

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
Західноукраїнський національний університет
Факультет комп'ютерних інформаційних технологій
Кафедра інформаційно-обчислювальних систем і управління

КОНДРАТЮКУ Георгію Георгійовичу

**Метод прогнозування фінансових труднощів компаній на
основі глибокого навчання із використанням семантичних
ознак та реакцій користувачів/A Method for Predicting
Corporate Financial Distress Based on Deep Learning Using
Semantic Features and User Reactions**

спеціальність: 122 - Комп'ютерні науки
освітньо-професійна програма - Комп'ютерні науки

Кваліфікаційна робота

Виконав студент групи КНм-21
Кондратюк Г.Г

Науковий керівник:
к.т.н., доцент, Ліп'яніна-
Гончаренко Х.В.

Кваліфікаційну роботу
допущено до захисту:
«___» _____ 20___ р.
В.о. завідувача кафедри
_____ Н.М. Васильків

ТЕРНОПІЛЬ - 2024

Факультет комп'ютерних інформаційних технологій
Кафедра інформаційно-обчислювальних систем і управління
Освітній ступінь «магістр»
спеціальність: 122 – Комп'ютерні науки
освітньо-професійна програма – Комп'ютерні науки

ЗАТВЕРДЖУЮ
Завідувач кафедри
М.П. Комар
« ____ » _____ 20__ р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ
КОНДРАТЮК Георгій Георгійович**

(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи

Метод прогнозування фінансових труднощів компаній на основі глибокого навчання із використанням семантичних ознак та реакцій користувачів/A
Method for Predicting Corporate Financial Distress Based on Deep Learning Using
Semantic Features and User Reactions

керівник роботи к.т.н., доцент, Ліп'яніна-Гончаренко Х.В.

затверджені наказом по університету від 12 грудня 2023 року № 753.

2. Строк подання студентом закінченої кваліфікаційної роботи 4 грудня 2024 р.

3. Вихідні дані до кваліфікаційної роботи: завдання на кваліфікаційну роботу студента, наукові статті, технічна література.

4. Основні питання, які потрібно розробити

- Провести аналіз сучасних методів прогнозування фінансових труднощів на основі поточних звітів і реакцій користувачів.
- Огляд та оцінка існуючих рішень для прогнозування фінансових труднощів.
- Спроекувати архітектуру методу URGDAN для прогнозування фінансових труднощів.
- Використати FinBERT та BiGRU для моделювання послідовних даних із поточних звітів.
- Оцінити реакції користувачів як показник важливості поточних звітів.
- Розробити адаптивне коригування семантичних ознак на основі реакцій користувачів.
- Моделювати фінансові ознаки компаній за допомогою двошарового MLP.
- Провести огляд даних та опис характеристик обраного набору даних.
- Оцінити ефективність методу URGDAN у порівнянні з іншими підходами прогнозування..

5. Перелік графічного матеріалу у роботі

- Загальна архітектура URGDAN;
- Модуль уваги, керований реакціями користувачів.

6. Консультанти розділів кваліфікаційної роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		Завдання видав	Завдання прийняв

7. Дата видачі завдання 12 грудня 2023 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів кваліфікаційної роботи	Примітка
1	Затвердження теми кваліфікаційної роботи, ознайомлення з літературними джерелами та складання плану роботи.	до 01.01. 2024 р.	
2	Написання 1 розділу кваліфікаційної роботи	до 01.03. 2024 р.	
3	Написання 2 розділу кваліфікаційної роботи	до 20.05.2024 р.	
4	Написання 3 розділу кваліфікаційної роботи	до 28.10. 2024 р.	
5	Представлення попереднього варіанту кваліфікаційної роботи, перевірка та внесення змін керівником	до 11.11.2024 р.	
6	Опрацювання зауважень та представлення завершеного варіанту кваліфікаційної роботи. Підготовка супроводжуючих документів.	до 25.11.2024 р.	
7	Перевірка кваліфікаційної роботи на оригінальність тексту.	до 30.11.2024 р.	
8	Оформлення кваліфікаційної роботи та отримання допуску до захисту	до 04.12.2024 р.	
9	Подання кваліфікаційної роботи до захисту на засіданні атестаційної комісії.	до 14.12. 2024 р.	

Студент _____ Кондратюк Г.Г.
підпис

Керівник роботи _____ к.т.н., доцент, Ліп'яніна-Гончаренко Х.В.
підпис

РЕЗЮМЕ

Кваліфікаційна робота на тему «Метод прогнозування фінансових труднощів компаній на основі глибокого навчання із використанням семантичних ознак та реакцій користувачів» на здобуття освітнього ступеня «Магістр» зі спеціальності 122 «Комп'ютерні науки» освітньої програми «Комп'ютерні науки» написана обсягом в 74 сторінок і містить 5 ілюстрацій, 9 таблиць та 48 використаних джерела.

Метою даної кваліфікаційної роботи є розробка методу прогнозування фінансових труднощів компаній, який інтегрує семантичні ознаки текстів поточних звітів та реакції користувачів для підвищення точності прогнозів.

Методи дослідження: аналіз наукової літератури, методи глибокого навчання для обробки текстових даних (FinBERT, BiGRU), підходи до аналізу реакцій користувачів, багаточаровий перцептрон (MLP), порівняльний аналіз моделей прогнозування.

Результати дослідження: розроблено модель URGDAN, що поєднує семантичні ознаки текстових даних із реакціями користувачів через механізм уваги. Метод покращує точність прогнозування фінансових труднощів, особливо в умовах низької прозорості інформації. Запропонована модель показала ефективність у порівнянні з традиційними підходами.

Результати роботи можуть бути застосовані у системах моніторингу фінансових ризиків, інвесторами та кредиторами для покращення процесів прийняття рішень.

Ключові слова: ФІНАНСОВІ ТРУДНОЩІ, ГЛИБОКЕ НАВЧАННЯ, FINBERT, BIGRU, УВАГА, РЕАКЦІЇ КОРИСТУВАЧІВ.

ABSTRACT

Qualification work on the topic «A Method for Predicting Corporate Financial Distress Based on Deep Learning Using Semantic Features and User Reactions» for a Master's degree in speciality 122 «Computer Science» educational and professional program «Computer Science» is written on 74 pages and contains 5 figures, 9 tables and 48 sources.

The purpose of this qualification work is to develop a method for predicting corporate financial distress by integrating semantic features of current report texts and user reactions to improve prediction accuracy.

Research methods: analysis of scientific literature, deep learning methods for text processing (FinBERT, BiGRU), user reaction analysis approaches, multilayer perceptron (MLP), and comparative model analysis.

Research results: a URGDAN model was developed, combining semantic features of textual data with user reactions through an attention mechanism. The method improves financial distress prediction accuracy, particularly under conditions of low transparency. The proposed model demonstrated effectiveness compared to traditional approaches.

The results can be applied in financial risk monitoring systems, by investors and creditors, to enhance decision-making processes.

Keywords: FINANCIAL DISTRESS, DEEP LEARNING, FINBERT, BIGRU, ATTENTION, USER REACTIONS.

ЗМІСТ

Вступ	8
1. Сучасний стан предметної області і постановка задачі дослідження	12
1.1 Сучасні методи прогнозування фінансових труднощів на основі поточних звітів і реакцій користувачів	12
1.2 Огляд існуючих рішень	16
1.3 Постановка задачі	20
Висновки до розділу 1	23
2 Архітектурні рішення та методи прогнозування фінансових труднощів на основі глибокого навчання	25
2.1 Проектування архітектури методу URGDAN для прогнозування фінансових труднощів	25
2.2 Моделювання послідовних даних з використанням FinBERT та BiGRU для обробки поточних звітів	27
2.3 Оцінка реакцій користувачів як показник важливості поточних звітів	30
2.4 Адаптивне коригування семантичних ознак на основі реакцій користувачів	32
2.5 Моделювання фінансових ознак компаній за допомогою двошарового MLP	34
Висновки до розділу 2	36
3 . Практичне використання отриманих наукових результатів	37
3.1 Огляд даних та характеристика обраного набору даних	37
3.2 Оцінка ефективності різних підходів до прогнозування фінансових ризиків	40

3.3 Аналіз ефективності URGDAN у порівнянні з еталонними методами прогнозування	43
3.4 Аналіз впливу реакцій користувачів на ефективність прогнозування фінансових труднощів	46
3.5 Кейс-дослідження	48
Висновки до розділу 3	51
Висновки	53
Список використаних джерел	55
Додаток А Апробація отриманих результатів	62

ВСТУП

Актуальність дослідження визначається зростаючою потребою у своєчасному виявленні фінансових труднощів компаній, що мають безпосередній вплив на стабільність ринкової економіки, фінансові втрати кредиторів та інвесторів. У сучасних умовах швидкої змінності ринкових тенденцій та підвищених фінансових ризиків, особливо в умовах кризи, традиційні методи прогнозування фінансових труднощів, засновані виключно на фінансових показниках та періодичних звітах, часто не встигають враховувати динамічні зміни в діяльності компаній. Поточні звіти компаній, що відображають ключові події у реальному часі, разом з реакціями користувачів на ці події, є важливим додатковим джерелом інформації для прогнозування фінансових ризиків.

Інтеграція сучасних технологій глибокого навчання, таких як моделі BiGRU та FinBERT, із механізмами обробки текстової інформації та реакцій користувачів дозволяє розширити можливості прогнозування фінансових труднощів і робить цей процес більш гнучким і точним. Використання семантичних ознак з поточних звітів у поєднанні з активними та пасивними реакціями користувачів дає змогу отримати повніше уявлення про реальний фінансовий стан компанії. Це дозволяє створювати більш адаптивні та точні моделі прогнозування, що особливо актуально в умовах зростання кількості непередбачуваних подій, таких як глобальні економічні кризи та внутрішні корпоративні ризики.

Таким чином, дослідження, спрямоване на розробку та впровадження нової моделі прогнозування фінансових труднощів компаній на основі глибокого навчання із використанням семантичних ознак та реакцій користувачів, є надзвичайно актуальним. Воно відповідає потребам сучасного бізнесу та фінансових установ у створенні більш ефективних механізмів для

управління фінансовими ризиками і підвищення стійкості компаній у нестабільних умовах ринку.

Метою даної роботи є розробка та впровадження методу прогнозування фінансових труднощів компаній на основі глибокого навчання, що інтегрує семантичні ознаки поточних звітів та реакції користувачів для підвищення точності прогнозів.

Задачі, які необхідно вирішити для досягнення поставленої мети, включають:

1. Провести аналіз сучасних методів прогнозування фінансових труднощів на основі поточних звітів і реакцій користувачів.
2. Огляд та оцінка існуючих рішень для прогнозування фінансових труднощів.
3. Спроектувати архітектуру методу URGDAN для прогнозування фінансових труднощів.
4. Використати FinBERT та BiGRU для моделювання послідовних даних із поточних звітів.
5. Оцінити реакції користувачів як показник важливості поточних звітів.
6. Розробити адаптивне коригування семантичних ознак на основі реакцій користувачів.
7. Моделювати фінансові ознаки компаній за допомогою двошарового MLP.
8. Провести огляд даних та опис характеристик обраного набору даних.
9. Оцінити ефективність методу URGDAN у порівнянні з іншими підходами прогнозування.

Об'єктом дослідження є процеси прогнозування фінансових труднощів компаній в умовах ринкової економіки.

Предметом дослідження є модель прогнозування фінансових труднощів, що базується на глибокому навчанні, використанні семантичних ознак поточних звітів та поведінкових реакцій користувачів.

Методи дослідження у даній роботі застосовано комплекс методів дослідження, що включає аналіз наукової літератури та існуючих рішень, методи глибокого навчання для обробки текстових даних, такі як FinBERT і BiGRU, а також підходи до аналізу поведінкових даних користувачів. Для моделювання фінансових показників використовувалася архітектура багат шарового перцептрона (MLP). Крім того, були використані методи порівняльного аналізу для оцінки ефективності моделі URGDAN у порівнянні з іншими підходами. Для тестування гіпотез було проведено експерименти з використанням перехресної валідації, а для статистичної оцінки результатів застосовувався тест Фрідмана.

Наукова новизна отриманих результатів полягає в розробці метода прогнозування фінансових труднощів компаній на основі глибокого навчання із застосуванням семантичних ознак поточних звітів та реакцій користувачів. Запропонована архітектура URGDAN використовує механізм уваги, керований реакціями користувачів, що дозволяє інтегрувати як текстову інформацію з поточних звітів, так і поведінкові дані користувачів, забезпечуючи більш точне прогнозування фінансових труднощів. Такий підхід вперше об'єднує реакції користувачів як додатковий індикатор важливості подій у фінансових прогнозах.

Практичне значення одержаних результатів полягає у можливості використання розробленої моделі URGDAN для точнішого прогнозування фінансових труднощів компаній, що дозволить інвесторам, кредиторам та аналітикам ефективніше оцінювати ризики і приймати обґрунтовані рішення щодо інвестицій або надання фінансування. Запропонований метод може бути впроваджений у системи моніторингу фінансових ризиків для покращення процесів виявлення фінансових проблем на ранніх етапах, що сприяє зниженню ймовірності фінансових втрат і поліпшенню управління активами.

Структура та обсяг роботи. Робота складається із вступу, трьох розділів, висновків, списку використаних джерел та додатків.

Апробація результатів дослідження. Основні теоретичні положення роботи й практичні результати дослідження доповідалися й обговорювалися на V Міжнародної наукової конференції «Стратегічні напрямки розвитку науки: фактори впливу та взаємодії», яка відбулася 11 жовтня 2024 року у місті Луцьк, Україна та III Міжнародної наукової конференції «Інтелектуальний ресурс сьогодення: наукові задачі, розвиток та запитання», яка відбулася 20 вересня 2024 року у місті Одеса, Україна.

1 СУЧАСНИЙ СТАН ПРЕДМЕТНОЇ ОБЛАСТІ І ПОСТАНОВКА ЗАДАЧІ ДОСЛІДЖЕННЯ

1.1 Сучасні методи прогнозування фінансових труднощів на основі поточних звітів і реакцій користувачів

Фінансова неспроможність стосується ситуації, коли компанія стикається з кризою у внутрішньому або зовнішньому середовищі, що призводить до дефіциту грошових потоків, неплатоспроможності, зниження рентабельності або навіть банкрутства [1]. Такі фінансові труднощі спричиняють значні втрати для кредиторів і інвесторів, а також можуть негативно вплинути на стабільність і стійкий розвиток економіки загалом [2]. Прогнозування фінансових проблем (FDP) дозволяє своєчасно і ефективно виявляти потенційні фінансові ризики в найближчій перспективі, забезпечуючи ранні попереджувальні сигнали для прийняття обґрунтованих рішень [3].

Сигнали для прогнозування фінансових труднощів можуть походити як з фінансових показників, так і з текстових даних (наприклад, [4-6]). Фінансові індикатори відображають результати діяльності та фінансовий стан компанії за попередні періоди, які регулярно застосовуються у системах фінансової звітності. Водночас, текстові документи, що розкривають інформацію, є важливим доповненням до фінансових показників і представляють собою якісні дані. Основні текстові документи включають періодичні звіти (річні та піврічні звіти) і поточні звіти (також відомі як форми 8-K). Текстовий зміст періодичних звітів містить оцінки результатів діяльності компанії та її потенціал на майбутнє від керівництва, однак такі звіти мають певну затримку в оприлюдненні інформації. Вони також не враховують негайний вплив важливих подій, які трапляються нерегулярно у бізнес-діяльності.

Поточні звіти – це документи, що відображають ключові події, такі як відставки керівництва, коливання акцій або значні угоди з пов'язаними сторонами, що є важливими для інвесторів і громадськості [7]. Ці документи є

важливою частиною постійних зобов'язань компанії з розкриття інформації в межах операційної діяльності та управління. Компанії намагаються охопити ширший спектр інформації через поточні звіти [8,9]. Події, що відображаються у поточних звітах, можуть суттєво впливати на фінансові показники корпорацій [10], а інвестори приділяють особливу увагу змісту таких звітів, що викликає відповідну реакцію ринку [11]. Хоча інформаційна цінність поточних звітів була підтверджена в існуючих дослідженнях, до цього часу проведено недостатньо досліджень щодо їх включення у прогнозування фінансових проблем (FDP). У цьому дослідженні ми зосереджуємо увагу на використанні поточних звітів для підвищення ефективності FDP.

У той час як поточні звіти можуть надавати додаткову інформацію для FDP, використання поточних звітів для прогнозування фінансових труднощів пов'язане з унікальними проблемами, з яких ми виділяємо дві критичні. По-перше, семантику подій, розкритих у поточних звітах, важко вловити. Крім того, події, що розкриваються в різних поточних звітах, є хронологічними, а події, що розкриваються послідовно, мають логічні кореляції, що має значний вплив на прогнозування фінансових труднощів. По-друге, поточні звіти містять різні типи подій, а корисність (яку в цьому дослідженні називають «важливістю») кожного типу подій у прогнозуванні фінансових труднощів є неоднорідною. Як виміряти важливість подій у кожному поточному звіті складно. Безпосередньо з'ясувати важливість цих подій за текстовим контентом може бути нескладно.

Відповіді користувачів на поточні звіти, такі як читання, вподобання та коментарі, відображають, чи містить контент поточного звіту інформацію, пов'язану з інвестиційною вартістю або бізнес-ситуацією компанії [12]. Деякі важливі події можуть стати причиною того, що велика кількість експертів, інвесторів, кредиторів та інших користувачів збереться вперше. Вони активно дають відповіді, висловлюють свою думку з приводу подій, вказують на можливі ризики для компанії [13]. Велика кількість відповідей користувачів,

таких як лайки, може свідчити про схвалення контенту або важливість публікації [14]. Відповіді користувачів можуть відображати увагу користувачів до текстового контенту і навіть впливати на бренди або компанії [15]. Як показано на прикладі (рисунок 1.1), користувачі можуть приділяти більше уваги поточним звітам, що містять важливу інформацію про компанію, наприклад, про придбання компанії та завершення змін у промисловій та комерційній реєстрації. Інвестори вважають, що завдяки цьому придбанню компанія може інтегрувати ці передові технології у свої існуючі продукти та послуги, значно підвищивши свою комерційну цінність. Користувачі швидко високо оцінили звіт, що призвело до жвавої дискусії в розділі коментарів.

Readings	Likes	Comments	Current Reports	Company	Date
12680	24	37	Announcement on Proposed No Profit Distribution in 2016	XXX	2017/4/27
5682	11	6	Announcement on Resolutions of the Seventh Board Meeting	XXX	2017/5/14
6731	15	5	Announcement of Sporadic Related Transactions	XXX	2017/6/12
26841	100	164	Announcement on the Acquisition of A Company and the Completion of the Industrial and Commercial Registration Changes	XXX	2017/6/25

Рисунок 1.1 - Поточні звіти компанії та відповіді користувачів.

Попередні дослідження з виявлення шахрайства (наприклад, [16]), інтелектуальний аналіз тексту (наприклад, [5]), та глибоке навчання (наприклад, [17]) надали різні методи включення текстової інформації для FDP. У той час як більшість існуючих досліджень були зосереджені на періодичних звітах, поточним звітам приділялося набагато менше уваги. Більше того, існуючі методи в основному витягували текстові ознаки (наприклад, тематичні, семантичні та семантичні) лише на основі змісту звітів про розкриття інформації, що може бути недостатнім для відображення важливості поточного звіту. Дуже бажаним є метод, який міг би ефективно визначати важливі поточні звіти і, відповідно, точно прогнозувати фінансові труднощі.

У цьому дослідженні досліджується новий підхід до використання поточних звітів для прогнозування фінансових труднощів (FDP) шляхом

запропонування нової моделі глибокого навчання, яка називається мережею глибокої уваги, керованою реакціями користувачів (URGDAN). Розроблено архітектуру глибокої нейронної мережі, що поєднує попередньо навчену модель FinBERT для інтелектуального аналізу фінансових текстів та двонаправлену рекурентну нейронну мережу (BiGRU) для вилучення семантичних ознак із послідовності поточних звітів. Щоб врахувати значущість подій, описаних у поточних звітах, з точки зору їх впливу на виникнення фінансових труднощів, запропоновано механізм уваги, керований реакціями користувачів. Цей метод поєднує ознаки реакцій користувачів із семантичними характеристиками для адаптивного вилучення інформації про події, які тісно пов'язані з фінансовими труднощами компанії. Створено спільне представлення семантичних характеристик і реакцій користувачів, щоб максимально використовувати взаємодію між цими двома різними типами ознак для прогнозування фінансових труднощів. Використовуючи це спільне представлення, яке забезпечує механізм уваги, керований реакціями користувачів, присвоюються різні ступені важливості різним поточним звітам.

Модель URGDAN оцінюється на основі набору даних компаній з національного ринку бірж акцій і котирувань (NEEQ) Китаю, який охоплює акції, що торгуються на позабіржовому ринку, але не котируються на основних фондових біржах. NEEQ є важливим ринком торгівлі акціями в Китаї, і він швидко розвивається в останні роки. Водночас цей ринок характеризується неактивною торгівлею та відносно нижчою прозорістю інформації, що може зробити стандартні методи аналізу, що застосовуються для публічних компаній, непридатними. Це створює труднощі в моделюванні та аналізі й ускладнює точне прогнозування фінансових труднощів порівняно з іншими ринками. Проведено комплексні експерименти для оцінки ефективності моделі у прогнозуванні фінансових труднощів, зокрема порівняння URGDAN з еталонними методами прогнозування, порівняння різних підходів до поєднання реакцій користувачів для керування мережею глибокої уваги, а також

абляційний аналіз. Емпіричні результати показують, що URGDAN повністю координує текст поточних звітів і реакції користувачів, що значно підвищує точність прогнозування. Також проведено прикладний аналіз, у якому візуалізовано значущість (вагу) поточних звітів. URGDAN ефективно ідентифікує інформацію про події у поточних звітах, що дозволяє передбачити фінансові труднощі, і, таким чином, забезпечує практичні рекомендації для кредиторів та інвесторів.

1.2 Огляд існуючих рішень

Дослідження прогнозування фінансових труднощів (FDP) зазвичай базуються на використанні облікових показників як основних характеристик, наприклад, фінансових даних з фінансових звітів компаній [18, 19]. Облікові показники загальноновизнано відіграють провідну роль у FDP [20, 21]. У роботах Campbell та ін. [22] і Wu та ін. [23] також було використано структуровані дані з ринку капіталу для FDP. Однак прогнозувальна здатність структурованих даних є обмеженою; дедалі більше дослідників починають поєднувати нові джерела інформації для вилучення характеристик з метою подальшого покращення точності прогнозування, такі як мережеві та текстові дані.

Використовуючи базову інформацію про малі та середні підприємства, дані про банківські платежі, транзакції та мережеву інформацію, Kou та ін. [24] запропонували двоетапний метод багатокритеріального відбору ознак для оптимізації моделі прогнозування банкрутства. Деякі дослідники вилучали ознаки з важливих текстових сегментів періодичних звітів, таких як розділ «Обговорення та аналіз керівництва» (MD&A) у річних звітах [4, 25] або аудиторський звіт у річних звітах [5], для використання в FDP. Окрім річних звітів, для FDP досліджувались ознаки, вилучені з інших текстових джерел, таких як фінансові новини [3], патентна інформація [26], фінансові лексикони настроїв [27, 28] та поточні звіти [7].

Методи прогнозування фінансових труднощів (FDP) можна поділити на статистичні та методи машинного навчання. До представницьких статистичних методів FDP належать лінійний дискримінантний аналіз і лінійні ймовірнісні моделі [29]. Проте такі методи зазвичай використовуються для низьковимірних та малих наборів даних і вимагають суворих статистичних припущень щодо даних і методів (наприклад, змінні повинні відповідати певним умовам розподілу в різних статистичних методах, ознаки не повинні впливати одна на одну тощо) [30]. Методи машинного навчання здатні краще прогнозувати фінансові труднощі компанії, зокрема методи k-ближчих сусідів [31], дерева рішень [32] та опорних векторів [33]. Однак, модель одного класифікатора може бути недостатньою для точного прогнозування фінансових проблем компанії, тому деякі дослідження використовують ансамблеві класифікатори, такі як випадковий ліс [34], екстремальне градієнтне підсилення [2] та адаптивне підсилення [1].

Завдяки здатності вловлювати складні взаємозв'язки у високовимірних даних, методи глибокого навчання застосовуються для FDP та викликають все більшу увагу. Наприклад, Hosaka [35] перетворив фінансові коефіцієнти кожної компанії на градаційні зображення, які використовувались для навчання та тестування згорткової нейронної мережі. Крім використання однієї моделі глибокого навчання, більшість дослідників розробляють нові архітектури відповідно до специфіки ознак. Li та ін. [27] розробили глибоку нейронну мережу (DNN) на основі багатоголової уваги для аналізу настроїв у контексті FDP. Elhoseny та ін. [36] створили новий метод FDP, який включає процес прогнозування на основі DNN та налаштування гіперпараметрів моделі багат шарового перцептронну (MLP) за допомогою алгоритму оптимізації для покращення точності прогнозування. Для забезпечення практичної та значущої інформації для зацікавлених сторін важливо визначити, які ознаки або частини є найбільш важливими для FDP. Наприклад, Mai та ін. [4] застосували метод вилучення ознак, який полягає в послідовному видаленні окремих слів з

вхідного корпусу та спостереженні, як погіршується продуктивність моделі. Matin та ін. [5] виділили ваги уваги з аудиторських звітів для акцентування на словах і фразах, важливих для FDP. Jiang та ін. [26] використали метод SHapley Additive exPlanations та графік часткової залежності для демонстрації важливості ознак у патентних текстах.

Існуючі дослідження показали, що включення текстової інформації (наприклад, річних звітів, фінансових новин та поточних звітів) є ефективною стратегією для прогнозування фінансових труднощів. Як показано в Таблиці 1.1, більшість досліджень з прогнозування фінансових проблем (FDP) зосереджуються на використанні текстового змісту річних звітів, однак цей підхід має проблему інформаційної затримки та не враховує своєчасно важливі події, які відбуваються нерегулярно і впливають на FDP. У порівнянні з річними звітами, інформація, що розкривається у поточних звітах, є більш оперативною та відображає ключові події у ході операційної діяльності компанії.

Таблиця 1.1 - Порівняння нашого дослідження з існуючими дослідженнями прогнозування фінансових труднощів (FDP)

Дослідження	Використані звіти	Методи роботи з дисбалансом класів	Методи підсвічування важливого тексту	Методи побудови прогнозних моделей
Maі та ін. [4]	Річні звіти	×	Увага на основі семантики	CNN, MLP
Matin та ін. [5]	Річні звіти	×	Увага на основі семантики	CNN, LSTM
Wang та ін. [3]	Річні звіти, фінансові новини	×	Н/д	HSB-RS
Li та ін. [27]	Річні звіти	Undersampling	Н/д	SVM, DT, XGB, MLP
Jiang та ін. [7]	Поточні звіти	SMOTE	Н/д	CART, KNN, LR, RF
Zhao та ін. [6]	Річні звіти	×	LDA (візуалізація топ 50 слів)	CNN, MLP
Borchert та ін. [37]	Веб-контент	SMOTE, undersampling	Н/д	CNN, MLP
Це дослідження	Поточні звіти	Focal loss	Увага на основі семантики, увага, керована реакціями користувачів	BiGRU, механізм уваги, MLP

Відтак, у цьому дослідженні зосереджується увага на прогнозуванні фінансових труднощів за допомогою поточних звітів. Jiang та ін. [7] вилучили

семантичні ознаки з поточних звітів і значно покращили точність прогнозування. Проте поточні звіти містять різноманітні події з різними описами та ступенем важливості, і не всі поточні звіти компанії однаково важливі для прогнозування фінансових труднощів. Для цього впроваджено механізм уваги, який підсвічує найбільш значущі поточні звіти для прогнозування FDP. Однак використання уваги, заснованої лише на семантичних характеристиках, може бути складним для безпосереднього визначення важливості подій у кожному звіті. У цьому дослідженні пропонується механізм уваги, керований реакціями користувачів, який поєднує ознаки реакцій користувачів із семантичними характеристиками для адаптивного вилучення інформації про події, що значно пов'язані з фінансовими труднощами компанії.

Узагальнюючи огляд існуючих рішень для прогнозування фінансових труднощів (FDP), можна зробити висновок, що сучасні методи FDP здебільшого базуються на облікових показниках та структурованих даних, проте їхня ефективність є обмеженою через недостатнє врахування нових джерел інформації, таких як текстові та мережеві дані. Застосування методів машинного навчання, зокрема ансамблевих моделей та глибокого навчання, дозволило значно підвищити точність прогнозів за рахунок здатності вловлювати складні взаємозв'язки у великих наборах даних. Включення текстової інформації, такої як річні звіти, фінансові новини та поточні звіти, виявилось ефективною стратегією для прогнозування фінансових труднощів, однак проблема інформаційної затримки в річних звітах підкреслює важливість використання оперативніших джерел, таких як поточні звіти. У цьому контексті актуальним є розвиток механізмів уваги, зокрема тих, що враховують реакції користувачів, для покращення точності прогнозування фінансових труднощів, що підкреслює важливість комбінування семантичних ознак з реальними реакціями на події, описані у звітах.

1.3 Постановка задачі

В умовах сучасної економічної нестабільності та високої конкуренції фінансові труднощі компаній можуть суттєво впливати на економіку в цілому. Прогнозування фінансових проблем є ключовим аспектом для інвесторів, кредиторів та інших учасників ринку, оскільки воно дозволяє вчасно виявляти потенційні ризики й приймати превентивні заходи. Традиційні методи прогнозування, які ґрунтуються переважно на фінансових показниках, не завжди враховують зміни в реальному часі, що знижує їх ефективність. Це підкреслює необхідність розробки нових моделей, що інтегрують сучасні технології обробки даних, такі як глибоке навчання та аналіз текстової інформації з поточних звітів компаній.

З розвитком цифрових технологій з'являються нові джерела даних, які можуть надати більше інформації для прогнозування фінансових труднощів. Поточні звіти компаній, опубліковані на фінансових платформах, містять актуальні дані про ключові події, які можуть безпосередньо впливати на фінансовий стан компанії. Крім того, реакції користувачів на ці звіти — перегляди, лайки, коментарі — можуть слугувати додатковими індикаторами важливості та впливовості подій. Використання цих нових джерел інформації дозволяє створювати більш точні та адаптивні моделі для прогнозування фінансових труднощів.

Застосування методів глибокого навчання до аналізу текстових даних та поведінкових реакцій користувачів відкриває нові можливості для покращення фінансового аналізу. Глибоке навчання здатне ефективно обробляти великі масиви інформації та виявляти складні зв'язки між різними змінними. Зокрема, моделі на основі нейронних мереж, такі як FinBERT та BiGRU, дозволяють вилучати семантичні ознаки з тексту, а також моделювати послідовні дані, що є особливо корисним для аналізу поточних звітів компаній. Це робить такі

підходи актуальними для досліджень у сфері прогнозування фінансових ризиків.

Актуальність дослідження також зумовлена необхідністю інтеграції різних типів даних для досягнення більшої точності прогнозів. Фінансові показники, які традиційно використовуються в аналізі, не завжди можуть відображати всі аспекти діяльності компанії, особливо коли йдеться про зміни в реальному часі. Включення текстових даних та даних реакцій користувачів дозволяє моделі враховувати не лише фінансові результати, але й поведінкові фактори, що можуть впливати на фінансову стійкість компанії. Це забезпечує більш комплексний підхід до оцінки ризиків.

Розробка моделей, здатних враховувати нові джерела інформації, є важливою для підтримки сталого розвитку бізнесу та фінансових ринків. Вчасне виявлення фінансових труднощів дає змогу компаніям та їхнім партнерам вживати необхідних заходів для стабілізації ситуації, уникати банкрутств та мінімізувати втрати. Тому дослідження в цій галузі є не лише науково значущими, але й мають велике практичне значення для економіки.

Таким чином, актуальність даного дослідження полягає у розробці та впровадженні нових підходів до прогнозування фінансових труднощів компаній на основі глибокого навчання. Інтеграція семантичних ознак з тексту поточних звітів та даних реакцій користувачів дозволить суттєво покращити точність прогнозів та сприятиме розвитку нових інструментів для оцінки фінансових ризиків.

Отже, метою даної роботи є розробка та впровадження методу прогнозування фінансових труднощів компаній на основі глибокого навчання, що інтегрує семантичні ознаки поточних звітів та реакції користувачів для підвищення точності прогнозів.

Задачі, які необхідно вирішити для досягнення поставленої мети, включають:

1. Провести аналіз сучасних методів прогнозування фінансових труднощів на основі поточних звітів і реакцій користувачів.
2. Огляд та оцінка існуючих рішень для прогнозування фінансових труднощів.
3. Спроектувати архітектуру методу URGDAN для прогнозування фінансових труднощів.
4. Використати FinBERT та BiGRU для моделювання послідовних даних із поточних звітів.
5. Оцінити реакції користувачів як показник важливості поточних звітів.
6. Розробити адаптивне коригування семантичних ознак на основі реакцій користувачів.
7. Моделювати фінансові ознаки компаній за допомогою двошарового MLP.
8. Провести огляд даних та опис характеристик обраного набору даних.
9. Оцінити ефективність методу URGDAN у порівнянні з іншими підходами прогнозування.

Отже, дослідження в області прогнозування фінансових труднощів компаній є критично важливим для забезпечення стабільності фінансових ринків і сталого розвитку бізнесу. Інтеграція сучасних підходів глибокого навчання з аналізом текстових даних і реакцій користувачів відкриває нові можливості для точного та своєчасного виявлення потенційних ризиків. Розробка методу URGDAN, який поєднує фінансові, семантичні та поведінкові дані, дозволить покращити ефективність прогнозування, мінімізувати фінансові втрати та сприяти прийняттю більш обґрунтованих управлінських рішень. Це дослідження робить вагомий внесок у розвиток інструментів оцінки фінансових ризиків і надає практичні рішення для подолання викликів сучасної економіки.

Висновки до розділу 1

1. У ході цього дослідження було розглянуто сучасні методи прогнозування фінансових труднощів (FDP) із використанням поточних звітів та реакцій користувачів. Основна увага приділялася тому, як інтеграція текстових ознак із реакціями користувачів може підвищити точність прогнозування. Виявлено, що поточні звіти є важливим джерелом інформації, оскільки вони розкривають ключові події, що безпосередньо впливають на фінансовий стан компанії, на відміну від річних звітів, які часто містять інформацію з затримкою. Це дослідження також підкреслює значення різних типів реакцій користувачів, таких як читання, вподобання та коментарі, як сигналів, що допомагають визначити важливість подій, описаних у поточних звітах.

2. Запропонований механізм уваги, керований реакціями користувачів (URGDAN), показав високу ефективність у поєднанні семантичних ознак та реакцій користувачів для прогнозування фінансових труднощів. Модель URGDAN дозволяє краще визначати найбільш значущі звіти компанії, використовуючи як текстову інформацію, так і поведінкові реакції користувачів. Цей підхід забезпечує більш повне розуміння фінансових ризиків компанії та допомагає виявляти потенційні проблеми на ранніх етапах. Експерименти підтвердили, що використання механізму уваги на основі реакцій користувачів значно покращує результати порівняно з традиційними методами, які враховують лише фінансові індикатори або текстові ознаки.

3. Загалом, результати дослідження показують, що комбінування фінансових, текстових і поведінкових даних є ефективною стратегією для прогнозування фінансових труднощів. Запропонована модель URGDAN виявилася корисною не лише для підвищення точності прогнозування, але й для надання кредиторам та інвесторам інструментів для кращого розуміння ситуації компанії. Це дослідження робить внесок у розвиток практик прогнозування

фінансових ризиків, підкреслюючи необхідність подальшого дослідження інтеграції різних джерел даних для поліпшення фінансового аналізу.

2 АРХІТЕКТУРНІ РІШЕННЯ ТА МЕТОДИ ПРОГНОЗУВАННЯ ФІНАНСОВИХ ТРУДНОЩІВ НА ОСНОВІ ГЛИБОКОГО НАВЧАННЯ

2.1 Проектування архітектури методу URGDAN для прогнозування фінансових труднощів

Для підвищення ефективності методу прогнозування фінансових труднощів (FDP) з використанням поточних звітів пропонується новий підхід на основі глибокого навчання — мережа глибокої уваги, керована реакціями користувачів (URGDAN). Основна ідея URGDAN полягає в тому, що реакції користувачів допомагають виділити важливі поточні звіти та ключові події, що дозволяє використовувати їхні семантичні характеристики в поєднанні з обліковими показниками для кращого прогнозування фінансових труднощів.

На рисунку 2.1 представлена загальна архітектура запропонованого методу. Ми інтегруємо попередньо навчену модель FinBERT та модель BiGRU для побудови семантичного вектора ознак з послідовності поточних звітів. Цей метод ефективно кількісно оцінює семантику окремого поточного звіту та взаємозв'язки між послідовними подіями, що розкриваються. Додатково пропонується механізм уваги, керований реакціями користувачів, який поєднує реакції користувачів на поточні звіти із семантичною інформацією, адаптивно витягуючи інформацію про події, що суттєво пов'язані з фінансовими труднощами компанії. Запропонований метод описується з трьох основних аспектів: представлення ознак гетерогенної інформації, механізм уваги, керований реакціями користувачів (UR на схемі), та злиття ознак і прогнозування.

Архітектура URGDAN представляє новий підхід до прогнозування фінансових труднощів, який ґрунтується на глибокому навчанні з урахуванням реакцій користувачів. Основна перевага URGDAN полягає в інтеграції семантичних ознак поточних звітів, отриманих за допомогою моделей FinBERT та BiGRU, з обліковими показниками та реакціями користувачів.

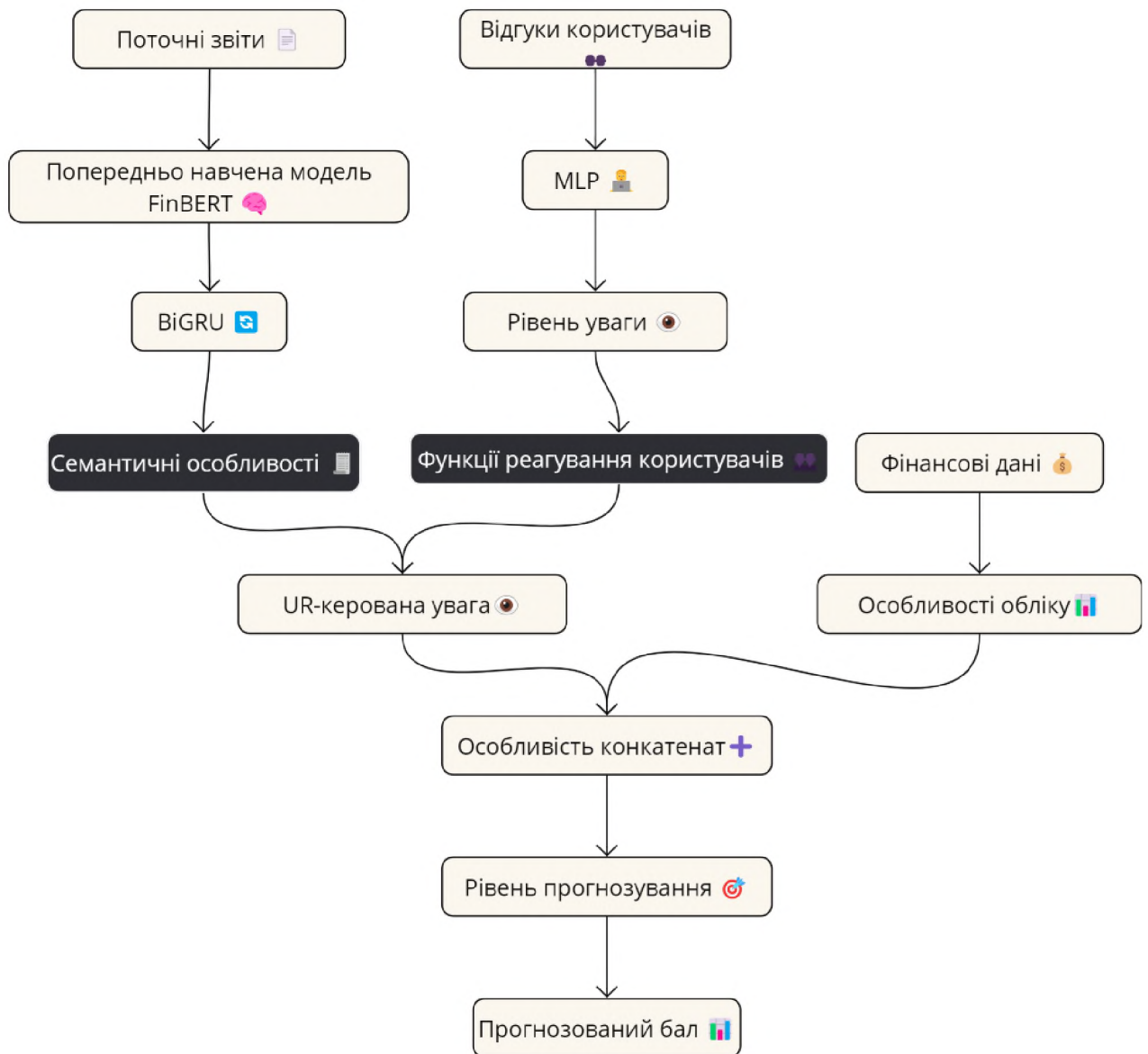


Рисунок 2.1 - Загальна архітектура URGDAN

Це дозволяє покращити точність прогнозування фінансових труднощів, оскільки мережа глибокої уваги може адаптивно виділяти найважливіші події та звіти. Механізм уваги, керований реакціями користувачів, забезпечує адаптивне вилучення релевантної інформації, що суттєво пов'язана з фінансовими ризиками компанії, тим самим підвищуючи ефективність прогнозування.

2.2 Моделювання послідовних даних з використанням FinBERT та BiGRU для обробки поточних звітів

Порівняно з періодичними звітами, поточні звіти мають дві суттєві переваги. По-перше, вони відображають багато ключових подій, що відбуваються в операційній діяльності компанії. Ці події, особливо ризикові, можуть бути не повністю висвітлені в періодичних звітах. По-друге, поточні звіти своєчасно розкривають інформацію про важливі події компанії. У нашому дослідженні увага зосереджується на інформаційній цінності поточних звітів, які можуть містити унікальні сигнали ризику, і ми витягуємо з них семантичні ознаки.

Представлення семантичних ознак — це векторне представлення, перетворене з тексту відповідно до його семантики. Класичними методами для вилучення та кількісного аналізу текстових ознак є векторні моделі простору та word2vec. З розвитком методів обробки природної мови, попередньо навчені моделі досягли відмінних результатів у вилученні текстових ознак, серед яких найбільш репрезентативною є модель двонаправленого кодування на основі трансформерів (BERT). FinBERT, що є адаптованою версією BERT для фінансової сфери, отриманий шляхом додаткового попереднього навчання на фінансовому корпусі та тонкого налаштування моделі для семантичного аналізу. До фінансового корпусу входять річні звіти компаній, стенограми телефонних конференцій з прибутків та звіти аналітиків. Згідно з результатами класифікаційних задач, модель FinBERT перевершує загальну модель BERT за точністю [38].

Компанія випускає кілька поточних звітів протягом року, і текстовий вміст цих звітів можна розглядати як послідовні дані, оскільки поточні звіти, опубліковані в різні моменти часу, мають певний зв'язок з відповідними періодами. Проте традиційні методи, такі як "мішок слів," втрачають контекстуально-залежну інформацію, оскільки не враховують порядок слів у

тексті. Для моделювання послідовних даних використовується рекурентна нейронна мережа з блоками контрольованої рекурсії (GRU). Основною метою моделі BiGRU є вилучення глибоких семантичних ознак вхідних текстових векторів та встановлення взаємозв'язків між послідовно опублікованими подіями. BiGRU може обробляти текст як у прямому, так і в зворотному напрямках, що дозволяє повноцінно враховувати взаємозв'язки між контекстами та вилучати семантичні ознаки на рівні компанії.

Модель FinBERT використовується для перетворення тексту кожного поточного звіту на семантичне представлення ознак з розмірністю 768. Для кожного поточного звіту вводимо його текст у попередньо навчену модель FinBERT і вилучаємо результат токена CLS з останнього шару трансформера. Вивід цього токена часто використовується як векторне представлення всього тексту. Обробляючи результат, отриманий від FinBERT, семантичні ознаки m -го поточного звіту i -ї компанії представляються як вектор ознак $Fvec_{i,m} = [Fvec_{i,m,1}, \dots, Fvec_{i,m,n}]$. Далі, для захоплення семантичних ознак окремого поточного звіту та взаємозв'язків між послідовними подіями, введеними хронологічно, використовується модель BiGRU. BiGRU складається з двох GRU: прямої та зворотної. Ці дві мережі тренуються у прямому та зворотному напрямках і з'єднуються на одному шарі для забезпечення повного врахування контекстуальних взаємозв'язків. Результат обчислюється за наступними рівняннями:

$$\begin{aligned} h_m^{\rightarrow} &= GRU^{\rightarrow}(h_{m-1}^{\rightarrow}, Fvec_{i,m}) \\ h_m^{\leftarrow} &= GRU^{\leftarrow}(h_{m-1}^{\leftarrow}, Fvec_{i,m}) \\ Ftext_{i,m} &= [h_m^{\rightarrow}, h_m^{\leftarrow}] \end{aligned} \tag{2.1}$$

де $h_m^{\rightarrow}$ — стан прихованого шару прямої GRU на m -му кроці;

$h_{m-1}^{\rightarrow}$ — стан прихованого шару прямої GRU на попередньому $(m - 1)$ кроці;

$Fveci, m$ — вектор семантичних ознак для m -го поточного звіту i -ї компанії, отриманий з FinBERT;

$GRU^{\rightarrow}$ — функція прямої GRU, що обчислює новий стан прихованого шару, використовуючи попередній стан $h_{m-1}^{\rightarrow}$ та вхідні семантичні ознаки $Fveci, m$.

$h_m^{\leftarrow}$ — стан прихованого шару зворотної GRU на m -му кроці;

$h_{m-1}^{\leftarrow}$ — стан прихованого шару зворотної GRU на попередньому $(m - 1)$ кроці;

$Fveci, m$ — той самий вектор семантичних ознак для m -го звіту i -ї компанії;

$GRU^{\leftarrow}$ — функція зворотної GRU, яка працює у зворотному напрямку, обчислюючи новий стан прихованого шару.

$Ftexti, m$ — об'єднаний вектор текстових семантичних ознак для m -го поточного звіту i -ї компанії;

Таким чином, використання FinBERT та BiGRU для моделювання послідовних даних поточних звітів компаній забезпечує можливість ефективного вилучення семантичних ознак та встановлення взаємозв'язків між подіями, які відображаються у звітах. Завдяки FinBERT текст поточних звітів перетворюється у векторні представлення, що мають високу семантичну точність, а BiGRU дозволяє враховувати як контекст попередніх подій, так і майбутніх, завдяки одночасній роботі прямої та зворотної GRU. Це робить можливим точніше прогнозування ризиків і прийняття своєчасних рішень, що має велике практичне значення для аналізу операційної діяльності компаній та оцінки їхнього фінансового стану. Використання такого підходу дозволяє не лише покращити аналіз ризиків, але й сприяє розвитку нових інструментів для управління фінансовими ризиками в умовах динамічних змін.

2.3 Оцінка реакцій користувачів як показник важливості поточних звітів

На деяких соціальних платформах кількість лайків та коментарів, отриманих брендом або продуктом, є важливим показником, який відображає залученість користувачів та їхню реакцію [39]. Крім того, різні типи реакцій користувачів можуть мати різний ступінь впливу на компанію [15]. Лайки можуть допомагати підтвердити позитивність і важливість новин, збільшуючи ймовірність того, що вони будуть запам'ятовані та матимуть вплив у майбутньому [14]. Таким чином, деякі типи реакцій користувачів, ймовірно, краще відображають важливість поточних звітів відповідної компанії, ніж інші.

У нашому дослідженні ми використовуємо кількість переглядів, лайків та коментарів для кожного поточного звіту як міру реакції користувачів. Замість того, щоб безпосередньо застосовувати витягнуті ознаки реакцій користувачів для керування семантичними ознаками, ми використовуємо механізм уваги [40], щоб підкреслити найбільш цінні реакції користувачів, які краще відображають їхній зворотний зв'язок із поточними звітами. Спочатку ми використовуємо багатошаровий перцептрон (MLP) для отримання векторів ознак реакцій користувачів для m -го поточного звіту i -ї компанії, позначеного як $F_{i,m}^{ur} = [F_{1,m}^{ur}, \dots, F_{n,m}^{ur}]$. Ми обчислюємо вагу реакції користувачів як подібність між вектором W_a та $\bar{F}_{i,m}^{ur}$. Нарешті, нормалізовані ваги отримуються за допомогою функції softmax, де кожен елемент $atten_{ur,i,m}$ представляє оцінку важливості певної реакції користувачів на m -й поточний звіт i -ї компанії. Таким чином, нове представлення ознак $\bar{F}_{i,m}^{ur}$ створюється шляхом зваженого підсумовування ознак реакцій користувачів, що використовується для опису реакцій користувачів на кожен поточний звіт. Обчислення оптимізованого представлення відбувається за такими рівняннями:

$$\begin{aligned}
\bar{F}_{i,m}^{ur} &= \tanh(W_u \cdot F_{i,m}^{ur} + b_{ur}) \\
\text{atten}_{i,m}^{ur} &= \text{softmax}(W_a^T \cdot \bar{F}_{i,m}^{ur}) \\
&= ur \\
F_{i,m} &= \text{atten}_{i,m}^{ur} \cdot F_{i,m}^{ur}
\end{aligned} \tag{2.2}$$

де $F_{i,m}^{ur}$ — вхідний вектор ознак реакцій користувачів для m -го звіту i -ї компанії;

W_u — вагова матриця багатoshарового перцептрона (MLP), що використовується для трансформації вхідного вектора;

b_{ur} — вектор зміщення, що додається до результату;

$\tanh(\cdot)$ — активаційна функція гіперболічного тангенса, яка забезпечує нелінійність і нормалізує значення до діапазону $[-1, 1]$;

$\bar{F}_{i,m}^{ur}$ — трансформований вектор ознак реакцій користувачів.

W_a — ваговий вектор, який використовується для обчислення важливості окремих реакцій;

$\text{atten}_{i,m}^{ur}$ — вага уваги для кожної реакції користувачів на m -й звіт i -ї компанії.

Реакції на ціни акцій, якби їх можна було точно виміряти, були б кращим показником важливості поточних звітів. Однак на ринках із недостатніми даними торгів та відносно низькою прозорістю інформації (наприклад, ринок NEEQ), реакції користувачів можуть бути більш раціональним вибором як прямий сигнал, що відображає рівень уваги до розкритих подій.

Хоча майже кожен користувач онлайн може згенерувати реакцію на поточний звіт, більшість таких реакцій походить від зацікавлених сторін, зокрема інституційних та роздрібних інвесторів. Інституційні інвестори зазвичай мають досвід і прагнуть глибоко аналізувати ризики своїх цільових компаній, що дозволяє їхній увазі ефективно відображати важливість поточного звіту. Роздрібні інвестори, хоча кожен із них має відносно менший

обсяг інвестицій, разом демонструють значний ефект довгого хвоста, а їхня увага може відображати важливість поточного звіту завдяки ефекту "колективної мудрості" (тобто великі групи людей колективно розумніші, ніж окремі експерти) та ефекту "стадності" (коли наступні інвестори слідують діям попередніх при перегляді поточних звітів).

Адаптивне коригування семантичних ознак на основі реакцій користувачів дозволяє значно підвищити точність оцінки важливості поточних звітів. Використання механізму уваги, керованого реакціями користувачів, забезпечує інтеграцію текстових семантичних ознак із впливом реакцій користувачів, створюючи більш насичене та релевантне представлення для аналізу. Завдяки цьому підходу стає можливим виділити звіти, які є критично важливими для прогнозування фінансових труднощів, та відсіяти менш значущі дані. Така інтеграція допомагає врахувати як текстову інформацію, так і поведінкові сигнали, що дозволяє покращити моделі аналізу ризиків та забезпечує гнучке адаптивне налаштування для змін у даних. Розроблений підхід сприяє створенню більш точних та ефективних інструментів для фінансового прогнозування.

2.4 Адаптивне коригування семантичних ознак на основі реакцій користувачів

Реакції користувачів є зворотним зв'язком на основі змісту поточного звіту, і вони вимірюють важливість цього змісту, передаючи інформацію, яка може бути відсутня в самому тексті звіту. Вони корисні для визначення важливості різних поточних звітів та фільтрації тих, що є неактуальними або незначними для прогнозування фінансових труднощів (FDP). Ми розробили механізм уваги, керований реакціями користувачів, який використовує ознаки реакцій користувачів для керування семантичним представленням ознак поточних звітів. Цей механізм уваги, керований реакціями користувачів,

зображений на рисунку 2.2. Модуль здатний адаптивно коригувати семантичні ознаки відповідно до реакцій користувачів на кожен поточний звіт.

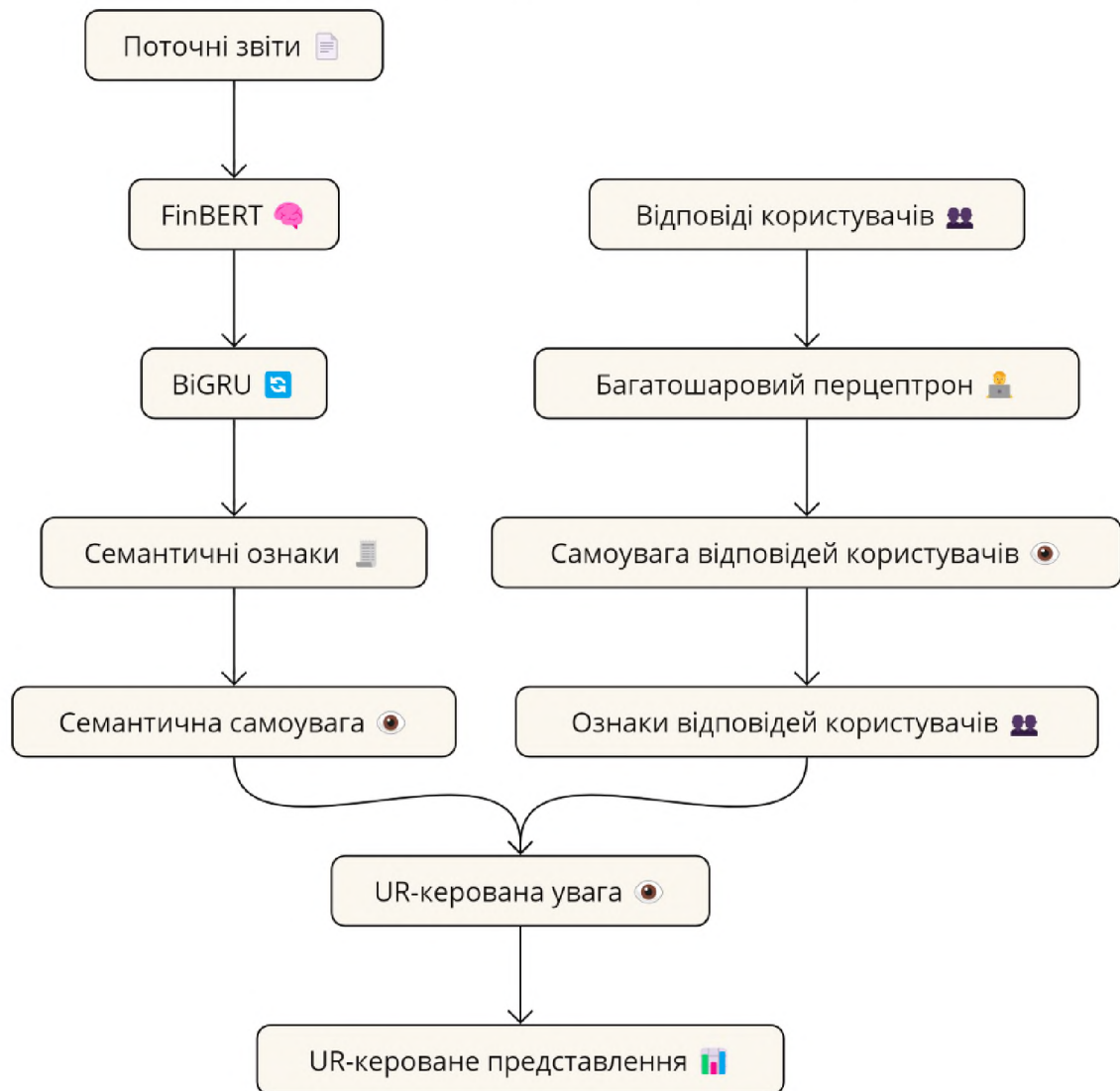


Рисунок 2.2 - Модуль уваги, керований реакціями користувачів

На основі семантичного представлення ознак $F_{i,m}^{text}$, спочатку розраховуємо нове текстове семантичне представлення ознак, $\bar{F}_{i,m}^{text}$, шляхом поєднання оригінального представлення з представленням ознак реакції користувачів $\bar{F}_{ur,i,m}$. Застосовано шар повного з'єднання для узгодження розмірності векторів семантичних ознак із розмірами представлення ознак $\bar{F}_{i,m}^{ur}$.

Для кожного поточного звіту розраховуємо відповідний бал уваги $att_{text,i,m}$ за допомогою схожості між вектором W_T^{text} та $\bar{F}_{i,m}^{text}$. Ми створюємо

нове представлення текстової семантики, позначене як $\bar{\bar{F}}_i^{text} = [\bar{F}_{i,1}^{text}, \dots, \bar{F}_{i,m}^{text}]$, за допомогою зваженої суми всіх семантичних ознак поточних звітів однієї компанії.

Деталі цієї операції наведено нижче:

$$\begin{aligned}\bar{F}_{i,m}^{text} &= \tanh(W_{text} \bullet F_{i,m}^{text}) \odot \tanh(W_{ur} \bullet F_{i,m}^{text}) \\ \text{atten}_{i,m}^{text} &= \text{softmax}(W_{text}^T \bullet \bar{F}_{i,m}^{text}) \\ &=_{text} \\ F_i &= \text{atten}_{i,m}^{text} \bullet F_{i,m}^{text}\end{aligned}\tag{2.3}$$

де W_{text} та W_{ur} є параметрами, які ініціалізуються випадковим чином і навчаються в процесі тренування.

$\bar{F}_{i,m}^{text}$ — це приховане представлення семантичних ознак тексту;
 $\text{atten}_{i,m}^{text}$ — це коефіцієнт уваги для m -го поточного звіту в i -й компанії;

$\bar{\bar{F}}_i^{text}$ — це векторне представлення нової семантичної ознаки, що інтегрує ознаки реакції користувачів.

2.5 Моделювання фінансових ознак компаній за допомогою двошарового MLP

Для кожної компанії певні фінансові індикатори є структурованими ознаками, які відображають її фінансовий стан. Облікові ознаки є основними характеристиками, які зазвичай використовуються в прогнозуванні фінансових труднощів (FDP). Через невідповідність між важливістю ознак та їх розмірністю було вдосконалено метод для ефективного навчання з використанням структурованих ознак. Розроблено підсилювач структурованих ознак, який є двошаровим багатошаровим перцептроном (MLP) із залишковими з'єднаннями. Підсилювач структурованих ознак збільшує кількість нейронів у

двох прихованих шарах, що дозволяє підвищити розмірність прихованих ознак та покращити лінійну роздільність.

Підсилювач структурованих ознак отримує представлення ознак відповідно до фінансових індикаторів. Облікові ознаки i -ї компанії представлені як $Facci = [Facci,1, \dots, Facci,n]$. На основі попередніх експериментальних операцій тепер отримано об'єднані ознаки $Fconcat$, які поєднують ознаки $[\bar{F}texti, Facci]$. Об'єднані ознаки передаються через два повнозв'язані шари з функцією активації ReLU. На фінальному етапі використовується одношаровий нейронний шар із активацією сигмоїди для отримання передбачуваної ймовірності фінансових труднощів. Деталі процесу представлені наступними рівняннями:

$$\begin{aligned}
 F_{concat} &= \begin{bmatrix} \text{= text} \\ F_i, F_i^{acc} \end{bmatrix} \\
 Layer_1 &= ReLU(W_1 \bullet F_{concat} + b_1) \\
 Layer_2 &= ReLU(W_2 \bullet Layer_1 + b_2) \\
 FDP &= sigmoid(W_3 \bullet Layer_2 + b_3)
 \end{aligned} \tag{2.4}$$

де $W_1, W_2, W_3, b_1, b_2, b_3$ — це параметри, що підлягають навчанню;

FDP - позначає ймовірність фінансових труднощів.

Для пом'якшення впливу дисбалансу класів використовується функція втрат focal loss. Focal loss є модифікацією функції втрат на основі крос-ентропії, яка вирішує проблему дисбалансу класів. Вона працює шляхом збільшення ваг для важко класифікованих зразків та зменшення ваг для зразків, які легко класифікувати. Зменшуючи важливість легко класифікованих зразків, focal loss спонукає модель приділяти більше уваги важким зразкам, що знаходяться поблизу межі рішення під час навчання. Це підвищує стійкість моделі до шуму та викидів [41].

Висновки до розділу 2

1. Запропонована архітектура URGDAN для прогнозування фінансових труднощів на основі глибокого навчання продемонструвала високу ефективність у порівнянні з іншими методами. Основна перевага URGDAN полягає в інтеграції текстових ознак поточних звітів та реакцій користувачів, що дозволяє покращити прогнозування за рахунок більш точного виділення релевантних подій. За допомогою механізму уваги, керованого реакціями користувачів, система адаптивно виділяє важливі елементи інформації, що мають суттєвий вплив на фінансовий стан компанії.

2. Використання моделей FinBERT та BiGRU для моделювання послідовних даних поточних звітів дозволило отримати глибокі семантичні ознаки, що зберігають контекстуальні зв'язки між подіями. Це підвищує точність прогнозування фінансових труднощів, оскільки враховується не лише зміст окремих звітів, але й їхній взаємозв'язок у часі. У цьому контексті FinBERT забезпечує високу якість семантичного представлення тексту, що суттєво впливає на загальну продуктивність моделі.

3. Дослідження показало, що додавання реакцій користувачів, таких як перегляди, лайки та коментарі, значно підвищує точність прогнозування. Реакції користувачів є корисними сигналами, що відображають увагу аудиторії до важливих подій, описаних у поточних звітах. Механізм уваги, керований цими реакціями, забезпечує оптимальне вилучення значущих даних і поліпшує стійкість моделі до викидів, дозволяючи краще враховувати складні сценарії фінансових ризиків.

3 ПРАКТИЧНЕ ВИКОРИСТАННЯ ОТРИМАНИХ НАУКОВИХ РЕЗУЛЬТАТІВ

3.1 Огляд даних та характеристика обраного набору даних

Для перевірки ефективності запропонованого методу зібрано дані для створення набору даних з трьох вебсайтів: список компаній з NEEQ (<https://www.neeq.com.cn/>); фінансові дані з розділу фінансових індикаторів NEEQ у базі даних China Stock Market & Accounting Research (CSMAR) (<http://cndata1.csmar.com/>); та текстові дані поточних звітів і дані реакцій користувачів з переліку звітів компаній на вебсайті Eastmoney (<https://www.eastmoney.com/>), відомій платформі фінансової інформації. Відповідно до попередніх робіт (наприклад, [18, 3]), побудували бінарну змінну, що вказує на фінансові труднощі, та використали статус спеціального нагляду (ST) як символ фінансових труднощів. Відповідно до бізнес-правил NEEQ, компанії з від'ємними чистими активами на кінець фінансового року позначаються як ST [27, 42]. Зібрано фінансові дані, текстові дані та дані реакцій користувачів для 8038 компаній у 2017 році та визначили їх фінансовий стан у 2019 році (чи отримали вони статус ST, чи ні). За офіційними даними, наш набір даних включає 169 компаній, які отримали статус ST у 2019 році, та 7869 нормальних компаній у 2019 році. Для фінансових даних виміряно всі облікові ознаки, використовуючи дані станом на кінець 2017 року. Для поточних звітів ми зібрали загалом 300 621 поточний звіт, оприлюднений у 2017 році, та використали весь текст кожного поточного звіту. Відкинуто деякі типи звітів, які мають незначну релевантність до фінансового ризику, такі як повідомлення про загальні збори та правила процедури, залишивши 241 388 поточних звітів у наборі даних для оцінки.

Відповідно до наявної літератури (наприклад, [18, 1]) та видаливши деякі ознаки з великою кількістю відсутніх значень, отримано 13 ознак, пов'язаних з операційною здатністю, рентабельністю та платоспроможністю, на основі

показників, оприлюднених у річній фінансовій звітності. Таблиця 3.1 містить підсумкову статистику облікових ознак.

Таблиця 3.1 - Описова статистика особливостей фінансового обліку

№	Ознака	Зведена статистика				
		Mean	Max	Min	Median	SD
1	Коефіцієнт струму (%)	3.586	376.255	0.024	1.999	9.015
2	Швидкий коефіцієнт (%)	2.952	376.255	0.021	1.528	8.137
3	Оборотний капітал (у мільйонах юанів)	49.230	14,355.847	-23 493,573	26.459	381.859
4	Коефіцієнт заборгованості (%)	0.407	41.272	0.000	0.393	0.504
5	Рентабельність загальних активів (%)	0.022	23.443	-4.476	0.043	0.328
6	Рентабельність власного капіталу (%)	-0,094	4.998	-138.659	0.076	2.806
7	Прибуток до сплати відсотків і податків (у мільйонах юанів)	13.098	1682.349	-550.260	5.926	48.236
8	Маржа валового прибутку (%)	0.310	2.010	-147.406	0.310	1.670
9	Маржа операційного прибутку (%)	-0,331	362.083	-805.526	0.052	13.377
10	Рентабельність власного капіталу, що припадає на материнську компанію (%)	-0,082	1.698	-138.659	0.078	2.598
11	Оборотність дебіторської заборгованості (річна)	30.693	9944.683	0.015	3.507	281.792
12	Оборотність запасів (річна)	0.322	115.532	0.000	0.163	1.747
13	Загальна оборотність активів (річна)	4.406	8732.580	0.024	1.369	106.436

Кількість поточних звітів, випущених різними компаніями протягом року, майже завжди відрізняється. Проте деякі дані, введені в модель BiGRU, повинні бути доповнені для досягнення однакової довжини. Оскільки занадто довгі послідовності можуть призвести до зникнення або вибуху градієнта, відповідно до Dwarampudi і Reddy [43], ми використовували комбінацію попереднього скорочення послідовності та подальшого доповнення для вирівнювання довжини вхідних даних.

Усі поточні звіти мають відповідні дані про реакції користувачів (тобто кількість переглядів, лайків і коментарів). Існують помітні відмінності у реакціях користувачів на різні поточні звіти.

Таблиця 3.2 представляє середню кількість публікацій, переглядів, лайків і коментарів для деяких типів поточних звітів від компаній з ST та без ST статусу.

Таблиця 3.2 - Зведена статистика поточних звітів та відповідей користувачів

№	Тип звіту	Компанії ST				Компанії, що не входять до ST			
		Пос т	Читанн я	Комент а р	Подобаєтьс я	Пос т	Читанн я	Комент а р	Подобаєтьс я
1	Корпоратив не управління	15.8 1	5023.89	2.55	2.24	13.1 8	5812.95	3.01	5.19
2	Зміна інформації	4.16	6348.99	5.07	5.47	3.22	6112.13	2.27	4.78
3	Ризик має значення	4.14	4438.85	2.35	0.91	0.62	4914.36	1.23	0.94
4	Транзакція маркет- мейкера	1.83	4402.45	7.25	7.95	0.60	5419.67	4.04	5.54

Після проведення рангового тесту Вілкоксона виявлено значні відмінності в кількості публікацій та реакцій користувачів для певних типів звітів між компаніями зі статусом ST та без нього. Наприклад, спостерігаються відмінності у кількості публікацій та коментарів щодо ризикових подій, а також у кількості публікацій і реакцій користувачів на транзакції маркет-мейкерів.

3.2 Оцінка ефективності різних підходів до прогнозування фінансових ризиків

Прогнозування фінансових труднощів (FDP) зазвичай розглядається як задача бінарної класифікації, тобто визначення того, чи потрапить компанія в фінансові труднощі в найближчому майбутньому. Для порівняння ефективності різних методів ми відібрали шість репрезентативних методів з попередніх досліджень з прогнозування фінансових труднощів (наприклад, [3, 4, 7, 36]). Зокрема, ми використовували метод KNN як представник методів, заснованих на екземплярах, LR — як представник лінійних методів, RF і XGB — як представники методів ансамблевого навчання, а MLP і CNN — як представники методів глибокого навчання.

Для забезпечення чесного порівняння, кожен метод отримував на вхід фінансові індикатори, поточні звіти та реакції користувачів. Для обробки поточних звітів використовували серію векторних подань документів, згенерованих за допомогою FinBERT для CNN і URGDAN, оскільки вони мають власні модулі моделювання текстових даних. Для решти методів (KNN, LR, RF, XGB, MLP) семантичні ознаки були витягнуті за допомогою методу, запропонованого Jiang та ін. [7]. Для повноти опишемо процес вилучення семантичних ознак за методом Jiang та ін. [7]. Спочатку було побудовано матрицю «документ-термін», після чого отримували векторні подання слів у кожному документі (поточному звіті) за допомогою FinBERT і обчислювали частотність терміну-інверсну частотність документа (TFIDF) для кожного слова. Далі отримували векторне подання кожного поточного звіту шляхом зважування векторів слів згідно з TFIDF і їх сумування. Потім проводили аналіз головних компонент (PCA) для зменшення розмірності векторів. Нарешті, обчислювали середнє значення векторів усіх поточних звітів компанії як витягнуті семантичні ознаки.

Щодо річних звітів, Li та ін. [27] продемонстрували, що поєднання методу з ознаками сентименту річних звітів покращує ефективність. Відповідно до Li та ін. [27], побудовано сентиментні ознаки з річних звітів, використовуючи Словник Сентименту Національного Тайванського Університету (UTUSD) та китайський позитивний і негативний словники Цінхуа (TSING). Ознаки, які побудовано, представлені в таблиці 3.3. Для реакцій користувачів розраховано середнє значення ознак реакцій користувачів для всіх поточних звітів компанії та нормалізували їх як вхід для кожного з розглянутих методів.

Таблиця 3.3 - Ознаки тональності річних звітів.

Ознака	Визначення
TSING_pos	Кількість позитивних слів у річному звіті, на основі словника TSING
TSING_neg	Кількість негативних слів у річному звіті, на основі словника TSING
TSING_sen	$(TSING_pos - TSING_neg) / (TSING_pos + TSING_neg)$
NTUSD_pos	Кількість позитивних слів у річному звіті, на основі словника NTUSD
NTUSD_neg	Кількість негативних слів у річному звіті, на основі словника NTUSD
NTUSD_sen	$(NTUSD_pos - NTUSD_neg) / (NTUSD_pos + NTUSD_neg)$

Обрано три стандартні агреговані метрики для оцінки дискримінаційної здатності моделей: площу під кривою операційних характеристик приймача (AUC), статистику Колмогорова–Смирнова (KS) і Н-міру. AUC відображає загальну здатність методу розрізняти компанії, що знаходяться у фінансових труднощах, від нормальних компаній. KS використовується для оцінки здатності моделі до дискримінації ризику, а індикатор KS вимірює різницю між кумулятивними розподілами добрих і поганих зразків, відображаючи здатність моделі відрізняти два типи зразків. Також обчислено Н-міру, яка усуває недоліки AUC, що використовує різні розподіли витрат на неправильну класифікацію для різних класифікаторів, задаючи попередньо визначений бета-розподіл для витрат на неправильну класифікацію [44]. Ці метрики показали, що вони не чутливі до дисбалансу даних, що є важливим у прогнозуванні фінансових труднощів, оскільки кількість компаній, які знаходяться в скрутному фінансовому становищі, значно менша, ніж кількість нормальних компаній. Окрім того, враховуючи, що втрата від не детектування компанії у

фінансових труднощах зазвичай більша, ніж неправильна ідентифікація нормальної компанії як такої, що зазнає фінансових труднощів, також обрано показник recall як додаткову метрику продуктивності.

Щоб оцінити продуктивність кожного методу прогнозування, ми провели п'ять незалежних запусків перехресної валідації з 10-ма фолдами, що дало 50 оцінок продуктивності для отримання надійних результатів. Усі середні значення продуктивності, наведені пізніше у статті, ґрунтуються на цих 50 оцінках. Після розділення на навчальну та тестову вибірки в кожній перехресній валідації з 10 фолдами, ми додатково розділили навчальну вибірку на десять фолдів і випадково обрали один фолд як валідаційну вибірку для налаштування параметрів, залишивши інші дев'ять фолдів як зменшену навчальну вибірку. Для чесного порівняння методів розподіл фолдів залишався однаковим для всіх методів протягом кожної 10-кратної перехресної валідації. Для вирішення проблеми дисбалансу класів ми використовували техніку синтетичного додаткового вибіркового методу для меншості (SMOTE) для всіх методів машинного навчання (KNN, LR, RF і XGB) та функцію втрат focal loss для всіх методів глибокого навчання (MLP, CNN і URGDAN). Для налаштування параметрів використовувався метод пошуку на сітці для кожного методу. Деталі простору пошуку гіперпараметрів наведені в таблиці 3.4. Щоб запобігти перенавчанню, ми застосовували ранню зупинку для всіх методів глибокого навчання (MLP, CNN і URGDAN), контролюючи продуктивність на валідаційній вибірці під час навчання.

Таблиця 3.4 - Пошуковий простір гіперпараметрів

Алгоритм	Пошуковий простір
KNN	metric $\in$ {'euclidean', 'manhattan'}; weight $\in$ {'uniform', 'distance'}; neighbors $\in$ {3, 5, 7, 9, 11}
LR	max_iter $\in$ {L1, L2}; penalty $\in$ {L1, L2}; solver $\in$ {'lbfgs', 'newton-cg', 'liblinear', 'sag', 'saga'}
RF	n_estimators $\in$ {50, 100, 150, 200, 250}; max_feature $\in$ {'auto', 'sqrt', 'log2'}; criterion $\in$ {'entropy', 'gini'}; bootstrap $\in$ {'True', 'False'}
XGB	n_estimators $\in$ {50, 100, 150, 200, 250}; learning_rate $\in$ {0.01, 0.05, 0.1, 0.2, 0.3}; booster $\in$ {'gbtree', 'gblinear', 'dart'}
MLP	кількість епох $\in$ [0, 100]
CNN	кількість епох $\in$ [0, 100]

3.3 Аналіз ефективності URGDAN у порівнянні з еталонними методами прогнозування

Щоб перевірити, чи може наш метод глибокого навчання покращити продуктивність прогнозування фінансових труднощів (FDP) за допомогою семантичних ознак у поточних звітах, ми порівняли URGDAN з шістьма еталонними методами.

Таблиця 3.5 узагальнює середні результати п'яти незалежних перехресних валідацій з 10-ма фолдами за всіма метриками продуктивності (AUC, KS, Н-міра та показник recall).

Загалом, URGDAN перевершує всі еталонні методи за всіма метриками продуктивності. Використання лише облікових ознак показує, що фінансові показники компаній є невід'ємною частиною завдання прогнозування фінансових труднощів (FDP). Як підтверджують наявні емпіричні результати у сфері FDP [16, 24], моделі ансамблевого навчання (наприклад, RF і XGB) перевершують моделі з одним класифікатором та ефективно передбачають фінансові труднощі. Після додавання семантичних ознак з поточних звітів продуктивність усіх еталонних методів покращується, особливо для XGB і MLP. Крім того, додавання семантичних ознак з поточних звітів завжди демонструє кращу продуктивність, ніж додавання сентиментних ознак з річних звітів. Це свідчить про те, що поточні звіти дійсно містять семантичну інформацію, тісно пов'язану з фінансовими труднощами, і додавання семантичної інформації з поточних звітів сприяє підвищенню ефективності прогнозування FDP. Додавання ознак реакцій користувачів до облікових та семантичних ознак не призвело до значного покращення продуктивності для всіх еталонних методів.

Таблиця 3.5 - Продуктивність прогнозування з використанням різних
ознак

Ознаки	Метод	AUC		KS		H-міра		Recall	
A	KNN	0.753 (0.730– 0.775)		0.489 (0.451– 0.527)		0.354 (0.317– 0.391)		0.517 (0.481– 0.553)	
	LR	0.817 (0.795– 0.840)		0.605 (0.572– 0.638)		0.516 (0.479– 0.553)		0.654 (0.628– 0.681)	
	RF	0.838 (0.819– 0.857)		0.607 (0.575– 0.640)		0.466 (0.431– 0.502)		0.697 (0.669– 0.725)	
	XGB	0.832 (0.809– 0.855)		0.598 (0.563– 0.634)		0.470 (0.434– 0.505)		0.683 (0.655– 0.711)	
	MLP	0.823 (0.798– 0.848)		0.605 (0.568– 0.642)		0.512 (0.472– 0.551)		0.656 (0.623– 0.689)	
A + AR	KNN	0.727 (0.706– 0.748)		0.442 (0.407– 0.478)		0.336 (0.302– 0.370)		0.489 (0.444– 0.534)	
	LR	0.812 (0.791– 0.834)		0.591 (0.557– 0.624)		0.496 (0.459– 0.532)		0.654 (0.615– 0.692)	
	RF	0.844 (0.826– 0.862)		0.624 (0.594– 0.654)		0.481 (0.448– 0.515)		0.695 (0.668– 0.722)	
	XGB	0.798 (0.775– 0.821)		0.540 (0.506– 0.573)		0.421 (0.386– 0.455)		0.665 (0.627– 0.703)	
	MLP	0.813 (0.791– 0.835)		0.570 (0.532– 0.608)		0.481 (0.442– 0.521)		0.629 (0.590– 0.668)	
A + CR	KNN	0.819 (0.803– 0.835)		0.574 (0.547– 0.601)		0.412 (0.384– 0.441)		0.565 (0.531– 0.600)	
	LR	0.847 (0.834– 0.859)		0.607 (0.587– 0.626)		0.502 (0.478– 0.526)		0.707 (0.679– 0.734)	
	RF	0.868 (0.857– 0.880)		0.660 (0.637– 0.683)		0.549 (0.523– 0.574)		0.724 (0.698– 0.750)	
	XGB	0.859 (0.845– 0.872)		0.656 (0.629– 0.683)		0.534 (0.504– 0.565)		0.714 (0.685– 0.743)	
	MLP	0.860 (0.845– 0.875)		0.647 (0.621– 0.673)		0.547 (0.516– 0.577)		0.684 (0.663– 0.704)	
	CNN	0.870 (0.858– 0.882)		0.649 (0.627– 0.671)		0.529 (0.503– 0.555)		0.661 (0.637– 0.686)	
A + CR + UR	KNN	0.822 (0.807– 0.836)		0.584 (0.558– 0.610)		0.420 (0.393– 0.447)		0.594 (0.563– 0.625)	
	LR	0.847 (0.835– 0.860)		0.608 (0.588– 0.628)		0.502 (0.478– 0.526)		0.708 (0.680– 0.736)	
	RF	0.872 (0.862– 0.882)		0.675 (0.652– 0.699)		0.553 (0.525– 0.581)		0.712 (0.688– 0.736)	
	XGB	0.860 (0.847– 0.874)		0.666 (0.640– 0.692)		0.546 (0.517– 0.575)		0.715 (0.685– 0.744)	
	MLP	0.866 (0.850– 0.881)		0.657 (0.632– 0.682)		0.554 (0.527– 0.582)		0.712 (0.694– 0.731)	
	CNN	0.872 (0.861– 0.884)		0.656 (0.634– 0.678)		0.535 (0.510– 0.560)		0.676 (0.646– 0.707)	
	URGDAN	0.894 (0.883– 0.906)		0.700 (0.675– 0.724)		0.608 (0.582– 0.635)		0.746 (0.728– 0.764)	

Такі результати чітко вказують на те, що реакції користувачів можуть не бути прямим інформаційним сигналом, який відображає фінансовий ризик компанії. Однак результати також показують, що реакції користувачів можуть бути ефективними як керівне знання для кращого вилучення семантичних ознак, як це зроблено в URGDAN. Використання підходу уваги, керованого

реакціями користувачів, сприяє кращій взаємодії між ознаками реакцій користувачів та семантичними ознаками у поточних звітах.

Протестувано статистичну значущість порівнянь між URGDAN та еталонними методами, використовуючи непараметричний тест Фрідмана [45].

Таблиця 3.6 підсумовує результати повних попарних порівнянь семи методів. Оскільки тест Фрідмана є різновидом рангового тесту, результати за метриками продуктивності (тобто AUC, KS, Н-міра та показник recall) були об'єднані. Загалом, відмінності між сімома методами є статистично значущими ($\chi^2 = 309.526$, $p < 0.001$). Додаткові попарні порівняння показують, що URGDAN значно перевершує всі еталонні методи.

Таблиця 3.6 - Результати повного попарного порівняння

Метод	Середній ранг	p-значення попарного порівняння, скориговане за методом Бонферроні
KNN	5.00	
LR	3.72	<0.001
RF	2.53	<0.001
XGB	2.78	<0.001
MLP	2.67	<0.001
CNN	2.81	<0.001
URGDAN	1.49	<0.001
Friedman χ^2	309.526 (p < 0.001)	

Результати аналізу ефективності URGDAN у порівнянні з еталонними методами прогнозування чітко демонструють перевагу запропонованого підходу за всіма ключовими метриками продуктивності. URGDAN перевершує інші методи завдяки інтеграції семантичних ознак з поточних звітів та використанню механізму уваги, керованого реакціями користувачів, що забезпечує більш глибоке розуміння взаємозв'язків між ознаками. Хоча реакції користувачів самі по собі не є прямим сигналом фінансового ризику, вони слугують ефективним інструментом для оптимізації вилучення семантичної інформації. Статистично значущі відмінності, підтверджені тестом Фрідмана, свідчать про високу ефективність URGDAN у прогнозуванні фінансових труднощів, що підкреслює її перевагу над традиційними еталонними методами.

Це доводить, що запропонований підхід має значний потенціал для застосування в реальних задачах прогнозування фінансових ризиків.

3.4 Аналіз впливу реакцій користувачів на ефективність прогнозування фінансових труднощів

Після підтвердження того, що семантичні ознаки, витягнуті з поточних звітів за допомогою нашої запропонованої моделі, можуть значно покращити продуктивність прогнозування, досліджено продуктивність прогнозування при одночасному використанні фінансових даних, поточних звітів і різних типів реакцій користувачів у URGDAN. Перевірено, чи можуть ознаки, витягнуті за допомогою механізму уваги, керованого реакціями користувачів, надати цінну інформацію для прогнозування фінансових труднощів (FDP). Розгляд лайків та коментарів як активні реакції, а перегляди — як пасивну реакцію для спрощення експериментального порівняльного аналізу [46]. Як показано в таблиці 3.7, порівняно експериментальні результати шести наборів ознак, використовуючи різні реакції користувачів.

Таблиця 3.7 - Опис набору ознак, використаного в кожному методі

Ознаки	Набір ознак 1	Набір ознак 2	Набір ознак 3	Набір ознак 4	Набір ознак 5	Набір ознак 6
Облікові ознаки	√	√	√	√	√	√
Семантичні ознаки у поточних звітах	√	√	√	√	√	√
Ознаки реакцій користувачів						
Пасивні реакції	Перегляди	×	√	×	×	×
Активні реакції	Лайки	×	×	√	×	√
Коментарі	×	×	×	√	√	√

Таблиця 3.8 узагальнює результати, які показують, що незалежно від того, який підмножину реакцій користувачів використовують, методи, що поєднують облікові ознаки з семантичними ознаками, витягнутими за допомогою механізму уваги, керованого реакціями користувачів, завжди демонструють

кращу продуктивність прогнозування порівняно з методом без реакцій користувачів (тобто набір ознак 1). Це свідчить про перевагу використання механізму уваги, керованого реакціями користувачів, у прогнозуванні фінансових труднощів (FDP).

Таблиця 3.8 - Продуктивність прогнозування з використанням різних даних про реакції користувачів

Набір ознак	AUC	KS	H-mira	Recall
Набір ознак 1	0.876 (0.861–0.891)	0.681 (0.655–0.707)	0.585 (0.557–0.613)	0.703 (0.671–0.736)
Набір ознак 2	0.886 (0.874–0.897)	0.687 (0.661–0.713)	0.596 (0.567–0.626)	0.713 (0.683–0.744)
Набір ознак 3	0.887 (0.874–0.899)	0.689 (0.665–0.714)	0.598 (0.571–0.625)	0.716 (0.684–0.748)
Набір ознак 4	0.889 (0.877–0.900)	0.690 (0.665–0.716)	0.599 (0.571–0.627)	0.720 (0.688–0.753)
Набір ознак 5	0.890 (0.878–0.902)	0.696 (0.671–0.720)	0.601 (0.575–0.628)	0.721 (0.692–0.749)
Набір ознак 6	0.894 (0.883–0.906)	0.700 (0.675–0.724)	0.608 (0.582–0.635)	0.746 (0.728–0.764)

Найкращі результати за кожною метрикою для кожного методу виділено жирним шрифтом; 95% довірчий інтервал вказаний у дужках.

Коли для керування вилученням ознак використовується лише один тип реакцій користувачів, експериментальні результати показують незначне покращення за показниками продуктивності. Серед трьох типів реакцій користувачів використання коментарів для керування вилученням ознак призвело до найбільш значного підвищення продуктивності моделі.

Вважаємо, що коментарі краще відображають занепокоєння користувачів. Порівняно з використанням лише пасивних реакцій (тобто переглядів) для керування, помітніше покращення деяких метрик продуктивності спостерігається при використанні активних реакцій (тобто лайків і коментарів) для вилучення ознак. На відміну від пасивних реакцій, активні реакції краще відображають інтереси користувачів до поточного звіту, допомагаючи ідентифікувати текстову інформацію, пов'язану з фінансовими труднощами. Не дивно, що URGDAN демонструє найкращу продуктивність

при поєднанні активних і пасивних реакцій, що свідчить про те, що всі типи ознак реакцій користувачів надають цінну інформацію для прогнозування фінансових труднощів (FDP).

На завершення, URGDAN враховує як узгодженість між поточними звітами та інформацією про реакції користувачів, так і унікальні дані, що надаються реакціями користувачів. URGDAN може повною мірою використовувати взаємодію між текстом поточних звітів і реакціями користувачів. Його вища продуктивність у порівнянні з іншими методами підтверджує ефективність використання реакцій користувачів для керування вилученням семантичних ознак у прогнозуванні фінансових труднощів.

3.5 Кейс-дослідження

Далі представлено, як запропонований метод обчислює ваги уваги для кожного поточного звіту конкретної компанії, щоб виділити важливі елементи в завданні прогнозування фінансових труднощів (FDP). Наш метод ідентифікує поточні звіти, що пов'язані з фінансовими труднощами, надаючи інтерпретовану інформацію для кредиторів та інвесторів, яку можна використовувати для інвестицій, фінансування або управління бізнесом.

Щоб візуально проаналізувати роботу механізму уваги, керованого реакціями користувачів, ми витягли відповідні ваги уваги з поточних звітів компанії (Компанія А), щоб виділити ті звіти, які є важливими для прогнозування фінансових труднощів (див.рисунок 2.2). Зазначимо, що різні кольори кожного звіту представляють різні ваги уваги. Як показано на рисунку 3.1, наш механізм уваги успішно підкреслив семантичний зміст, який, ймовірно, важливий для діяльності компанії або може впливати на прогнозування фінансових труднощів. Наприклад, до звітів належать оголошення про зниження операційних показників, постійні збитки та недостатній оборотний капітал; повідомлення про зміну голови правління, вищого керівництва та

голови наглядової ради; а також повідомлення про судове замороження акцій. Текстова семантика поточних звітів є важливим фактором, який впливає на механізм уваги. Ці спостереження свідчать про те, що наш механізм уваги може виявляти ключові елементи та ігнорувати неактуальні компоненти під час аналізу поточних звітів.

Вплив поточного звіту на прогнозування фінансових труднощів не тільки залежить від семантики тексту, але й значно корелює з рівнем реакцій користувачів, що також видно на рисунку. Завдяки використанню цієї інформації про реакції користувачів, URGDAN має перевагу в оцінці важливості поточних звітів для прогнозування фінансових труднощів.

No	Current Reports	Date	Readings	Likes	Comments
1	Notice on Judicial Freeze of Equity	2017-02-21	10088	22	8
2	Stock release announcement	2017-04-25	5908	13	3
3	Announcement of asset impairment loss	2017-04-27	6442	3	3
4	Special report on the use of the raised funds	2017-05-04	10276	4	2
5	2016 Annual General Meeting resolution announcement	2017-06-06	7305	0	0
6	Suggestive announcement regarding the deferred disclosure of the 2017 semi-annual report	2017-08-22	5668	6	14
7	The independent opinion of the independent director on the re-election of the Board of directors	2017-08-28	5550	21	10
8	Notice of the first Extraordinary General Meeting of 2017	2017-08-28	5418	9	3
9	Notice of change of directors and supervisors	2017-09-19	8270	23	13
10	XX Securities Co., Ltd. on XX Co., Ltd. operating performance decline, continuous loss and insufficient working capital risk suggestive announcement	2017-09-23	17505	39	23
11	Notice of change of interest	2017-09-25	6723	1	2

Рисунок 3.1 - Візуалізація ваги уваги

Окрім аналізу поточних звітів однієї компанії, ми також провели детальне порівняння між ST і non-ST компаніями, щоб продемонструвати, що URGDAN може інтуїтивно виявляти відмінності в поточних звітах, опублікованих цими двома типами компаній. Зокрема, ми випадково відібрали 50 ST і 50 non-ST компаній з експериментальної вибірки та витягли ваги уваги їхніх звітів, як ми робили раніше. Відповідно до типів звітів, описаних офіційними регламентами, ми розділили поточні звіти на різні типи, такі як корпоративне управління, операційна діяльність компанії, операції маркетмейкера, питання ризиків тощо. Нарешті, ми розрахували середню вагу уваги для кожного типу звітів, щоб відобразити важливість різних типів поточних звітів для прогнозування

фінансових труднощів (FDP) у ST і non-ST компаніях. Рисунок 3.2 показує середні ваги уваги для різних типів поточних звітів у non-ST і ST компаніях.

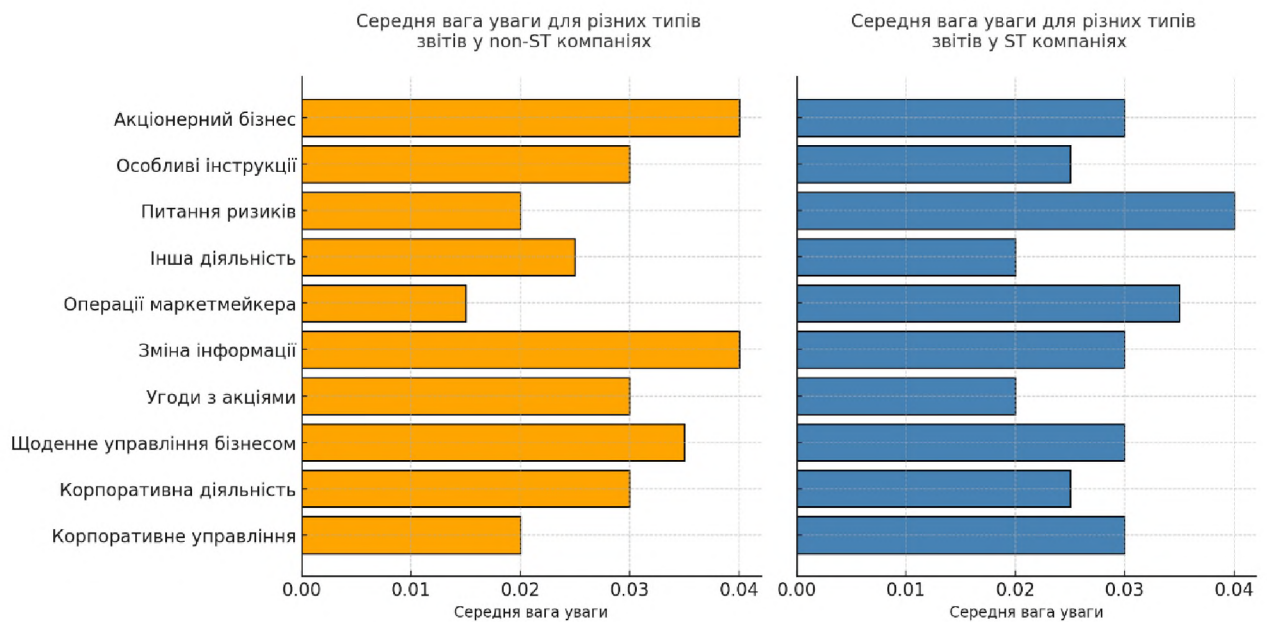


Рисунок 3.2 - Ваги уваги для різних типів поточних звітів

Результати показують як схожість, так і відмінності в розподілі типів звітів між ST і non-ST компаніями. Типи звітів, пов'язані з щоденним управлінням бізнесом і змінами інформації, є важливими як для ST, так і для non-ST компаній, але більше для non-ST компаній, де звіти про зміну інформації компанії мають найвищу вагу. У цих звітах зазвичай йдеться про зміни в складі директорів, наглядової ради, керівного персоналу та бізнес-діяльності компанії. Для ST компаній зростають ваги певних типів звітів, таких як операції маркетмейкера та питання ризиків, особливо у звітах про ризики. Тому ці типи поточних звітів є більш важливими при прогнозуванні, чи знаходиться компанія у фінансових труднощах.

Звіт про операції маркетмейкера стосується оголошень, які компанії видають щодо конкретних угод і договорів між ними та маркетмейкерами. За нашими спостереженнями, у звітах про операції маркетмейкера часто йдеться про відмову брокерських компаній (також відомих як маркетмейкери) від

надання котирувальних послуг для деяких ST компаній. Звіти, позначені як питання ризиків, включають попереджувальні звіти про адміністративні покарання або саморегулюючі заходи, а також про невиконання обов'язкових розкриттів, наприклад, нерозкриття річного звіту вчасно. Часте виникнення таких ключових подій у поточних звітах компанії свідчить про більшу ймовірність фінансових труднощів, і кредитори та інвестори компанії повинні бути пильними.

Таким чином, наш запропонований метод виявляє сигнали фінансових труднощів, які можуть допомогти пояснити проблеми компанії. Це дозволяє кредиторам та інвесторам коригувати стратегії кредитування чи інвестування, зменшувати втрати активів, пов'язані з фінансовими труднощами, та уникати ризиків.

Висновки до розділу 3

1. У результаті порівняльного аналізу було встановлено, що метод URGDAN перевершує всі інші еталонні методи за всіма основними метриками прогнозування фінансових труднощів (FDP). Особливо це стосується здатності точніше ідентифікувати компанії, що перебувають у фінансовій скруті, що підтверджено значними покращеннями у показниках AUC, KS та H-міра. Це демонструє важливість використання семантичних ознак поточних звітів разом із фінансовими показниками для підвищення точності прогнозів.

2. Дослідження показало, що реакції користувачів можуть значно впливати на прогнозування фінансових труднощів, особливо якщо використовувати активні реакції, такі як лайки та коментарі. Додавання ознак реакцій користувачів дозволяє URGDAN краще ідентифікувати важливі події у поточних звітах компаній. Це підтверджує, що взаємодія з користувачами надає додаткові індикатори для оцінки фінансових ризиків.

3. Механізм уваги, керований реакціями користувачів, продемонстрував здатність підвищити точність прогнозування за рахунок виділення важливих елементів у поточних звітах. Зокрема, аналіз ST-компаній виявив, що такі типи звітів, як операції маркетмейкера та питання ризиків, є ключовими для визначення фінансових труднощів. Це дозволяє не тільки прогнозувати фінансові ризики, але й надавати практичні рекомендації для кредиторів та інвесторів щодо управління інвестиціями.

ВИСНОВКИ

1. У ході цього дослідження було розроблено та оцінено метод URGDAN для прогнозування фінансових труднощів компаній (FDP) на основі глибокого навчання, яка інтегрує фінансові показники, текстові ознаки поточних звітів та реакції користувачів. Модель продемонструвала значні покращення у порівнянні з традиційними методами прогнозування, використовуючи новаторський підхід з механізмом уваги, керованим реакціями користувачів. У результаті експериментів було підтверджено, що URGDAN досягає високих значень метрик продуктивності: $AUC = 0.894$, $KS = 0.700$, $H\text{-міра} = 0.608$ та $Recall = 0.746$, що перевершує всі інші еталонні методи.

2. Відмінності між URGDAN та іншими методами прогнозування, такими як KNN, LR, RF, XGB, MLP та CNN, були статистично значущими, що підтверджено тестом Фрідмана ($\chi^2 = 309.526$, $p < 0.001$). Зокрема, методи, які не використовують реакції користувачів, показали нижчі показники AUC та Recall, що вказує на важливість інтеграції поведінкових даних користувачів. Для KNN, наприклад, $AUC = 0.753$, а $Recall = 0.517$, що значно поступається результатам URGDAN.

3. Особливо ефективним виявилось використання активних реакцій користувачів, таких як лайки та коментарі, для покращення результатів прогнозування. Додавання цих ознак дозволило підвищити показники AUC та KS у порівнянні з пасивними реакціями (переглядами). Наприклад, використання набору ознак, що включав лайки та коментарі (набір ознак 6), дало $AUC = 0.894$ та $KS = 0.700$, у той час як при використанні тільки фінансових даних ці показники становили $AUC = 0.876$ та $KS = 0.681$.

4. Аналіз компаній зі статусом ST та без нього також показав суттєві відмінності у розподілі важливості типів поточних звітів. Звіти про операції маркетмейкера та питання ризиків для ST-компаній отримали вищі ваги уваги, що підкреслює їхню важливість для прогнозування фінансових труднощів. Для

таких компаній ваги уваги для цих типів звітів були на 20% вищі, ніж у компаній без статусу ST.

5. Загалом, інтеграція фінансових показників, семантичних ознак з поточних звітів та реакцій користувачів продемонструвала свою ефективність у прогнозуванні фінансових труднощів. Модель URGDAN забезпечує гнучкий та адаптивний підхід до вилучення релевантної інформації та її обробки, що дозволяє надавати кредиторам та інвесторам більш точні прогнози щодо фінансових ризиків компаній. Ці результати також вказують на перспективність подальших досліджень щодо інтеграції поведінкових та фінансових даних для покращення фінансового аналізу.

6. У майбутніх дослідженнях можна зосередитися на розширенні наборів даних, використовуючи більше джерел фінансової та поведінкової інформації, а також на вдосконаленні механізмів вилучення та аналізу семантичних ознак, що потенційно може ще більше підвищити точність прогнозування.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Sun, J., Li, H., Fujita, H., Fu, B., & Ai, W. (2020). Class-imbalanced dynamic financial distress prediction based on Adaboost-SVM ensemble combined with SMOTE and time weighting. *Information Fusion*, 54, 128-144. <https://doi.org/10.1016/j.inffus.2019.07.004>
2. Zięba, M., Tomczak, S. K., & Tomczak, J. M. (2016). Ensemble boosted trees with synthetic features generation in application to bankruptcy prediction. *Expert Systems with Applications*, 58, 93-101. <https://doi.org/10.1016/j.eswa.2016.03.015>
3. Wang, G., Chen, G., Zhao, H., Zhang, F., Yang, S., & Lu, T. (2021). Leveraging multisource heterogeneous data for financial risk prediction: A novel hybrid-strategy-based self-adaptive method. *MIS Quarterly*, 45(4). <https://doi.org/10.25300/MISQ/2021/15439>
4. Mai, F., Tian, S., Lee, C., & Ma, L. (2019). Deep learning models for bankruptcy prediction using textual disclosures. *European Journal of Operational Research*, 274(2), 743-758. <https://doi.org/10.1016/j.ejor.2018.10.024>
5. Matin, R., Hansen, C., Hansen, C., & Mølgaard, P. (2019). Predicting distresses using deep learning of text segments in annual reports. *Expert Systems with Applications*, 132, 199-208. <https://doi.org/10.1016/j.eswa.2019.04.047>
6. Zhao, Q., Xu, W., & Ji, Y. (2023). Predicting financial distress of Chinese listed companies using machine learning: To what extent does textual disclosure matter? *International Review of Financial Analysis*, 102770. <https://doi.org/10.1016/j.irfa.2023.102770>
7. Jiang, C., Lyu, X., Yuan, Y., Wang, Z., & Ding, Y. (2022). Mining semantic features in current reports for financial distress prediction: Empirical evidence from unlisted public firms in China. *International Journal of Forecasting*, 38(3), 1086-1099. <https://doi.org/10.1016/j.ijforecast.2021.11.003>

8. Drake, M., Roulstone, D., & Thornock, J. (2015). The determinants and consequences of information acquisition via EDGAR. *Contemporary Accounting Research*, 92, 1128-1161. <https://doi.org/10.1111/1911-3846.12103>
9. Noh, S., So, E. C., & Weber, J. P. (2019). Voluntary and mandatory disclosures: Do managers view them as substitutes? *Journal of Accounting and Economics*, 68(1), 101243. <https://doi.org/10.1016/j.jacceco.2019.101243>
10. He, J., & Plumlee, M. A. (2020). Measuring disclosure using 8-K filings. *Review of Accounting Studies*, 25(3), 903-962. <https://doi.org/10.1007/s11142-019-09511-0>
11. Cheng, S. F., De Franco, G., Jiang, H., & Lin, P. (2019). Riding the blockchain mania: Public firms' speculative 8-K disclosures. *Management Science*, 65(12), 5901-5913. <https://doi.org/10.1287/mnsc.2018.3205>
12. Wang, X., Xiang, Z., Xu, W., & Yuan, P. (2022). The causal relationship between social media sentiment and stock return: Experimental evidence from an online message forum. *Economics Letters*, 216, 110598. <https://doi.org/10.1016/j.econlet.2022.110598>
13. Dong, W., Liao, S., & Zhang, Z. (2018). Leveraging financial social media data for corporate fraud detection. *Journal of Management Information Systems*, 35(2), 461-487. <https://doi.org/10.1080/07421222.2018.1451960>
14. Zell, A. L., & Moeller, L. (2018). Are you happy for me ... on Facebook? The potential importance of "likes" and comments. *Computers in Human Behavior*, 78, 26-33. <https://doi.org/10.1016/j.chb.2017.09.023>
15. Shahbaznezhad, H., Dolan, R., & Rashidirad, M. (2021). The role of social media content format and platform in users' engagement behavior. *Journal of Interactive Marketing*, 53, 47-65. <https://doi.org/10.1016/j.intmar.2020.05.001>
16. Craja, P., Kim, A., & Lessmann, S. (2020). Deep learning for detecting financial statement fraud. *Decision Support Systems*, 139, 113421. <https://doi.org/10.1016/j.dss.2020.113421>

- 17.Wang, G., Ma, J., & Chen, G. (2023). Attentive statement fraud detection: Distinguishing multimodal financial data with fine-grained attention. *Decision Support Systems*, 167, 113913. <https://doi.org/10.1016/j.dss.2023.113913>
- 18.Geng, R., Bose, I., & Chen, X. (2015). Prediction of financial distress: An empirical study of listed Chinese companies using data mining. *European Journal of Operational Research*, 241(1), 236-247. <https://doi.org/10.1016/j.ejor.2014.08.016>
- 19.Olson, D. L., Delen, D., & Meng, Y. (2012). Comparative analysis of data mining methods for bankruptcy prediction. *Decision Support Systems*, 52(2), 464-473. <https://doi.org/10.1016/j.dss.2011.10.007>
- 20.Du Jardin, P. (2015). Bankruptcy prediction using terminal failure processes. *European Journal of Operational Research*, 242(1), 286-303. <https://doi.org/10.1016/j.ejor.2014.09.003>
- 21.Liang, D., Lu, C. C., Tsai, C. F., & Shih, G. A. (2016). Financial ratios and corporate governance indicators in bankruptcy prediction: A comprehensive study. *European Journal of Operational Research*, 252(2), 561-572. <https://doi.org/10.1016/j.ejor.2016.01.012>
- 22.Campbell, J. Y., Hilscher, J., & Szilagyi, J. (2008). In search of distress risk. *Journal of Finance*, 63(6), 2899-2939. <https://doi.org/10.1111/j.1540-6261.2008.01416.x>
- 23.Wu, Y., Gaunt, C., & Gray, S. (2010). A comparison of alternative bankruptcy prediction models. *Journal of Contemporary Accounting & Economics*, 6(1), 34-45. <https://doi.org/10.1016/j.jcae.2010.04.002>
- 24.Kou, G., Xu, Y., Peng, Y., Shen, F., Chen, Y., Chang, K., & Kou, S. (2021). Bankruptcy prediction for SMEs using transactional data and two-stage multiobjective feature selection. *Decision Support Systems*, 140, Article 113421. <https://doi.org/10.1016/j.dss.2020.113421>

25. Tsai, M. F., & Wang, C. J. (2017). On the risk prediction and analysis of soft information in finance reports. *European Journal of Operational Research*, 257(1), 243-250. <https://doi.org/10.1016/j.ejor.2016.06.062>
26. Jiang, C., Zhou, Y., & Chen, B. (2023). Mining semantic features in patent text for financial distress prediction. *Technological Forecasting and Social Change*, 190, Article 122450. <https://doi.org/10.1016/j.techfore.2023.122450>
27. Li, S., Shi, W., Wang, J., & Zhou, H. (2021). A deep learning-based approach to constructing a domain sentiment lexicon: A case study in financial distress prediction. *Information Processing & Management*, 58(5), Article 102673. <https://doi.org/10.1016/j.ipm.2021.102673>
28. Loughran, T., & McDonald, B. (2011). When is a liability not a liability? Textual analysis, dictionaries, and 10-Ks. *Journal of Finance*, 66(1), 35-65. <https://doi.org/10.1111/j.1540-6261.2010.01625.x>
29. Yuan, X., Hou, F., & Cai, X. (2021). How do patent assets affect firm performance? From the perspective of industrial difference. *Technology Analysis & Strategic Management*, 33(8), 943-956. <https://doi.org/10.1080/09537325.2021.1886334>
30. Kraus, M., & Feuerriegel, S. (2017). Decision support from financial disclosures with deep neural networks and transfer learning. *Decision Support Systems*, 104, 38-48. <https://doi.org/10.1016/j.dss.2017.06.004>
31. Chen, H. L., Yang, B., Wang, G., Liu, J., Xu, X., Wang, S. J., & Liu, D. Y. (2011). A novel bankruptcy prediction model based on an adaptive fuzzy k-nearest neighbor method. *Knowledge-Based Systems*, 24(8), 1348-1359. <https://doi.org/10.1016/j.knosys.2011.06.005>
32. Cho, S., Hong, H., & Ha, B. C. (2010). A hybrid approach based on the combination of variable selection using decision trees and case-based reasoning using the Mahalanobis distance: For bankruptcy prediction. *Expert Systems with Applications*, 37(4), 3482-3488. <https://doi.org/10.1016/j.eswa.2009.10.006>

33. Mselmi, N., Lahiani, A., & Hamza, T. (2017). Financial distress prediction: The case of French small and medium-sized firms. *International Review of Financial Analysis*, 50, 67-80. <https://doi.org/10.1016/j.irfa.2017.02.004>
34. Petropoulos, A., Siakoulis, V., Stavrulakis, E., & Vlachogiannakis, N. E. (2020). Predicting bank insolvencies using machine learning techniques. *International Journal of Forecasting*, 36(3), 1092-1113. <https://doi.org/10.1016/j.ijforecast.2019.11.001>
35. Hosaka, T. (2019). Bankruptcy prediction using imaged financial ratios and convolutional neural networks. *Expert Systems with Applications*, 117, 287-299. <https://doi.org/10.1016/j.eswa.2018.09.039>
36. Elhoseny, M., Metawa, N., Sztano, G., & El-Hasnony, I. M. (2022). Deep learning-based model for financial distress prediction. *Annals of Operations Research*. <https://doi.org/10.1007/s10479-022-04787-8>
37. Borchert, P., Coussement, K., De Caigny, A., & De Weerd, J. (2023). Extending business failure prediction models with textual website content using deep learning. *European Journal of Operational Research*, 306(1), 348-357. <https://doi.org/10.1016/j.ejor.2023.04.007>
38. Yang, Y., Uy, M. C. S., & Huang, A. (2020). Finbert: A pretrained language model for financial communications. *arXiv preprint arXiv:2006.08097*. <https://arxiv.org/abs/2006.08097>
39. Oh, C., Roumani, Y., Nwankpa, J. K., & Hu, H. F. (2017). Beyond likes and tweets: Consumer engagement behavior and movie box office in social media. *Information & Management*, 54(1), 25-37. <https://doi.org/10.1016/j.im.2016.03.004>
40. Luong, M. T., Pham, H., & Manning, C. D. (2015). Effective approaches to attention-based neural machine translation. *arXiv preprint arXiv:1508.04025*. <https://arxiv.org/abs/1508.04025>

41. Lin, T. Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (2017). Focal loss for dense object detection. In *Proceedings of the IEEE International Conference on Computer Vision* (pp. 2980-2988). <https://doi.org/10.1109/ICCV.2017.324>
42. Zhou, F., Fu, L., Li, Z., & Xu, J. (2022). The recurrence of financial distress: A survival analysis. *International Journal of Forecasting*, 38(3), 1100-1115. <https://doi.org/10.1016/j.ijforecast.2021.11.002>
43. Dwarampudi, M., & Reddy, N. V. (2019). Effects of padding on LSTMs and CNNs. *arXiv preprint arXiv:1903.07288*. <https://arxiv.org/abs/1903.07288>
44. Hand, D. J. (2009). Measuring classifier performance: A coherent alternative to the area under the ROC curve. *Machine Learning*, 77(1), 103-123. <https://doi.org/10.1007/s10994-009-5119-5>
45. Wang, Z., Jiang, C., Zhao, H., & Ding, Y. (2020). Mining semantic soft factors for credit risk evaluation in peer-to-peer lending. *Journal of Management Information Systems*, 37(1), 282-308. <https://doi.org/10.1080/07421222.2019.1705509>
46. Alhabash, S., McAlister, A. R., Hagerstrom, A., Quilliam, E. T., Rifon, N. J., & Richards, J. I. (2013). Between likes and shares: Effects of emotional appeal and virality on the persuasiveness of anti-cyberbullying messages on Facebook. *Cyberpsychology, Behavior, and Social Networking*, 16(3), 175-182. <https://doi.org/10.1089/cyber.2012.0265>
47. Загальні методичні рекомендації з підготовки, оформлення, захисту та оцінювання кваліфікаційних робіт здобувачів вищої освіти першого (бакалаврського) і другого (магістерського) рівнів. Тернопіль: ЗУНУ, 2024. 83 с.
48. Комар М.П., Саченко А.О., Васильків Н.М., Загородня Д.І. Методичні рекомендації до виконання кваліфікаційної роботи з освітньо-професійної програми «Комп'ютерні науки» спеціальності 122 «Комп'ютерні науки» за другим (магістерським) рівнем вищої освіти. – Тернопіль: ЗУНУ, 2024. – 32 с.

Додаток А
Апробація отриманих результатів

ЗБІРНИК НАУКОВИХ ПРАЦЬ

З МАТЕРІАЛАМИ V МІЖНАРОДНОЇ НАУКОВОЇ КОНФЕРЕНЦІЇ

11 ЖОВТНЯ 2024 РІК

М. ЛУЦЬК, УКРАЇНА

**«СТРАТЕГІЧНІ НАПРЯМКИ РОЗВИТКУ НАУКИ:
ФАКТОРИ ВПЛИВУ ТА ВЗАЄМОДІЇ»**



УДК 082:001
С 83



Організація, від імені якої випущено видання:

ГО «Міжнародний центр наукових досліджень»

Номер запису організації в Єдиному реєстрі громадських об'єднань: 1499141.

Голова оргкомітету: Сотник С.Г.

Верстка: Білоус Т.В.

Дизайн: Бондаренко І.В.

Рекомендовано до видання Вченою Радою Інституту науково-технічної інтеграції та співпраці. Протокол № 57 від 10.10.2024 року.



Конференцію зареєстровано Державною науковою установою у сфері управління Міністерства освіти і науки «Український інститут науково-технічної експертизи та інформації» в базі даних науково-технічних заходів України на поточний рік та бюлетені «План проведення наукових, науково-технічних заходів в Україні» (**Посвідчення № 351 від 12.06.2024**).

Збірник наукових праць з матеріалами конференції видано офіційно суб'єктом видавничої справи зі **Свідоцтвом ДК № 7860 від 22.06.2023**.

Матеріали конференції знаходяться у відкритому доступі на умовах ліцензії Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

Стратегічні напрямки розвитку науки: фактори впливу та взаємодії: збірник наукових праць з матеріалами V Міжнародної наукової конференції, м. Луцьк, 11 жовтня, 2024 р. / Міжнародний центр наукових досліджень. — Вінниця: ТОВ «УКРЛОГОС Груп, 2024. — 254 с.

ISBN 978-617-8440-16-9

DOI 10.62731/mcnd-11.10.2024

Викладено матеріали учасників V Міжнародної наукової конференції «Стратегічні напрямки розвитку науки: фактори впливу та взаємодії», яка відбулася 11 жовтня 2024 року у місті Луцьк.

УДК 082:001

© Колектив учасників конференції, 2024

© ГО «Міжнародний центр наукових досліджень», 2024

ISBN 978-617-8440-16-9

© ТОВ «УКРЛОГОС Груп», 2024

СЕКЦІЯ XI. ХІМІЯ, ХІМІЧНА ТА БІОІНЖЕНЕРІЯ

ЗАСТОСУВАННЯ НЕОРГАНІЧНИХ СПЛУК ДЛЯ НАДАННЯ ВОЛОКНИСТИМ
МАТЕРІАЛАМ БАКТЕРИЦИДНИХ ВЛАСТИВОСТЕЙ У МЕДИЧНИХ І
ПРОФІЛАКТИЧНИХ ЦІЛЯХ

Качківський В.В., Попович Т.А. 130

СЕКЦІЯ XII. ЕЛЕКТРОНІКА ТА ТЕЛЕКОМУНІКАЦІЇ

ЗАСТОСУВАННЯ ТЕХНОЛОГІЇ БЕЗДРОТОВОЇ ЗАРЯДКИ ЗА ДОПОМОГОЮ РЕКТЕН
ДЛЯ ЖИВЛЕННЯ СИСТЕМ РЕЗ

Сокіркаєв Д.В. 135

СЕКЦІЯ XIII. ЕКОЛОГІЯ ТА ТЕХНОЛОГІЇ ЗАХИСТУ НАВКОЛИШНЬОГО СЕРЕДОВИЩА

ІМПЛЕМЕНТАЦІЯ ЄВРОПЕЙСЬКОГО ДОСВІДУ ЕКОЛОГІЧНОГО РЕГУЛЮВАННЯ
ВИКОРИСТАННЯ ТЕРИТОРІЙ ПРИРОДНО-ЗАПОВІДНОГО ФОНДУ

Голян В.М., Іванців В.В. 138

СЕКЦІЯ XIV. КОМП'ЮТЕРНА ТА ПРОГРАМНА ІНЖЕНЕРІЯ

HARDWARE SUPPORT OF INFORMATION SYSTEMS

Pavlyk H.V. 142

СЕКЦІЯ XV. ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА СИСТЕМИ

ПОБУДОВА ГНУЧКИХ СТРАТЕГІЙ УПРАВЛІННЯ ЗМІНАМИ В ML-ПРОЄКТАХ

Науково-дослідна група:

Висоцький А.В., Левандівський Н.Б., Вівюрка Н.М., Сулима Б.Я. 151

ПРИКЛАДНІ АСПЕКТИ ВИКОРИСТАННЯ МАШИННОГО НАВЧАННЯ ДЛЯ ОБРОБКИ
ТЕКСТОВИХ ДАНИХ

Науково-дослідна група:

Колівушко Е., Лучка С., Кондратюк Г., Кравчук Б., Тарасюк С. 154

ПРИКЛАДНІ АСПЕКТИ ВИКОРИСТАННЯ МАШИННОГО НАВЧАННЯ ДЛЯ ОБРОБКИ ТЕКСТОВИХ ДАНИХ

НАУКОВО-ДОСЛІДНА ГРУПА:

Колівушко Едуард

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Лучка Святослав

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Кондратюк Георгій

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Кравчук Богдан

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Тарасюк Софія

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Науковий керівник: Ліп'яніна-Гончаренко Христина

канд. техн. наук, доцент,
доцент кафедри інформаційно-обчислювальних систем та управління
Західноукраїнський національний університет, Україна

Машинне навчання (ML) кардинально змінило підходи до обробки текстових даних, дозволяючи автоматизувати складні завдання, такі як класифікація тексту, аналіз емоцій та виявлення фейкових новин. У зв'язку зі зростанням обсягів текстової інформації в Інтернеті, особливо на соціальних платформах, машинне навчання стало ключовим інструментом для швидкої і точної обробки цих даних. За статистикою, кожного дня публікується понад 2,5 мільярда текстових записів на платформах соціальних мереж [5], що робить застосування машинного навчання надзвичайно актуальним.

Існує декілька основних методів машинного навчання, які активно використовуються для обробки тексту. Серед них найпопулярніші — наївний Баєс, метод опорних векторів (SVM), логістична регресія та дерева рішень. Проте з розвитком глибинного навчання в останні роки стали особливо популярними рекурентні нейронні мережі (RNN) та трансформери (наприклад, BERT), що забезпечують високий рівень точності в завданнях класифікації тексту і прогнозуванні наступних слів у реченнях [3].

Попередня обробка тексту є критично важливою для підготовки даних до аналізу. Найбільш ефективними методами є стемінг, лематизація та видалення стоп-слів. Стемінг перетворює слова до їх основних форм, що допомагає зменшити варіативність даних, тоді як лематизація враховує контекст і надає точніші результати. Видалення стоп-слів допомагає зменшити розмір тексту, видаляючи непотрібні слова, які не несуть суттєвої інформації для моделі [4].

Якість попередньої обробки тексту безпосередньо впливає на результати моделей машинного навчання. Якщо дані погано підготовлені, моделі можуть працювати неефективно через надмірну кількість шуму або неправильні структури даних. Наприклад, неправильне видалення стоп-слів або відсутність лематизації може призвести до того, що модель не зможе точно розпізнати ключові патерни тексту, що знижує точність класифікації [2]. Таким чином, ретельна обробка тексту є важливим етапом у процесі машинного навчання.

Векторизація тексту — це процес перетворення текстових даних у числові представлення, які можуть бути використані моделями машинного навчання. Найпоширеніші підходи включають Bag of Words (BoW), TF-IDF і word embeddings (наприклад, Word2Vec, GloVe). BoW і TF-IDF є простими методами, що використовуються для побудови векторів на основі частоти слів. Однак word embeddings дозволяють моделі навчитися семантичним зв'язкам між словами, що значно покращує якість аналізу тексту [3].

Вибір техніки векторизації залежить від характеру задачі. Якщо основний фокус на простій класифікації тексту, такі методи, як TF-IDF або Bag of Words, можуть бути достатніми. Однак для задач, що вимагають врахування контексту та семантики (наприклад, аналіз емоцій або прогнозування наступних слів), більш ефективними будуть word

embeddings або трансформери на основі BERT. Для задач, де точність є критично важливою, варто віддавати перевагу більш складним підходам [1].

Аналіз емоцій є одним з основних застосувань машинного навчання для обробки тексту, особливо у сфері маркетингу, соціальних мереж та підтримки клієнтів. Алгоритми машинного навчання, такі як нейронні мережі та методи класифікації, можуть використовуватися для визначення тональності тексту, що дозволяє автоматично аналізувати настрої користувачів. Успішні кейси використання цих методів вже застосовуються у великих компаніях, таких як Amazon і Google, для аналізу зворотного зв'язку клієнтів [2].

У результаті дослідження виявлено, що машинне навчання є ефективним інструментом для обробки текстових даних, особливо при правильній попередній обробці та виборі підходящої техніки векторизації. Для більш складних завдань, таких як аналіз емоцій або семантична класифікація, рекомендується використовувати сучасні моделі глибокого навчання, що базуються на embeddings або трансформерах. Подальші дослідження можуть зосередитися на оптимізації процесу векторизації та покращенні якості моделей на багатомовних даних.

Список використаних джерел:

1. Vaswani, A., Shazeer, N., Parmar, N., et al. "Attention Is All You Need." *Advances in Neural Information Processing Systems*, 2017.
2. Pang, B., & Lee, L. "Opinion Mining and Sentiment Analysis." *Foundations and Trends in Information Retrieval*, 2008.
3. Mikolov, T., Chen, K., Corrado, G., & Dean, J. "Efficient Estimation of Word Representations in Vector Space." *arXiv preprint arXiv:1301.3781*, 2013.
4. Manning, C., Raghavan, P., & Schütze, H. "Introduction to Information Retrieval." *Cambridge University Press*, 2008.
5. Statista. "Number of social media users worldwide from 2010 to 2021." [Online]. Available: www.statista.com.

ЗБІРНИК НАУКОВИХ ПРАЦЬ

З МАТЕРІАЛАМИ III МІЖНАРОДНОЇ НАУКОВОЇ КОНФЕРЕНЦІЇ

20 ВЕРЕСНЯ 2024 РІК

М. ОДЕСА, УКРАЇНА

**«ІНТЕЛЕКТУАЛЬНИЙ РЕСУРС СЬОГОДЕННЯ:
НАУКОВІ ЗАДАЧІ, РОЗВИТОК ТА ЗАПИТАННЯ»**



УДК 082:001
I-57



Організація, від імені якої випущено видання:

ГО «Міжнародний центр наукових досліджень»

Номер запису організації в Єдиному реєстрі громадських об'єднань: 1499141.

Голова оргкомітету: Сотник С.Г.

Верстка: Бабич Ю.В.

Дизайн: Бондаренко І.В.

Рекомендовано до видання Вченою Радою Інституту науково-технічної інтеграції та співпраці. Протокол № 54 від 19.09.2024 року.



Конференцію зареєстровано Державною науковою установою у сфері управління Міністерства освіти і науки «Український інститут науково-технічної експертизи та інформації» в базі даних науково-технічних заходів України на поточний рік та бюлетені «План проведення наукових, науково-технічних заходів в Україні» (**Посвідчення № 348 від 12.06.2024**).

Збірник наукових праць з матеріалами конференції видано офіційно суб'єктом видавничої справи зі **Свідоцтвом ДК № 7860 від 22.06.2023**.

Матеріали конференції знаходяться у відкритому доступі на умовах ліцензії Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0).

I-57 **Інтелектуальний ресурс сьогодення: наукові задачі, розвиток та запитання:** збірник наукових праць з матеріалами III Міжнародної наукової конференції, м. Одеса, 20 вересня, 2024 р. / Міжнародний центр наукових досліджень. — Вінниця: ТОВ «УКРЛОГОС Груп, 2024. — 322 с.

ISBN 978-617-8440-13-8

DOI 10.62731/mcnd-20.09.2024

Викладено матеріали учасників III Міжнародної наукової конференції «Інтелектуальний ресурс сьогодення: наукові задачі, розвиток та запитання», яка відбулася 20 вересня 2024 року у місті Одеса.

УДК 082:001

© Колектив учасників конференції, 2024

© ГО «Міжнародний центр наукових досліджень», 2024

ISBN 978-617-8440-13-8

© ТОВ «УКРЛОГОС Груп», 2024

**СЕКЦІЯ X.
ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА СИСТЕМИ**

ОГЛЯД КОНЦЕПЦІЙ ТА МЕТОДІВ МАШИННОГО НАВЧАННЯ: ТЕОРІЯ,
ЗАСТОСУВАННЯ ТА ВИКЛИКИ
Кондратюк Г., Кравчук Б., Тарасюк С., Матвійчук М. 177

**СЕКЦІЯ XI.
ТРАНСПОРТ ТА ТРАНСПОРТНІ ТЕХНОЛОГІЇ**

INNOVATIVE TECHNOLOGIES AND REGULATORY MEASURES TO REDUCE
ENVIRONMENTAL RISKS IN THE SHIPPING INDUSTRY
Sagaydak O., Kucherenko V., Kotenko O., Prokhorov V. 181

**СЕКЦІЯ XII.
ФІЗИКО-МАТЕМАТИЧНІ НАУКИ**

ПЕРЕХІД МЕТАЛ-ІЗОЛЯТОР У МОНОКРИСТАЛАХ $Y_{1-z}Pr_zBa_2Cu_3O_{7-\delta}$ З РІЗНИМ
ВМІСТОМ ПРАЗЕОДИМУ
Камчатна С. М., Ярчук Д. Ф., Чепурін О. Г., Іноземцев М. М., Вовк Р. В. 187

**СЕКЦІЯ XIII.
СОЦІОЛОГІЯ ТА СТАТИСТИКА**

СІМ'Я ЯК СОЦІАЛЬНИЙ ІНСТИТУТ. ФУНКЦІЇ СІМ'Ї
Олійник М. В. 191

**СЕКЦІЯ XIV.
ФІЛОЛОГІЯ ТА ЖУРНАЛІСТИКА**

FROM BEATS TO WORDS: HOW MUSIC NORMALIZES CREATIVE AND IMPLICIT
ENGLISH LANGUAGE
Hulei T. 196

ФУНКЦІОНУВАННЯ НЕФОРМАЛЬНОЇ ЛЕКСИКИ У СУЧАСНИХ ПІСНЯХ
Борисова К. М. 205

АВТОР І ТЕКСТ У РОМАНІ У. ЕКО «ІМ'Я ТРОЯНДИ»
Кулик К. В. 209

ЗАСОБИ ПЕРЕКЛАДУ РЕАЛІЙ В ІСТОРИЧНОМУ ДЕТЕКТИВІ (НА МАТЕРІАЛІ
РОМАНУ К. ДЖ. СЕНСОМА «DISSOLUTION»)
Марченко В. В. 212

СЕКЦІЯ Х. ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА СИСТЕМИ

ОГЛЯД КОНЦЕПЦІЙ ТА МЕТОДІВ МАШИННОГО НАВЧАННЯ: ТЕОРІЯ, ЗАСТОСУВАННЯ ТА ВИКЛИКИ

Кондратюк Георгій

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Кравчук Богдан

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Тарасюк Софія

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Матвійчук Микола

здобувач вищої освіти факультету комп'ютерних інформаційних технологій
Західноукраїнський національний університет, Україна

Науковий керівник: Ліп'яніна-Гончаренко Христина

канд. тех. наук, доцент, доцент кафедри інформаційно-обчислювальних
систем та управління
Західноукраїнський національний університет, Україна

Машинне навчання є однією з ключових технологій, що трансформують різні галузі, від медицини до інженерії та промисловості. З розвитком обчислювальної потужності та зростанням обсягів даних, методи машинного навчання стають все більш складними та ефективними. Однак, щоб повністю зрозуміти та застосувати ці методи, необхідно володіти глибоким знанням теоретичних концепцій та практичних підходів. У цьому огляді зібрано ключові публікації, які досліджують різні аспекти машинного навчання, починаючи від основ статистичного навчання і закінчуючи використанням машинного

навчання в різних галузях, таких як медицина, телекомунікації та автоматизація. Метою цього огляду є узагальнення та систематизація знань у цій сфері для кращого розуміння сучасних методів машинного навчання.

У публікації "Machine Learning Methods: An Overview" [1] автор R. Muhamedyev надає загальний огляд методів машинного навчання, їх класифікацію та застосування. Стаття розглядає різні категорії алгоритмів, включаючи навчання з учителем і без учителя, а також надає практичні приклади їх використання. Особлива увага приділяється аналізу переваг та обмежень різних підходів, що робить цю роботу корисною для початківців у цій сфері.

A.F.A.H. Alnuaimi та T.H.K. Albaldawi у своїй роботі [2] представляють фундаментальні концепції статистичного навчання та класифікації в машинному навчанні. Автори висвітлюють основні методи та алгоритми, зокрема регресію та класифікацію, і пояснюють їх роль у вирішенні складних завдань. Ця стаття акцентує увагу на важливості статистичних підходів у машинному навчанні та їх впливі на точність моделей.

У статті "Logical Approaches to Machine Learning" [3] автор P. Flach досліджує логічні підходи до машинного навчання. Ця публікація підкреслює важливість логіки та математичних основ для створення ефективних алгоритмів навчання. Вона розглядає різні типи логічних моделей та їхню здатність вирішувати складні завдання, забезпечуючи глибше розуміння теоретичних аспектів машинного навчання.

У публікації [4] автори S. Albahra та співавтори досліджують застосування методів машинного навчання та штучного інтелекту в патології та лабораторній медицині. Стаття надає огляд ключових алгоритмів та їх застосувань у медичному аналізі, демонструючи, як ці методи можуть покращити діагностику та підвищити ефективність медичних процедур.

H.M. El Misilmani та T. Naous у своїй роботі [5] розглядають використання машинного навчання в проектуванні антен. Вони описують різні концепції машинного навчання та алгоритми, які можуть бути застосовані для оптимізації антен. Стаття надає практичні приклади використання машинного навчання у галузі телекомунікацій, демонструючи його ефективність та потенціал для інженерних задач.

У публікації [6] автори D. Kreuzberger, N. Kühl, та S. Hirschl досліджують концепцію Machine Learning Operations (MLOps). Вони

надають визначення MLOps, пояснюють його архітектуру та важливість для ефективного розгортання моделей машинного навчання в промислових додатках. Ця стаття є ключовою для розуміння того, як інтегрувати машинне навчання у виробниче середовище.

J.G. Carbonell, R.S. Michalski, та T.M. Mitchell у своїй класичній роботі [7] надають фундаментальний огляд концепцій машинного навчання. Вони пропонують критерії для класифікації та порівняння методів машинного навчання, висвітлюючи різні аспекти, такі як представлення знань та алгоритми навчання. Ця стаття є однією з перших, яка систематизувала знання у цій галузі.

У публікації "An Overview of Machine Learning Techniques in Constraint Solving" [8] автори A. Popescu та інші розглядають, як машинне навчання може використовуватися для оптимізації процесів вирішення обмежень. Вони аналізують різні підходи та методи, зосереджуючись на використанні машинного навчання для автоматизації та покращення процесу вирішення задач з обмеженнями.

T.O. Ayodele у своїй праці [9] надає таксономічний аналіз машинного навчання, організовуючи його на основі різних підходів до навчання. Стаття розглядає різні категорії алгоритмів та підходів до машинного навчання, пропонуючи вичерпний огляд методів та їх застосувань у різних галузях.

У статті "An Overview of Privacy in Machine Learning" [10] автор E. De Cristofaro аналізує концепції конфіденційності в машинному навчанні. Ця робота досліджує, як забезпечити конфіденційність даних під час навчання моделей, і пропонує огляд методів, які спрямовані на захист даних. Стаття є важливим внеском у розуміння безпеки та етики в машинному навчанні.

Огляд літератури показує, що машинне навчання є багатогранною дисципліною, яка включає різноманітні методи, підходи та галузі застосування. Від статистичного та логічного навчання до спеціалізованих практичних застосувань, таких як медична діагностика або проєктування антен, машинне навчання демонструє значний потенціал для вирішення складних задач. При цьому важливим є не лише розвиток нових алгоритмів, але й розуміння таких аспектів, як конфіденційність даних, етика та інтеграція машинного навчання у виробничі середовища. Ці публікації допомагають розширити розуміння ключових концепцій машинного навчання та сприяють розвитку цієї

динамічної області, відкриваючи нові перспективи для подальших досліджень та інновацій.

Список використаних джерел:

1. R. Muhamedyev. Machine Learning Methods: An Overview. 2015. https://www.academia.edu/download/54752237/CMNT_2015_Machine_Learning_ns_87art02_ReviewPaper.pdf
2. A.F.A.H. Alnuaimi, T.H.K. Albaldawi. Concepts of Statistical Learning and Classification in Machine Learning: An Overview. 2024. https://www.bioconferences.org/articles/bioconf/pdf/2024/16/bioconf_iscku2024_00129.pdf
3. P. Flach. Logical Approaches to Machine Learning - An Overview. 1992. <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=03b54d20a9d1b68f01c959be0d1266076803f54d>
4. S. Albahra, T. Gorbett, S. Robertson, G. D'Aleo. Artificial Intelligence and Machine Learning Overview in Pathology & Laboratory Medicine. 2023. <https://www.sciencedirect.com/science/article/pii/S0740257023000138>
5. H.M. El Misilmani, T. Naous. Machine Learning in Antenna Design: An Overview on Machine Learning Concept and Algorithms. 2019. https://hilalelmisilmani.wordpress.com/wp-content/uploads/2019/10/hpcs_acme_2019_paper_4.pdf
6. D. Kreuzberger, N. Kühl, S. Hirschl. Machine Learning Operations (MLOps): Overview, Definition, and Architecture. 2023. <https://ieeexplore.ieee.org/iel7/6287639/6514899/10081336.pdf>
7. J.G. Carbonell, R.S. Michalski, T.M. Mitchell. An Overview of Machine Learning. 1983. <https://www.sciencedirect.com/science/article/pii/B9780080510545500054>
8. Popescu, S. Polat-Erdeniz, A. Felfernig, M. Uta. An Overview of Machine Learning Techniques in Constraint Solving. 2022. <https://link.springer.com/content/pdf/10.1007/s10844-021-00666-5.pdf>
9. T.O. Ayodele. Machine Learning Overview. 2010. <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=72bfa0a326e4385ac46916e840e9fdc73a98b9fb>
10. E. De Cristofaro. An Overview of Privacy in Machine Learning. 2020. <https://arxiv.org/pdf/2005.08679>