таблиці 1, демонструє чітку ієрархію ефективності досліджених підходів. Базова модель ARIMA показала точність 49.6%, що статистично еквівалентно випадковому вгадуванню та підтверджує обмеженість лінійних методів для прогнозування фінансових часових рядів. Застосування нейронної мережі LSTM забезпечило незначне покращення до 51.2%. Гібридні архітектури Wavelet+LSTM (рис. 2) та ARIMA+LSTM досягли точності 53–54%, що свідчить про ефективність попередньої фільтрації сигналу. Найвищу точність на окремому активі продемонструвала модель CNN+LSTM (56–57%, рис. 3), проте архітектура CNN+PatchTST (53–55%) забезпечує кращу узагальнювальну здатність при тестуванні на портфелі акцій.

Важливо зазначити, що перші моделі тестувалися на одному активі (ADBE), тоді як гібрид CNN та PatchTST було протестовано на 42 тікерах одночасно, що є значно складнішим завданням (табл. 2). Саме тому високий результат моделі CNN-LSTM (57.14%) є дещо оманливим і потребує критичної інтерпретації, оскільки нейромережі дуже легко «запам'ятати» характер руху однієї акції за короткий період і може свідчити про перенавчання, тоді як гібрид CNN та PatchTST обробляє 42 різні компанії з різною волатильністю однією нейромережею, що доводить її здатність до узагальнення.

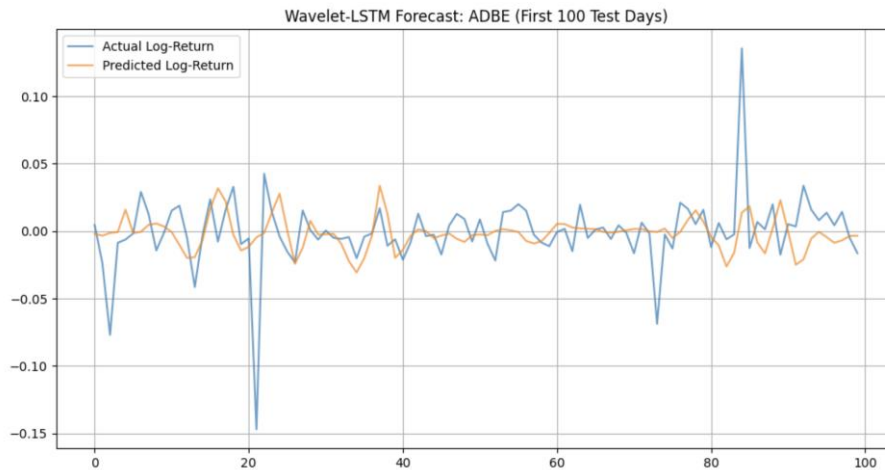


Рисунок 2. Порівняння фактичних та прогнозованих логарифмічних доходностей моделі Wavelet-LSTM (перші 100 днів тестової вибірки, тикер ADBE)

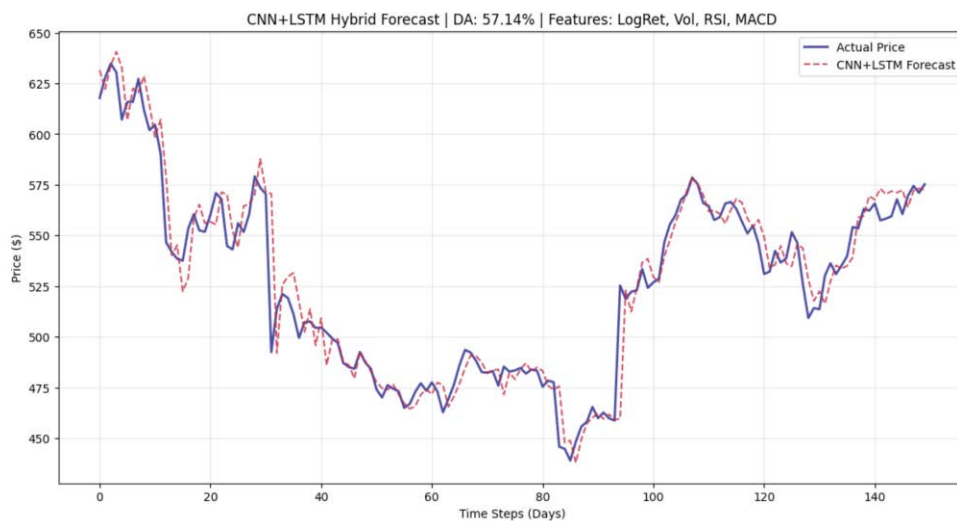


Рисунок 3. Реконструкція ціни акції ADBE на основі прогнозів моделі CNN-LSTM

Таблиця 2. Зведена таблиця результатів гібридних моделей

Характеристика	Wavelet-LSTM	CNN-LSTM	Hybrid CNN + PatchTST
Архітектура	Очищення шуму (Wavelet) + рекурентна мережа	Екстракція ознак (CNN) + моделювання часових залежностей (LSTM)	Паралельна (Локальні патерни + Глобальний тренд)
Охоплення даних	1 Тікер (ADBE)	1 Тікер (ADBE)	42 Тікери (Портфель)
Точність (Accuracy)	53.85%	57.14% (локальний максимум)	52.24% (Загальна) / 54.00% (Sniper)
Перевірка (Validation)	Random/Simple Split	Random/Simple Split	Temporal Split + Rolling Norm (Без витоків)

4. ВИСНОВОК

У дослідженні було обґрунтовано ефективність гібридних підходів на основі глибокого навчання для прогнозування напрямку руху цін акцій фондового ринку США.

Експериментально підтверджено, що класичні статистичні моделі (ARIMA) демонструють точність прогнозування напрямку на рівні випадкового вгадування, що узгоджується з гіпотезою слабкої форми ефективності ринку та обґрунтовує необхідність застосування більш складних підходів.

Гібридні архітектури, що поєднують статистичну передобробку часових рядів із методами глибокого навчання, стабільно перевершують як чисто статистичні, так і нейромережеві підходи. Моделі Wavelet-LSTM та CNN-LSTM досягли точності 53.85% та 57.14% відповідно на окремому активі (ADBE).

Реалізовано архітектуру Hybrid CNN+PatchTST, що поєднує багатомасштабні згорткові фільтри для виявлення локальних патернів із Transformer-енкодером для моделювання довгострокових залежностей. Застосування механізму незалежної обробки каналів забезпечує інваріантність до масштабу ознак та покращує узагальнювальну здатність моделі. Тестування запропонованого методу пройшло успішно, на портфелі з 42 акцій підтвердило здатність архітектури Hybrid CNN+PatchTST до узагальнення.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Box G.E.P., Jenkins G.M. Time Series Analysis: Forecasting and Control. – Holden-Day, 1970. – 575p. URL: https://www.researchgate.net/publication/299459188_Time_Series_Analysis_Forecasting_and_Control5th_Edition_by_George_E_P_Box_Gwilym_M_Jenkins_Gregory_C_Reinsel_and_Greta_M_Ljung_2015_Published_by_John_Wiley_and_Sons_Inc_Hoboken_New_Jersey_pp_712_ISBN_.
2. Nie Y., Nguyen N. H., Sinthong P., Kalagnanam J. A Time Series is Worth 64 Words: Long-term Forecasting with Transformers. International Conference on Learning Representations (ICLR). 2023. URL: <https://openreview.net/forum?id=Jbdc0vTOcol>.
3. Vaswani A., Shazeer N., Parmar N. et al. Attention Is All You Need. Advances in Neural Information Processing Systems (NeurIPS). 2017. Vol. 30. P. 5998–6008. URL: <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.
4. Fischer T., Krauss C. Deep learning with long short-term memory networks for financial market predictions. European Journal of Operational Research. 2018. Vol. 270, No. 2. P. 654–669. URL: <https://doi.org/10.1016/j.ejor.2017.11.054>.
5. Sezer O. B., Gudelek M. U., Ozbayoglu A. M. Financial time series forecasting with deep learning: A systematic literature review. Applied Soft Computing. 2020. Vol. 90. Article 106181. URL: <https://doi.org/10.1016/j.asoc.2020.106181>.
6. Zhang G. P. Time series forecasting using a hybrid ARIMA and neural network model. Neurocomputing. 2003. Vol. 50. P. 159–175. URL: [https://doi.org/10.1016/S0925-2312\(01\)00702-0](https://doi.org/10.1016/S0925-2312(01)00702-0).
7. Fama E. F. Efficient Capital Markets: A Review of Theory and Empirical Work. The Journal of Finance. 1970. Vol. 25, No. 2. P. 383–417. URL: <https://doi.org/10.2307/2325486>.
8. Yahoo Finance Historical Data [Електронний ресурс] – Режим доступу: <https://finance.yahoo.com/>.
9. LeCun Y., Bengio Y., Hinton G. Deep learning. Nature. 2015. Vol. 521, No. 7553. P. 436–444. URL: <https://doi.org/10.1038/nature14539>.

АДАПТИВНИЙ МЕТОД СТРУКТУРНОГО ПРУНІНГУ ДЛЯ ОПТИМІЗАЦІЇ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ

Швець В.О.¹, Шаповал Н.В.²

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського», Київ, Україна

¹ shvets.vitaliy@lkl.kpi.ua, ² shovgun@gmail.com [0000-0002-8509-6886]

Метою дослідження є підвищення ефективності розгортання великих мовних моделей архітектури Transformer на ресурсно-обмежених пристроях шляхом розробки адаптивного методу структурного прунінгу. Запропоновано та реалізовано метод Adaptive 2SSP (Two-Stage Structured Pruning) Reversed, який поєднує повне видалення блоків уваги на основі метрики косинусу подібності та адаптивне стиснення шарів MLP з урахуванням їх індивідуальної надмірності. Експериментальна перевірка на моделі Llama-3.2-3B демонструє зменшення споживання відеопам'яті на 35.1% (з 5.98 GB до 3.88 GB) та прискорення генерації токенів на 34.8% (з 92 до 124 TPS) при коефіцієнті прунінгу 0.4 та покращення середньої точності на бенчмарках в порівнянні з іншими методами. Новизна полягає у розробці механізму динамічного розподілу коефіцієнтів стиснення між шарами на основі метрики Block Influence та зворотному порядку оптимізації компонентів моделі в порівнянні з оригінальним 2SSP. Результати можуть використовуватись для оптимізації розгортання LLM на споживчому обладнанні з обмеженими ресурсами.

Ключові слова: великі мовні моделі, структурний прунінг, оптимізація, LLM, адаптивний прунінг.

1. ВСТУП

Стрімкий розвиток обробки природної мови, зумовлений появою великих мовних моделей (Large Language Models, LLM) таких як GPT, LLaMA, Mistral тощо, відкрив нові можливості у генерації тексту, підтримці прийняття рішень та автоматизації інтелектуальних задач [1].

Проте моделі з мільярдами параметрів вимагають колосальних обчислювальних ресурсів. Навчання GPT-3 (175 млрд параметрів) коштувало близько 4.6 млн доларів та спричинило викиди CO₂ еквівалентні 552 тоннам [2]. Основною перешкодою для широкого впровадження LLM є висока обчислювальна складність інференсу та значні вимоги до відеопам'яті (VRAM), що ускладнює їх використання на локальних користувацьких пристроях.

Існуючі методи компресії мають свої обмеження. Такі методи, як SparseGPT [3] та Wanda [4] використовують неструктурований прунінг, який заміняє надлишкові ваги на 0, що насправді не зменшує розмір моделі та вимагає спеціалізованого апаратного забезпечення для прискорення роботи моделі. Такі методи слугують радше прикладами максимально можливого збереження ефективності моделей при прунінгу, ніж ефективними методами для зменшення використовуваних ресурсів. Метод SliceGPT [5], його адаптивна модифікація Dynamic Slicing [6] та 2SSP [7] виконують видалення рядків і колонок у матрицях ваг або цілих

блоків мережі, що призводить до фактичного зменшення моделі та її прискорення. Однак більшість підходів використовують фіксовані коефіцієнти стиснення або не враховують повною мірою різну важливість окремих блоків.

Дослідження показують, що глибші шари трансформерних моделей часто мають вищу надмірність (layer redundancy) і можуть бути більш агресивно оптимізовані без суттєвої втрати якості [6, 8]. Це створює підґрунтя для розробки адаптивних методів стиснення, які враховують індивідуальні характеристики кожного шару.

У даній роботі пропонується вдосконалений підхід Adaptive 2SSP Reversed, який базується на гіпотезі про надмірність глибоких шарів та видаляє параметри з урахуванням специфіки компонентів моделі.

2. ОПИС ЗАПРОПОНОВАНОГО МЕТОДУ

2.1. Метрика оцінки надмірності (Block Influence)

Для кількісної оцінки важливості кожного блоку введено метрику Block Influence, що базується на косинусі подібності між вхідними та вихідними станами блоку. Оскільки в архітектурі Transformer вихід блоку Y формується як сума входу X та перетворення $F(X)$, тобто $Y = X + F(X)$, малий вплив функції $F(X)$ свідчить про низьку інформативність шару. Надмірність $R(L_i)$ для i -го шару обчислюється як косинусна подібність між вхідним та вихідним тензорами на калібрувальній вибірці:

$$R(L_i) = \frac{1}{N} \sum_{j=1}^N \text{cossim}(x_j, y_j) = \frac{1}{N} \sum_{j=1}^N \frac{x_j \cdot y_j}{\|x_j\| \|y_j\|}$$

де N – кількість зразків у вибірці для калібрування, x_j – вхідний вектор прихованих станів, y_j – вихідний вектор після додавання residual connection.

Значення $R(L_i) \rightarrow 1$ вказує на те, що шар майже не змінює векторне представлення токенів, а отже, є кандидатом на видалення або агресивне стиснення.

2.2. Алгоритм Adaptive 2SSP Reversed

Запропонований метод Adaptive Two-Stage Structured Pruning Reversed відрізняється від стандартного підходу 2SSP зворотним порядком оптимізації (що засновано на логіці: видаляємо повністю блоки уваги, які мають набагато меншу кількість параметрів за MLP блоки і при цьому модель втрачає невелику частину своєї «якості» навіть при прунінгу 30% таких шарів [9]) та адаптивністю розподілу коефіцієнтів стиснення. Алгоритм складається з чотирьох послідовних етапів:

Етап 1. Прунінг блоків уваги (Attention Pruning)

На першому етапі аналізується косинусна подібність між вхідними та вихідними станами блоків уваги на калібрувальній вибірці (датасет C4, 32 зразки по 2048 токенів). Блоки, що вносять найменший вклад у перетворення сигналу (висока подібність), вилучаються повністю, що зменшує загальну глибину мережі.

Етап 2. Обчислення метрики Block Influence

Для кожного MLP шару розраховується метрика надмірності $R(L_i)$, яка нормалізується в діапазон $[0, 1]$ для коректного порівняння та використання в адаптивному механізмі розподілу.

Етап 3. Адаптивний розподіл коефіцієнтів прунінгу

Замість рівномірного видалення нейронів, алгоритм розраховує індивідуальний коефіцієнт збереження для кожного шару за формулою:

$$Keep_ratio(L_i) = 1 - [R(L_i) \times (target_rate) / mean(R)],$$

де $target_rate$ – цільовий середній коефіцієнт прунінгу, $mean(R)$ – середня надмірність по всіх шарах. Шари з меншою інформативністю піддаються агресивнішому прунінгу, тоді як критично важливі шари зберігаються.

Етап 4. Структурований прунінг MLP нейронів

Для шарів Feed-Forward Networks оцінка важливості нейронів базується на L_2 -норми проміжних активацій у MLP блоці:

$$s_j = \frac{1}{N} \sum_{j=1}^N \|z_c^j\|_2,$$

де N – кількість зразків у вибірці для калібрування, z_c^j – активація j -го нейрона для s -ї послідовності в калібрувальному наборі даних.

Нейрони з найнижчими значеннями важливості вилучаються відповідно до розрахованих індивідуальних коефіцієнтів прунінгу для кожного шару. Видалення відбувається шляхом вирізання відповідних рядків та стовпців у матрицях ваг, що забезпечує реальне зменшення розміру моделі.

3. РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

3.1. Метрики якості генерації

Для комплексної оцінки якості було використано набір різнопланових задач, що перевіряють різні аспекти "інтелекту" моделі:

1. WikiText-2 (PPL) – оцінює здатність моделі передбачати наступний токен. Менше значення означає, що модель більш впевнена у своїх прогнозах і має високе значення прогнозу для одного токена, а не багато схожих ймовірностей для багатьох токенів.

2. ARC-Easy (AI2 Reasoning Challenge) – це підмножина бенчмарку ARC, яка містить завдання початкового рівня зі шкільної науки. Вона перевіряє базові наукові знання, елементарний здоровий глузд та здатність робити прості логічні висновки.

3. BoolQ (Boolean Questions) – датасет з відповідями «так/ні» на основі фрагментів тексту. Він перевіряє розуміння контексту, здатність робити логічні висновки, інтерпретацію тверджень і контекстних протиріч.

4. HellaSwag – вибір логічного завершення ситуації, де неправильні варіанти спеціально згенеровані, щоб бути "схожими" на правильні. Він перевіряє здатність моделі "читати між рядків" та здоровий глузд.

Порівняння запропонованого методу проводилося на моделі Llama-3.2-3B з методами GLU Aware Pruning (простий метод видалення рядків та колонок матриць вагів, заснований на магнітудах ваг [10]), Dynamic Slicing та стандартною версією 2SSP.

Аналіз результатів (табл. 1) показує, що запропонований метод перевершує конкурентні, особливо при високих ступенях стиснення (30-40%). Це підтверджує ефективність збереження нейронів у шарах, що є критичними для завдань логічного висновку (BoolQ, ARC), за рахунок видалення надлишковості в інших частинах мережі.

Варто також зауважити, що хоч Dynamic Slicing показує кращі значення Perplexity, при відсотках прунінгу від 0.1 до 0.3, отримані результати по бенчмаркам є значно нижчими і показують, що моделі не вдалось зберегти здатність логічного мислення та таку ж кількість знань, як при використанні запропонованого алгоритму.

Таблиця 1. Метрики якості

Відсоток прунінгу	Модель / Метод	WikiText2 (PPL)	ARC-Easy (Acc, %)	BoolQ (Acc, %)	HellaSwag (Acc, %)	AVG Acc (без PPL)
0%	Llama-3.2-3B (Baseline)	7,81	74,0%	74,0%	48,5%	65,50%
10%	GLU Aware Pruning (magnitude)	11,12	66,5%	48,9%	44,8%	53,40%
	Dynamic slicing	9,51	63,4%	45,9%	42,8%	50,70%
	2SSP	9,64	68,3%	65,6%	47,0%	60,30%
	Adaptive 2SSP	9,58	67,8%	68,6%	47,0%	61,13%
20%	GLU Aware Pruning (magnitude)	17,15	52,7%	41,4%	41,7%	45,27%
	Dynamic slicing	12,24	52,1%	38,1%	37,8%	42,67%
	2SSP	13,46	60,6%	59,5%	42,5%	54,20%
	Adaptive 2SSP	12,81	62,1%	61,3%	45,1%	56,17%
30%	GLU Aware Pruning (magnitude)	34,90	47,4%	40,9%	35,9%	41,40%
	Dynamic slicing	17,69	41,3%	37,5%	35,0%	37,93%
	2SSP	19,35	52,4%	56,5%	38,1%	49,00%
	Adaptive 2SSP	18,00	54,4%	60,7%	42,5%	52,53%
40%	GLU Aware Pruning (magnitude)	73,67	36,3%	52,3%	33,1%	40,57%
	Dynamic slicing	26,75	35,3%	37,5%	33,2%	35,33%
	2SSP	28,38	43,3%	61,9%	35,9%	47,03%
	Adaptive 2SSP	26,35	45,2%	63,5%	37,5%	48,73%

3.2. Ефективність компресії та швидкодії

Експериментальна перевірка проводилася на моделі Llama-3.2-3B (табл. 2). Оцінювалися споживання відеопам'яті (VRAM) та швидкість генерації (TPS – tokens per second) для різних значень коефіцієнта прунінгу. У якості базового методу порівняння використано оригінальну модель (Dense) та підходи Dynamic Slicing.

Таблиця 2. Порівняння ефективності компресії та швидкодії

Відсоток прунінгу	Споживання VRAM (GB)	Зменшення пам'яті	TPS	Приріст швидкості
0%	5,98	-	92	-
10%	5,45	8.9%	99	7.6%
20%	4,93	17.6%	108	17.4%
30%	4,40	26.4%	118	28.2%
40%	3,88	35.1%	124	34.8%

Метод дозволяє досягти значного зростання швидкодії зі збільшенням ступеня стиснення. При коефіцієнті 0.4 споживання пам'яті знижується до рівня <4 GB, що є критичним порогом для багатьох мобільних GPU, а швидкість генерації зростає на ~ 35%.

4. ВИСНОВКИ

У роботі розроблено та досліджено метод адаптивного структурного прунінгу Adaptive 2SSP Reversed. Наукова новизна полягає у ефективному поєднанні механізму динамічного розподілу коефіцієнтів стиснення між шарами на основі метрики Block Influence, що враховує індивідуальну надмірність компонентів моделі, видалення повних блоків механізмів уваги та видаленні нейронів MLP блоків з мережі.

Експериментальні результати на моделі Llama-3.2-3B демонструють ефективність запропонованого підходу в порівнянні з іншими методами структурного прунінгу. Також показане фактичне прискорення інференсу моделі (до 34.8% при 40% прунінгу) та зменшення її розміру (до 35.1% при 40% прунінгу), що зменшує потрібні обчислювальні ресурси для її запуску.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. A. Vaswani et al., «Attention Is All You Need», arXiv.org, 2017. <https://arxiv.org/abs/1706.03762> (дата звернення: 21.11.2025)
2. D. A. Patterson et al., «Carbon Emissions and Large Neural Network Training», arXiv, Apr. 2021, doi: <https://doi.org/10.48550/arxiv.2104.10350> (дата звернення: 21.11.2025)
3. E. Frantar and D. Alistarh, «SparseGPT: Massive Language Models Can Be Accurately Pruned in One-Shot», Jan. 2023, doi: <https://doi.org/10.48550/arxiv.2301.00774>. (дата звернення: 21.11.2025)
4. M. Sun, Z. Liu, A. Bair, and K. J. Zico, «A Simple and Effective Pruning Approach for Large Language Models», arXiv (Cornell University), Jun. 2023, doi: <https://doi.org/10.48550/arxiv.2306.11695>. (дата звернення: 21.11.2025)
5. S. Ashkboos, M. L. Croci, M. Gennari, T. Hoefler, and J. Hensman, «SliceGPT: Compress Large Language Models by Deleting Rows and Columns», arXiv (Cornell University), Jan. 2024, doi: <https://doi.org/10.48550/arxiv.2401.15024>. (дата звернення: 21.11.2025)
6. R.-G. Dumitru, P.-I. Clotan, V. Yadav, D. Peteleaza, and M. Surdeanu, «Change Is the Only Constant: Dynamic LLM Slicing based on Layer Redundancy», arXiv (Cornell University), Nov. 2024, doi: <https://doi.org/10.48550/arxiv.2411.03513>. (дата звернення: 21.11.2025)
7. F. Sandri, E. Cunegatti, and G. Iacca, «2SSP: A Two-Stage Framework for Structured Pruning of LLMs», arXiv (Cornell University), Jan. 2025, doi: <https://doi.org/10.48550/arxiv.2501.17771>. (дата звернення: 21.11.2025)
8. A. Gromov, K. Tirumala, H. Shapourian, P. Glorioso, and D. A. Roberts, «The Unreasonable Ineffectiveness of the Deeper Layers», arXiv.org, Mar. 26, 2024. <https://arxiv.org/abs/2403.17887> (дата звернення: 21.11.2025)
9. S. He, G. Sun, Z. Shen, and A. Li, «What Matters in Transformers? Not All Attention is Needed», arXiv (Cornell University), Jun. 2024, doi: <https://doi.org/10.48550/arxiv.2406.15786>. (дата звернення: 21.11.2025)
10. «Making LLMs Smaller Without Breaking Them: A GLU-Aware Pruning Approach» Huggingface.co, 2024. <https://huggingface.co/blog/oopere/making-llms-smaller-without-breaking-them> (дата звернення: 21.11.2025)