

**НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ
імені ІГОРЯ СІКОРСЬКОГО»**

Факультет прикладної математики

Кафедра програмного забезпечення комп'ютерних систем

«На правах рукопису»
УДК 004.912

«До захисту допущено»
Завідувач кафедри
_____ Євгенія СУЛЕМА
«__» _____ 2025 р.

Магістерська дисертація

на здобуття ступеня магістра

освітньо-науковою програмою

**«Інженерія програмного забезпечення мультимедійних та
інформаційно-пошукових систем»**

**зі спеціальності 121 Інженерія програмного забезпечення
на тему: «Спосіб і програмне забезпечення для визначення
потенційної популярності текстових постів в соціальних мережах»**

Виконав:

Студент II курсу, групи КП-41мп
Довженко Роман Олександрович _____

Керівник:

Доцент кафедри ПЗКС, к.т.н., доцент,
Заболотня Тетяна Миколаївна _____

Консультант з нормоконтролю:

Доцент кафедри ПЗКС, к.т.н., доцент,
Онай Микола Володимирович _____

Рецензент:

Доцент кафедри СПіСКС, к.т.н., доцент,
Тарасенко-Клятченко Оксана Володимирівна _____

Засвідчую, що у цій магістерській
дисертації немає запозичень з праць
інших авторів без відповідних
посилань.

Студент _____

Київ – 2025 року

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Факультет прикладної математики

Кафедра програмного забезпечення комп'ютерних систем

Рівень вищої освіти – другий (магістерський)

Спеціальність (спеціалізація) – 121 «Інженерія програмного забезпечення»

Освітньо-професійна програма «Інженерія програмного забезпечення мультимедійних та інформаційно-пошукових систем»

ЗАТВЕРДЖУЮ

Завідувач кафедри

_____ Євгенія СУЛЕМА

«__» _____ 2024 р.

ЗАВДАННЯ
на магістерську дисертацію студенту

Довженко Роману Олександровичу

1. Тема дисертації «Спосіб і програмне забезпечення для визначення потенційної популярності текстових постів в соціальних мережах», науковий керівник дисертації Заболотня Тетяна Миколаївна, к.т.н., доцент, затверджені наказом №4836-С по університету від «6» листопада 2025 р.
2. Термін подання студентом дисертації «19» грудня 2025 р.
3. Об'єкт дослідження: процес інтелектуального аналізу текстового контенту для оцінювання його потенційної популярності.
4. Предмет дослідження: способи, алгоритми та програмні засоби прогнозування популярності текстових публікацій в соціальних мережах.
5. Перелік завдань, які потрібно розробити:
 - провести аналіз способів прогнозування популярності текстового контенту в соціальних;
 - проаналізувати алгоритми машинного навчання, що застосовуються для задач класифікації та прогнозування популярності текстових публікацій;
 - сформулювати гіпотези щодо доцільності комбінованого використання текстових ознак і метаданих соціальних мереж;
 - розробити спосіб визначення потенційної популярності текстових постів у соціальних мережах;
 - виконати порівняльний аналіз отриманих результатів з існуючими підходами та сформулювати висновки щодо ефективності запропонованого способу.
6. Орієнтовний перелік графічного (ілюстративного) матеріалу:
 - графічне представлення узагальненої структури запропонованого способу;
 - графічне представлення t-sne структури векторних представлень BERT;
 - графічне представлення кореляційної матриці оброблених ознак;

- графічне представлення готового програмного рішення;
- графічне представлення бізнес-моделі проєкту у форматі Lean Canvas;
- графік точності способів на тестовому наборі даних.

7. Орієнтовний перелік публікацій:

- тези доповіді “Спосіб оцінювання потенційної популярності текстових публікацій в соціальних мережах”.

8. Консультанти розділів дисертації

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
Нормоконтроль	Онай М.В., доцент кафедри ПЗКС		

9. Дата видачі завдання «04» жовтня 2024 р.

Календарний план

№ з/п	Назва етапів виконання магістерської дисертації	Термін виконання етапів магістерської дисертації	Примітка
1.	Грунтовне ознайомлення з предметною галуззю	17.10.2024	
2.	Визначення структури магістерської дисертації; вивчення літератури, пошук додаткової літератури, патентний пошук	04.12.2024	
3.	Робота над першим розділом магістерської дисертації; проведення наукового дослідження	15.02.2025	
4.	Проведення наукового дослідження; робота над другим розділом магістерської дисертації; розроблення програмного забезпечення	05.04.2025	
5.	Проведення наукового дослідження; робота над третім розділом магістерської дисертації; розроблення програмного забезпечення	15.05.2025	
6.	Проведення наукового дослідження; робота над четвертим і п'ятим розділом магістерської дисертації	15.06.2025	
7.	Завершення роботи над основною частиною магістерської дисертації; підготовка ілюстративного матеріалу; підготовка матеріалів доповіді на конференції ПМК-2025	05.11.2025	
8.	Оформлення текстової і графічної частини магістерської дисертації	04.12.2025	

Студент

Роман ДОВЖЕНКО

Науковий керівник

Тетяна ЗАБОЛОТНЯ

РЕФЕРАТ

Актуальність теми. У сучасному цифровому середовищі соціальні мережі є одним із основних джерел генерації текстового контенту, обсяги якого щоденно зростають. Текстові публікації відіграють ключову роль у формуванні інформаційного простору, маркетингових стратегій, рекомендаційних систем та взаємодії між користувачами. У зв'язку з цим актуальною є задача автоматизованого визначення потенційної популярності текстових постів у соціальних мережах, що дозволяє прогнозувати рівень залученості аудиторії ще на етапі створення контенту. Популярність текстових публікацій залежить від багатьох чинників, серед яких семантичний зміст тексту, емоційне забарвлення, тематична спрямованість, час публікації, а також характеристики автора. Ручний аналіз таких багатовимірних даних є практично неможливим, що зумовлює необхідність застосування методів оброблення природної мови та машинного навчання.

Існуючі способи до прогнозування популярності контенту, зокрема статистичні методи та класичні алгоритми машинного навчання, мають обмежені можливості врахування глибоких семантичних зв'язків у тексті. З іншого боку, сучасні трансформерні моделі демонструють високу точність аналізу текстових даних, проте їх ефективне використання у задачах прогнозування популярності потребує інтеграції з аналізом метаданих та поведінкових характеристик користувачів. Це зумовлює необхідність розроблення комбінованих способів, що поєднують переваги різних способів.

Об'єктом дослідження є процес інтелектуального аналізу текстового контенту для оцінювання його потенційної популярності.

Предметом дослідження є способи, алгоритми та програмні засоби прогнозування популярності текстових публікацій у соціальних мережах.

Мета роботи: підвищення точності визначення потенційної популярності текстових постів у соціальних мережах шляхом розроблення способу та програмного забезпечення, що інтегрує методи оброблення природної мови, аналіз метаданих та сучасні трансформерні архітектури.

Методи дослідження. У роботі застосовано методи попереднього оброблення текстових даних, включаючи лінгвістичну нормалізацію, токенізацію та лематизацію, методи векторного подання тексту (TF-IDF, семантичні ембеддинги BERT), алгоритми класичного машинного навчання (логістична регресія, SVM, градієнтний бустинг), рекурентні нейронні мережі (LSTM), а також методи зниження розмірності та нормалізації ознак. Для аналізу результатів використано статистичні метрики якості класифікації та методи інтерпретації моделей SHAP і LIME.

Наукова новизна. Вперше запропоновано спосіб оцінювання потенційної популярності текстових публікацій у соціальних мережах, який відрізняється від існуючих поєднанням аналізу емоційного забарвлення тексту на основі трансформерних моделей, аналізу настрою та статистичного оброблення метаданих публікацій. Запропонований спосіб дозволяє враховувати як глибинні контекстуальні залежності тексту, так і часові та соціальні характеристики контенту, що забезпечує підвищення точності прогнозування популярності на 8.5% у порівнянні з базовими методами.

Практична значущість. Розроблено програмне забезпечення для визначення потенційної популярності текстових постів у соціальних мережах з використанням мови програмування Python та сучасних бібліотек машинного навчання і оброблення природної мови. Реалізований програмний прототип може бути використаний як окремий аналітичний інструмент або інтегрований у системи контент-менеджменту, маркетингової аналітики та рекомендаційні платформи. Результати експериментальних досліджень підтверджують ефективність запропонованого способу та його практичну придатність.

Апробація роботи. Основні положення та результати роботи були представлені та обговорювались на науковій конференції магістрантів та аспірантів «Прикладна математика та комп'ютинг ПМК-2025» (Київ, 19-21 листопада 2025 р.).

Структура та обсяг роботи. Магістерська дисертація складається зі вступу, п'яти розділів, висновків та додатків.

У вступі магістерської дисертації обґрунтовано актуальність задачі прогнозування потенційної популярності текстових публікацій у соціальних мережах в умовах зростання обсягів неструктурованого контенту та підвищення вимог до ефективності аналітичних систем. Сформульовано мету дослідження, яка полягає у розробленні та експериментальній перевірці способу автоматизованого прогнозування популярності текстових постів на основі методів оброблення природної мови та машинного навчання. Визначено об'єкт і предмет дослідження, а також окреслено сукупність методів, що включають статистичний аналіз, алгоритми машинного та глибинного навчання, трансформерні моделі та методи оцінювання якості прогнозування.

У першому розділі здійснено огляд сучасних способів до аналізу текстових даних і прогнозування популярності контенту у соціальних мережах. Розглянуто основні етапи оброблення тексту, методи формування векторних представлень, класичні алгоритми машинного навчання та сучасні моделі глибинного навчання, зокрема трансформерні архітектури. Проаналізовано переваги та обмеження існуючих способів, а також визначено ключові проблеми, пов'язані з контекстною залежністю текстів, дисбалансом даних і ресурсомісткістю моделей.

У другому розділі описано розроблений спосіб оцінювання потенційної популярності текстових публікацій у соціальних мережах та обґрунтовано його архітектуру. Запропонований спосіб базується на комбінованому використанні семантичних ознак, отриманих за допомогою трансформерних моделей, і статистичних характеристик метаданих

соціальних платформ. Розглянуто основні етапи формування ознак, механізми інтеграції різнорідних даних та принципи побудови моделі прогнозування.

У третьому розділі наведено реалізацію програмного забезпечення для прогнозування популярності текстового контенту та описано використані інструменти і технології. Обґрунтовано вибір мови програмування та бібліотек для оброблення текстових даних, навчання моделей машинного навчання й організації модульної архітектури системи. Розглянуто структуру програмного комплексу та взаємодію його основних компонентів.

У четвертому розділі представлено результати експериментальних досліджень ефективності запропонованого способу прогнозування популярності текстових публікацій. Проведено порівняльний аналіз отриманих результатів із базовими методами та існуючими способами на основі стандартних метрик якості класифікації. Показано, що запропонований спосіб забезпечує підвищення точності прогнозування та стабільність результатів на різних підмножинах даних.

У п'ятому розділі розглянуто питання бізнес-моделювання способу прогнозування потенційної популярності текстових публікацій у соціальних мережах. Проаналізовано проблему комерційного впровадження розробленого програмного рішення, визначено ключові виклики, пов'язані з монетизацією, масштабуванням та інтеграцією з існуючими цифровими платформами. Проведено ідентифікацію та аналіз зацікавлених сторін проєкту, що дозволило оцінити їхній вплив, рівень зацікавленості та пріоритетність у контексті реалізації стартапу.

У висновках узагальнено основні результати проведеного дослідження, сформульовано наукові та практичні висновки щодо ефективності запропонованого способу, а також визначено напрями подальших досліджень, зокрема розширення способу за рахунок

мультимодальних даних, адаптації моделей у реальному часі та персоналізації прогнозів.

У додатках наведено графічні матеріали до пояснювальної записки, лістинг коду розробленого програмного забезпечення та скріншоти презентації.

Робота виконана на 139 аркушах, містить 3 додатки та посилання на список використаних літературних джерел з 11 найменувань. У роботі наведено 6 рисунки, 8 таблиці.

Ключові слова: прогнозування популярності, соціальні мережі, обробка природної мови, трансформерні моделі, BERT, машинне навчання, аналіз текстових даних.

ABSTRACT

Theme urgency. In the modern digital environment, social networks are one of the main sources of textual content generation, the volume of which is constantly increasing. Textual posts play a key role in shaping the information space, marketing strategies, recommendation systems, and user interaction. In this context, the task of automated prediction of the potential popularity of textual posts in social networks becomes especially relevant, as it enables forecasting audience engagement at the content creation stage. The popularity of textual publications depends on multiple factors, including semantic content, emotional tone, thematic orientation, publication time, and author characteristics. Manual analysis of such multidimensional data is practically impossible, which necessitates the use of natural language processing and machine learning methods.

Existing approaches to content popularity prediction, including statistical methods and classical machine learning algorithms, have limited ability to capture deep semantic relationships within text. At the same time, modern transformer-based models demonstrate high accuracy in textual data analysis; however, their effective application to popularity prediction tasks requires integration with metadata analysis and user behavioral characteristics. This determines the need to develop combined approaches that integrate the advantages of different methods.

Object of research is the process of intelligent analysis of textual content for evaluating its potential popularity.

Subject of research is methods, algorithms, and software tools for predicting the popularity of textual publications in social networks.

Research objective is to increase the accuracy of predicting the potential popularity of textual posts in social networks by developing a method and software that integrate natural language processing techniques, metadata analysis, and modern transformer architectures.

Research methods. The study applies methods of textual data preprocessing, including linguistic normalization, tokenization, and

lemmatization; vector-based text representation methods (TF-IDF, BERT semantic embeddings); classical machine learning algorithms (logistic regression, SVM, gradient boosting); recurrent neural networks (LSTM); as well as feature normalization and dimensionality reduction techniques. Statistical classification metrics and model interpretability methods such as SHAP and LIME are used for result analysis.

Scientific novelty. For the first time, a method for evaluating the potential popularity of textual publications in social networks is proposed, which differs from existing approaches by the comprehensive integration of transformer-based semantic text analysis, sentiment analysis, and statistical processing of publication metadata. The proposed method enables capturing both deep contextual dependencies of text and temporal and social characteristics of content, thereby improving prediction accuracy at 8.5% compared to baseline methods.

Practical value. Software has been developed for predicting the potential popularity of textual posts in social networks using the Python programming language and modern machine learning and natural language processing libraries. The implemented software prototype can be used as a standalone analytical tool or integrated into content management systems, marketing analytics platforms, and recommendation systems. Experimental results confirm the effectiveness of the proposed method and its practical applicability.

Approbation. The main provisions and results of the research were presented and discussed at the scientific conference of master's and postgraduate students "Applied Mathematics and Computing PMC-2025" (Kyiv, November 19–21, 2025).

Structure and content of the thesis. The master's thesis consists of an introduction, five chapters, conclusions, and appendices.

The introduction substantiates the relevance of the problem of predicting the potential popularity of textual publications in social networks under conditions of growing volumes of unstructured content and increasing requirements for analytical system efficiency. The research objective is formulated, which consists

in developing and experimentally validating an automated method for predicting the popularity of textual posts based on natural language processing and machine learning techniques. The object and subject of research are defined, and the set of applied methods is outlined, including statistical analysis, machine learning and deep learning algorithms, transformer models, and quality evaluation techniques.

The first chapter provides an overview of modern approaches to textual data analysis and content popularity prediction in social networks. The main stages of text processing, vector representation methods, classical machine learning algorithms, and modern deep learning models, particularly transformer architectures, are considered. The advantages and limitations of existing approaches are analyzed, as well as key challenges related to contextual dependence of texts, data imbalance, and model computational complexity.

The second chapter describes the developed method for evaluating the potential popularity of textual publications in social networks and substantiates its architecture. The proposed method is based on the combined use of semantic features obtained using transformer models and statistical characteristics of social platform metadata. The main stages of feature formation, mechanisms for integrating heterogeneous data, and principles of prediction model construction are presented.

The third chapter presents the implementation of the software for predicting the popularity of textual content and describes the applied tools and technologies. The choice of the programming language and libraries for text processing, machine learning model training, and modular system architecture is justified. The structure of the software system and the interaction of its core components are discussed.

The fourth chapter presents the results of experimental studies evaluating the effectiveness of the proposed popularity prediction method. A comparative analysis with baseline methods and existing approaches is conducted using standard classification quality metrics. The results demonstrate that the proposed

method provides improved prediction accuracy and stable performance across different data subsets.

The fifth chapter addresses business modeling aspects of the method for predicting the potential popularity of textual publications in social networks. The challenges of commercial implementation of the developed software solution are analyzed, including monetization, scalability, and integration with existing digital platforms. Key project stakeholders are identified and analyzed, allowing assessment of their influence, level of interest, and prioritization in the context of startup implementation.

The conclusions summarize the main research results, formulate scientific and practical findings regarding the effectiveness of the proposed method, and outline directions for further research, including multimodal data integration, real-time model adaptation, and prediction personalization.

The appendices contain graphical materials for the explanatory note, source code listings of the developed software, and presentation screenshots.

The work is performed on 139 sheets, contains 3 appendices, and references 11 literature sources. The work includes 6 figures and 8 tables.

Keywords: popularity prediction, social networks, natural language processing, transformer models, BERT, machine learning, textual data analysis.

ЗМІСТ

СПИСОК ТЕРМІНІВ, СКОРОЧЕНЬ ТА ПОЗНАЧЕНЬ	4
ВСТУП.....	8
1. АНАЛІЗ ПРЕДМЕТНОЇ ГАЛУЗІ ТА ІСНУЮЧИХ РІШЕНЬ.....	9
1.1. Проблематика аналізу текстів у соціальних мережах.....	9
1.2. Основні методи оброблення природної мови та прогнозування популярності	11
1.3. Огляд існуючих програмних рішень у сфері аналізу текстового контенту.....	21
1.4. Аналіз переваг і недоліків існуючих рішень	27
1.5. Висновки до розділу 1	28
2. ЗАПРОПОНОВАНИЙ СПОСІБ ПРОГНОЗУВАННЯ ПОТЕНЦІЙНОЇ ПОПУЛЯРНОСТІ ТЕКСТОВИХ ПУБЛІКАЦІЙ	30
2.1. Обґрунтування вибору підходу до прогнозування популярності текстових публікацій в соціальних мережах	30
2.2. Теоретичні основи прогнозування потенційної популярності публікацій у соціальних мережах	32
2.3. Особливості роботи з багатомовними та різножанровими текстовими публікаціями.....	39
2.4. Спосіб оцінки потенційної популярності текстових постів	43
2.5. Висновки до розділу 2.....	47
3. ПРОГРАМНА РЕАЛІЗАЦІЯ СИСТЕМИ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ ПОПУЛЯРНОСТІ ПОСТІВ	48
3.1. Обґрунтування вибору засобів реалізації.....	48
3.2. Архітектура розробленого програмного забезпечення.....	56
3.3. Реалізація основних функціональних компонентів системи.....	63
3.4. Висновки до розділу 3.....	65
4. АНАЛІЗ ЕФЕКТИВНОСТІ ЗАПРОПОНОВАНОГО СПОСОБУ	66
4.1. Методологія експериментального дослідження.....	66

4.2. Підготовка та характеристика експериментальних даних	68
4.3. Порівняльний аналіз ефективності різних способів	70
4.4. Висновки до розділу 4.....	76
5. БІЗНЕС МОДЕЛЮВАННЯ СПОСОБУ ПРОГНОЗУВАННЯ ПОТЕНЦІЙНОЇ ПОПУЛЯРНОСТІ ТЕКСТОВИХ ПУБЛІКАЦІЙ СОЦІАЛЬНИХ МЕРЕЖ.....	77
5.1. Опис проблеми.....	77
5.2. Зацікавлені сторони.....	78
5.3. Комерційне рішення. Основні характеристики.....	81
5.4. Доходи та витрати	82
5.5. Бізнес-модель	84
ВИСНОВКИ	86
СПИСОК ВИКОРИСТАНИХ ЛІТЕРАТУРНИХ ДЖЕРЕЛ	88
ДОДАТКИ	90

СПИСОК ТЕРМІНІВ, СКОРОЧЕНЬ ТА ПОЗНАЧЕНЬ

BERT (Bidirectional Encoder Representations from Transformers) – глибока нейронна модель для обробки природної мови, що використовує двобічне кодування контексту за допомогою архітектури трансформера для точного розуміння семантики слів у реченнях.

NLP (Natural Language Processing) – галузь штучного інтелекту та комп'ютерних наук, що займається автоматичним аналізом, розпізнаванням та генерацією природної мови для забезпечення ефективної взаємодії між людьми та комп'ютерами.

Dataset – структурована сукупність даних, зібраних та підготовлених для аналізу, навчання, валідації й тестування моделей машинного навчання.

BoW (Bag of Words) – метод представлення тексту у вигляді множини слів без урахування порядку, що дозволяє перетворювати тексти на числові вектори для машинного навчання.

TF-IDF (Term Frequency-Inverse Document Frequency) – статистичний метод для оцінки значущості слова у документі відносно колекції текстів, який підсилює рідкісні та пригнічує часті терміни.

Word Embeddings – числові векторні представлення слів, що зберігають семантичні відношення між словами у багатовимірному просторі.

Features (ознаки) – характеристики або властивості даних, що використовуються моделлю машинного навчання для побудови прогнозів або класифікації.

Flesch Reading Ease – метрика, що оцінює легкість сприйняття тексту на основі середньої довжини речень і складності слів. Чим вищий показник, тим текст легше читати, що дозволяє кількісно аналізувати читацьку доступність матеріалу.

Flesch–Kincaid Grade Level – показник читацької складності тексту, що оцінює його відповідність рівню освіти у США, необхідному для розуміння тексту, базуючись на довжині речень і складності слів.

Latent Dirichlet Allocation (LDA) – статистична модель тематичного моделювання, що дозволяє виявляти приховані теми в колекції документів шляхом представлення кожного документа як комбінації тем і кожної теми як розподілу слів.

Non-negative Matrix Factorization (NMF) – метод факторизації матриці на дві або більше матриць з обмеженням на невід’ємні елементи, що застосовується для виявлення прихованих структур у даних, зокрема для тематичного моделювання текстів.

Attention-based – підхід у нейронних мережах, що дозволяє моделі фокусуватися на релевантних частинах вхідних даних під час обробки послідовностей.

Boosting Machines (GBM) – ансамблевий метод машинного навчання, який послідовно об’єднує слабкі моделі для покращення точності прогнозів шляхом зменшення похибок попередніх моделей.

Feedforward Neural Networks (FNN) – клас штучних нейронних мереж, де інформація рухається тільки вперед від вхідного шару до вихідного без зворотних зв’язків.

Recurrent Neural Networks (RNN) – нейронні мережі для обробки послідовних даних, які зберігають інформацію про попередні стани через рекурентні зв’язки.

Трансформери (Transformer-based architectures) – архітектура нейронних мереж, що використовує механізм self-attention для паралельної обробки послідовностей і забезпечує високу ефективність у NLP.

Convolutional Neural Networks (CNN) – нейронні мережі з конволюційними шарами, оптимізовані для обробки даних із локальною структурою, таких як зображення або текст.

RoBERTa (Robustly optimized BERT approach) – модифікація BERT з покращеною продуктивністю завдяки оптимізації процесу передтренування та збільшенню обсягу навчальних даних.

NSP-задачі (Next Sentence Prediction) – завдання передтренування моделей, що передбачає правильне визначення логічного порядку речень у тексті.

F1-score – гармонійне середнє точності та повноти моделі, що використовується для оцінки якості класифікації, особливо при незбалансованих класах.

Confusion Matrix (Матриця неточностей) – таблиця, що показує кількість правильних і неправильних прогнозів моделі для кожного класу, допомагаючи аналізувати її ефективність.

GPU (Graphics Processing Unit) – обчислювальний пристрій, оптимізований для паралельної обробки великих обсягів даних.

Sentiment Analysis – автоматичне визначення емоційної тональності тексту (позитивна, негативна, нейтральна) за допомогою алгоритмів NLP.

Hugging Face – платформа та бібліотека для роботи з моделями NLP, що забезпечує доступ до передтренованих трансформерів та інструментів для обробки природної мови.

SVM (Support Vector Machine) – алгоритм навчання з учителем, який формулює задачу класифікації як оптимізаційну задачу, знаходячи гіперплощину з максимальним зазором (margin) між класами у просторі ознак, з можливістю використання ядрових функцій для відображення даних у високовимірний простір для розділення нелінійно роздільних класів.

Random Forest – ансамблевий метод машинного навчання, що складається з великої кількості випадково побудованих дерев рішень.

Logistic Regression – статистична модель та алгоритм машинного навчання для бінарної або багатокласової класифікації, що прогнозує ймовірність належності об'єкта до певного класу за допомогою логістичної функції (сигмоїди), перетворюючи лінійну комбінацію ознак у ймовірнісну оцінку.

LSTM (Long Short-Term Memory) – різновид рекурентних нейронних мереж, що використовує спеціальні комірки пам'яті та механізми гейтів для ефективного збереження та обробки довгострокових залежностей у послідовних даних, зменшуючи проблему зникнення градієнта.

GRU (Gated Recurrent Unit) – різновид рекурентних нейронних мереж, що спрощує архітектуру LSTM, об'єднуючи механізми пам'яті у два гейти (update та reset), що дозволяє ефективно моделювати послідовні дані та зберігати довгострокові залежності при меншій обчислювальній складності.

REST API (Representational State Transfer Application Programming Interface) – архітектурний стиль для створення вебсервісів, що дозволяє клієнтам взаємодіяти з сервером через стандартизовані HTTP-запити для отримання, створення, оновлення або видалення ресурсів, представлених у вигляді структурованих даних.

Docker – платформа для контейнеризації застосунків, що забезпечує ізольоване середовище та спрощує розгортання програмного забезпечення.

Parquet – колонково-орієнтований формат зберігання даних, оптимізований для аналітичних задач, що забезпечує ефективну компресію та швидкий доступ до окремих стовпців у великих масивах даних.

CLI (Command Line Interface) – інтерфейс взаємодії користувача з комп'ютером через текстові команди.

Gradient Boosting – ансамблевий метод машинного навчання, який послідовно будує слабкі моделі (зазвичай дерева рішень), кожна з яких коригує помилки попередніх, оптимізуючи функцію втрат за допомогою градієнтного спуску для підвищення точності прогнозів.

ВСТУП

З розвитком цифрових технологій та стрімким збільшенням обсягів текстової інформації в Інтернеті та соціальних мережах, виникає нагальна потреба у створенні ефективних методів автоматизованого аналізу та оброблення текстових даних. Серед численних завдань цієї галузі особливе місце займає проблема визначення та прогнозування популярності текстових повідомлень.

Текстові повідомлення в соціальних мережах містять як структуровану, так і неструктуровану інформацію, яка впливає на їх сприйняття аудиторією та рівень поширення. Ручний аналіз таких обсягів даних є неможливим, тому застосовуються методи машинного навчання, статистичного оброблення та оброблення природної мови (NLP), які дозволяють автоматизувати процес оцінки популярності та виявлення ключових чинників, що на неї впливають.

Однак, існуючі способи часто мають обмеження через складність семантичних зв'язків, неоднорідність мовного матеріалу та багатовимірність чинників, що визначають популярність.

Метою даної магістерської роботи є розроблення способу та програмного забезпечення для визначення потенційної популярності текстових публікацій у соціальних мережах, що враховує складні лінгвістичні та поведінкові характеристики контенту. Запропонований спосіб передбачає застосування комплексного способу, який інтегрує традиційні методи оброблення тексту з сучасними трансформерними архітектурами, для покращення точності прогнозування потенційної популярності текстових публікацій соціальних мереж.

Результати дослідження можуть бути використані у створенні інструментів маркетингової аналітики, систем рекомендацій та платформах соціальних мереж для більш ефективного управління контентом і взаємодії з аудиторією.

1. АНАЛІЗ ПРЕДМЕТНОЇ ГАЛУЗІ ТА ІСНУЮЧИХ РІШЕНЬ

1.1. Проблематика аналізу текстів у соціальних мережах

У сучасному інформаційному суспільстві соціальні мережі відіграють ключову роль у поширенні думок, новин, емоцій та подій. Зростання обсягів користувацького контенту породжує необхідність ефективного аналізу текстових повідомлень для виявлення закономірностей, трендів, потреб аудиторії та передбачення реакцій. Особливої актуальності набуває завдання аналізу потенційної популярності текстових публікацій як інструменту для прогнозування охоплення, вподобайок, репостів чи коментарів. Це має практичну цінність для маркетингу, політики, журналістики, соціології та інших сфер.

Однак аналіз текстів у соціальних мережах супроводжується низкою серйозних викликів:

- Неструктурованість та неформальність тексту. Контент у соціальних мережах часто є неструктурованим, включає сленг, емодзі, граматичні помилки, скорочення, жаргон, візуальні символи та нестандартні форми мови. Це ускладнює застосування традиційних методів оброблення природної мови (Natural Language Processing, NLP) [1], які зазвичай орієнтовані на літературні або формальні тексти.
- Короткий обсяг повідомлень. Багато платформ, зокрема Twitter (від липня 2023 – X), мають обмеження на кількість символів у повідомленні. Це змушує користувачів викладати думки стисло, іноді жертвуючи граматиною чи контекстом. Для алгоритмів машинного навчання така обмеженість інформації може ускладнювати формування репрезентативного вектору ознак.
- Висока динаміка лексики. Мова у соцмережах змінюється дуже швидко: з'являються нові меми, хештеги, тренди, неологізми. Це означає, що лінгвістичні моделі повинні постійно оновлюватися,

адаптуватися до нових слів і контекстів. Статичні словники або попередньо натреновані моделі швидко втрачають актуальність.

- Контекстуальність та полісемія. Один і той самий вираз може мати різні значення залежно від контексту або інтонації. Наприклад, фраза «класно, просто супер» може бути як щирою, так і саркастичною. Визначення емоційного забарвлення повідомлення вимагає глибокого контекстуального розуміння, що є складним завданням навіть для великих трансформерних моделей.
- Вплив візуального та мультимедійного контексту. Пости часто супроводжуються зображеннями, відео, гіфками або гіперпосиланнями, які мають значний вплив на популярність, але сам текст при цьому може бути коротким або взагалі другорядним. Неможливо повно оцінити популярність такого контенту лише на основі тексту.
- Аудиторія та час публікації. Популярність поста залежить не лише від змісту, але й від зовнішніх факторів: кількості підписників, часу публікації, дня тижня, поточних трендів. Це створює мультифакторну природу прогнозування, де текст є лише однією з багатьох складових.

Підсумовуючи вищезазначене, аналіз текстів у соціальних мережах є складним завданням через неформальність мовлення, стрімку зміну лексики, контекстуальність висловів та вплив зовнішніх факторів.

Сукупність цих викликів ускладнює прогнозування популярності публікацій, оскільки текст рідко є єдиним чинником, що визначає реакцію аудиторії.

У цій роботі буде розглянуто спосіб, що комбінує лінгвістичні та контекстуальні характеристики текстових публікацій в соціальних мережах.

1.2. Основні методи оброблення природної мови та прогнозування популярності

У рамках завдання прогнозування популярності повідомлень у соціальних медіа ключову роль відіграють методи оброблення природної мови (Natural Language Processing, NLP) та алгоритми машинного й глибокого навчання, які дозволяють формалізувати взаємозв'язок між текстовими ознаками та рівнем залученості користувачів (engagement).

Нижче розглянемо основні методи оброблення природної мови. Зокрема, виділимо етапи попередньої обробки тексту, способи його векторизації, формування ознак, алгоритми машинного та глибокого навчання, а також ключові метрики оцінки якості моделей.

1.2.1. Попередня обробка тексту (*Text preprocessing*)

Попереднє оброблення є першим і необхідним етапом в NLP-завданнях. Для текстів із соціальних мереж вона зазвичай включає наступні кроки:

- Нормалізація тексту: перетворення всіх символів у нижній регістр, видалення пунктуації, html-розмітки, емодзі.
- Токенізація: розбиття тексту на слова або підслова (tokens).
- Стемінг та лематизація: приведення слів до базової форми для зменшення розмірності словника.
- Видалення стоп-слів: виключення часто вживаних, але малозначущих слів ("і", "та", "але" тощо).
- Обробка емодзі та скорочень: заміна символів на текстовий опис (наприклад, "❤️" → "love").
- Видалення шуму: URL-адрес, хештегів, згадок користувачів (@username), якщо вони не несуть інформативності.

1.2.2. Векторизація тексту

Після очищення тексту його необхідно перетворити у числове представлення, що підходить для алгоритмів машинного навчання. Найбільш поширені методи:

- Bag-of-Words (BoW) [2]: представлення тексту як частотного розподілу слів. Це один з найпростіших способів, що ґрунтується на підрахунку кількості входжень кожного слова в документі. Розрахунок вектора розмірності словника виконується за формулою:

$$x_i = [f_1, f_2, \dots, f_n], \quad (1)$$

де f_j – частота слова j у документі i ; n – кількість унікальних слів у корпусі.

- TF-IDF (Term Frequency-Inverse Document Frequency) [3]. Покращення методу BoW, що враховує не лише кількість входжень слова, але й його унікальність у всьому корпусі представлено у формулі:

$$TF - IDF(t, d, D) = TF(t, d) \cdot \log\left(\frac{N}{|\{d \in D : t \in d\}|}\right), \quad (2)$$

де $TF(t, d)$ – частота терміну t у документі i ; N – загальна кількість документів; $|\{d \in D : t \in d\}|$ – кількість документів, що містять термін t .

- Word Embeddings (word2vec, GloVe, fastText) [4] методи побудови щільних векторів (embedding) для слів на основі їх контексту. Дають змогу вловити семантичні зв'язки між словами. Наприклад, word2vec навчає вектори слів так, щоб схожі за контекстом слова були близькими у векторному просторі. Формули для представлення Word Embeddings:

$$CBOW: P(w_t | w_{t-k}, \dots, w_{t+k}), \quad (3)$$

$$Skip - gram: P(w_{t-k}, \dots, w_{t+k} | w_t). \quad (4)$$

- BERT, RoBERTa, DistilBERT: трансформерні моделі, що дозволяють формувати контекстно-залежні векторні представлення слів або цілих речень. Doc2Vec – узагальнення методу Word2Vec, яке дозволяє отримувати векторне представлення для всього документа. Sentence-BERT – модифікація моделі BERT, призначена для формування семантично значущих векторів речень [5].

1.2.3. Побудова ознак для прогнозу популярності

Процес прогнозування популярності текстового контенту в соціальних мережах вимагає ретельного конструювання вхідних ознак (features), що подаються на вхід моделі машинного або глибокого навчання. На відміну від задачі аналізу тональності чи класифікації тематик, тут необхідно враховувати не лише зміст повідомлення, але й різноманітні мета- та поведінкові характеристики. Сукупність таких ознак формується у вектор ознак (feature vector), що є числовим представленням спостереження у багатовимірному просторі.

До першої групи ознак відносяться лінгвістичні метрики. Серед них – кількість слів, кількість унікальних слів, середня довжина речення, а також співвідношення слів з позитивною або негативною емоційною забарвленістю. Широко застосовуються також індекси читабельності, такі як Flesch Reading Ease та Flesch–Kincaid Grade Level [6], представлені у формулі:

$$FRE = 206.835 - 1.015 \left(\frac{total\ words}{total\ sentences} \right) - 84.6 \left(\frac{total\ syllables}{total\ words} \right). \quad (5)$$

Ці метрики дозволяють оцінити складність тексту та його потенційну доступність для читача, що може впливати на рівень взаємодії з постом.

До другої групи належать семантичні та тематичні ознаки. Сюди входить аналіз емоційної тональності (sentiment analysis), що дозволяє визначити полярність тексту (позитивну, негативну або нейтральну). Для оцінки тематики застосовуються тематичні моделі, зокрема Latent Dirichlet

Allocation (LDA), Non-negative Matrix Factorization (NMF) або сучасні способи, побудовані на трансформерах, як-от BERTopic. Семантичні ознаки також можуть включати наявність у тексті певних ключових слів, згадок відомих брендів або політичних тем, що зазвичай мають вищий рівень взаємодії з аудиторією [5].

До третьої категорії відносяться метадані та поведінкові характеристики. Важливими є довжина повідомлення в символах і словах, наявність хештегів, емодзі, згадок інших користувачів, посилань або мультимедійних вкладень. Також слід враховувати дату і час публікації, зокрема день тижня та годину доби, оскільки активність аудиторії залежить від цих чинників. Додатково враховується тип платформи, що використовується, оскільки, наприклад, характер поширення дописів у Twitter відрізняється від Instagram.

Окрему увагу слід приділити інформації про автора. До цієї групи ознак можна віднести кількість підписників, середню активність підписників (відношення реакцій до аудиторії), частоту публікацій, а також наявність верифікації (вплив статусу "підтвердженого" акаунта). Часто розраховується коефіцієнт взаємодії (engagement rate), який визначається формулою:

$$Engagement Rate = \frac{Likes + Shares + Comments}{Followers}. \quad (6)$$

Окрім поточних публікацій, до ознак можуть залучатися історичні дані: середні показники популярності попередніх дописів, зміна активності з часом, реакція аудиторії на аналогічний контент. Такі історичні патерни особливо корисні при використанні моделей з пам'яттю, зокрема рекурентних або attention-based архітектур.

Таким чином, ознаки для прогнозування популярності охоплюють широкий спектр параметрів – від текстових характеристик до часових, соціальних та контекстуальних факторів. Їх поєднання в єдиному векторі дозволяє підвищити точність моделі та забезпечити глибше розуміння

закономірностей взаємодії користувачів із контентом. У сучасних системах машинного навчання дедалі частіше використовуються мультимодальні способи, де різні типи ознак обробляються окремими нейронними мережами з подальшою конкатенацією їх виходів для побудови остаточного прогнозу.

1.2.4. Методи машинного навчання

У процесі побудови системи прогнозування популярності дописів у соціальних мережах важливу роль відіграє вибір математичної моделі, яка дозволяє виявити закономірності між вхідними ознаками та цільовими змінними. Залежно від постановки задачі регресії або класифікації застосовуються відповідні способи: від базових лінійних моделей до складних глибоких нейронних мереж і мультимодальних архітектур.

Лінійні моделі є фундаментом у статистичному аналізі даних та забезпечують інтерпретованість результатів. У випадку, коли необхідно передбачити числову популярність (наприклад, кількість лайків, репостів або коментарів), застосовується лінійна регресія, яка встановлює лінійну залежність між вектором ознак x та цільовою змінною y . Стандартна модель представлена формулою:

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n + \varepsilon, \quad (7)$$

де β_i – коефіцієнти регресії, які відображають важливість відповідної ознаки; ε – випадкова похибка.

У випадках, коли потрібно здійснити класифікацію (наприклад, поділ постів на «популярні» та «непопулярні» за певним порогом), використовується логістична регресія, яка оцінює ймовірність належності до певного класу та представлення формулою:

$$P(y = 1 | x) = \frac{1}{1 + e^{-(\beta_0 + \beta^T x)}}. \quad (8)$$

Ця модель широко використовується завдяки простоті, швидкості навчання та високій інтерпретованості.

Дерева рішень та ансамблеві методи є наступним кроком у підвищенні продуктивності моделей для оброблення табличних або гетерогенних даних. Дерева рішень будують послідовність умовних переходів, як вказано у формулі (9), що дозволяє формувати інтерпретовані правила для прогнозу. Проте, як відомо, окреме дерево може бути нестабільним, тому застосовуються ансамблі моделей, зокрема Random Forest та Gradient Boosting Machines (GBM). Останні реалізуються через ефективні бібліотеки, такі як XGBoost, LightGBM, CatBoost. GBM навчає моделі послідовно, мінімізуючи функцію втрат на основі помилок попередніх моделей.

$$F_m(x) = F_{m-1}(x) + \gamma_m h_m(x), \quad (9)$$

де $h_m(x)$ – нова модель, яка апроксимує негативний градієнт функції втрат; γ_m – коефіцієнт оновлення.

Нейронні мережі надають змогу обробляти складні, високорозмірні та структуровано неявні ознаки. Для задач прогнозування популярності можуть застосовуватись кілька типів архітектур:

- Feedforward Neural Networks (FNN) – базова архітектура багат шарової перцептронної мережі, яка ефективно працює на збагаченому векторі ознак, зокрема якщо вони включають як числові, так і категоріальні атрибути.
- Recurrent Neural Networks (RNN), а також її покращені варіанти LSTM (Long Short-Term Memory) та GRU (Gated Recurrent Unit) – дозволяють моделювати часову або послідовну структуру тексту. Завдяки механізмам пам'яті ці архітектури здатні враховувати контекст попередніх слів, що підвищує якість моделювання семантики.
- Трансформери, зокрема BERT, GPT та RoBERTa, які завдяки механізму self-attention, представленого формулою (10), обробляють текст у глобальному контексті, що дозволяє враховувати взаємозв'язки між усіма словами в реченні незалежно від їх порядку.

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V, \quad (10)$$

де Q , K , V – матриці запитів, ключів та значень відповідно;
 d_k – розмірність простору ключів.

Мультимодальні способи прогнозування потенційної популярності текстових публікацій соціальних мереж є сучасним напрямом у задачах прогнозування популярності, оскільки популярність контенту зумовлюється не лише текстовим описом, але й зображенням, метаданими, а також особливостями взаємодії з користувачами. У таких випадках застосовуються моделі, які поєднують окремі нейронні архітектури для різних типів входів (текстових, візуальних, табличних), як вказано у формулі:

$$f(x_{text}, x_{image}, x_{meta}) \rightarrow y. \quad (11)$$

Така архітектура дозволяє ефективно агрегувати інформацію з різнорідних джерел за допомогою механізмів конкатенації, уваги (attention fusion) або крос-модального навчання.

Загалом, вибір моделі визначається як природою задачі, так і доступними обчислювальними ресурсами. Для невеликих наборів даних доцільними є прості регресійні або ансамблеві способи, тоді як для оброблення великих корпусів текстів і зображень у реальному часі перевага надається глибоким та мультимодальним нейронним мережам.

1.2.5. Глибинне навчання та трансформери

Розвиток глибинного навчання (deep learning) суттєво змінив підхід до оброблення природної мови (Natural Language Processing, NLP), дозволяючи моделювати складні семантичні та синтаксичні структури тексту. Особливої популярності набули нейронні архітектури, засновані на трансформерах (Transformer), які забезпечують як високу точність, так і можливість масштабування на великі обсяги даних.

Класичні рекурентні нейронні мережі (Recurrent Neural Networks, RNN) є базовими архітектурами для роботи з послідовностями, однак вони схильні до проблеми зникання градієнта. Для подолання цього недоліку були запропоновані поліпшені варіанти – LSTM (Long Short-Term Memory) та GRU (Gated Recurrent Unit) [4].

Архітектура LSTM включає в себе механізми контролю інформаційного потоку через входні, вихідні та забувальні "вікна" (gates), що дозволяє зберігати важливу інформацію протягом довгих послідовностей. GRU є спрощеним варіантом LSTM і має меншу кількість параметрів, що прискорює навчання. Основна ідея GRU полягає в об'єднанні механізмів забування та введення в одне "оновлювальне вікно".

Згорткові нейронні мережі (Convolutional Neural Networks, CNN) здебільшого використовуються в комп'ютерному зорі, вони також можуть ефективно застосовуватись для оброблення тексту. У цьому випадку текст подається як послідовність векторів (ембедінгів), до яких застосовуються фільтри (kernels) різної ширини, що дозволяє виділяти локальні шаблони (наприклад, біграми чи триграми).

Прорив у NLP пов'язаний із появою моделі Transformer, запропонованої у роботі "Attention is All You Need" (Vaswani et al., 2017) [7]. На відміну від RNN, трансформери дозволяють моделювати глобальний контекст у тексті завдяки механізму самоуваги.

На базі архітектури Transformer побудовано низку високопродуктивних моделей:

- BERT – навчається в режимі маскованого моделювання мови (Masked Language Modeling, MLM) та враховує контекст з обох боків (бідірекційно). У задачах прогнозування популярності тексту BERT часто застосовується у режимі fine-tuning – донавчання на специфічній задачі, наприклад, класифікації популярності поста.
- GPT – модель автогенеративного типу, що використовується для побудови послідовностей, особливо в генеративних способах

модельовання популярності (наприклад, генерації варіантів дописів із різним очікуваним впливом).

- RoBERTa – модифікація BERT, що оптимізована для кращої продуктивності за рахунок покращеного навчального процесу (усунення NSP-задачі, збільшення розміру batch, більші обсяги даних тощо).

Завдяки використанню попередньо навчених ембедінгів ці моделі здатні витягувати високорівневі ознаки, що істотно підвищує якість прогнозу навіть на невеликих вибірках.

1.2.6. Метрики оцінки якості прогнозування

Для оцінки ефективності моделей прогнозування використовуються різні метрики, що залежать від типу задачі – класифікації або регресії. Вибір метрики впливає на інтерпретацію результатів та допомагає вибрати найкращу модель.

Метрики для задач класифікації:

- Ассурасу (Точність) – частка правильно класифікованих прикладів представлена формулою:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (12)$$

де TP – кількість істинно позитивних; TN – істинно негативних; FP – хибно позитивних; FN – хибно негативних прогнозів.

- Precision (Точність для позитивного класу) – частка істинно позитивних серед усіх позитивних прогнозів як вказано в формулі:

$$Precision = \frac{TP}{TP + FP} \quad (13)$$

- Recall (Повнота) – частка істинно позитивних, серед усіх фактичних позитивних:

$$Recall = \frac{TP}{TP + FN} \quad (14)$$

- F1-score – гармонійне середнє Precision і Recall, що наведено формулою:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (15)$$

Ці метрики особливо корисні при класифікації постів на популярні/непопулярні, коли класи можуть бути дисбалансними.

Метрики для багатокласової класифікації:

- Macro-F1 – середнє F1-score по всіх класах без урахування частоти, що розраховується за допомогою формули:

$$Macro - F1 = \frac{1}{K} \sum_{k=1}^K F1_k, \quad (16)$$

де K – кількість класів.

- Confusion Matrix (Матриця неточностей) – таблиця, що відображає кількість правильних та неправильних класифікацій по кожному класу. Вона допомагає детально оцінити типи помилок.

1.2.7. Проблеми та обмеження сучасних способів

Незважаючи на значний прогрес у галузі оброблення природної мови (Natural Language Processing, NLP) та використання сучасних моделей, таких як трансформери, існує низка суттєвих проблем і обмежень, що впливають на якість та застосовність моделей прогнозування популярності текстового контенту.

Однією з ключових проблем є перенавчання моделей на тренувальних даних, тобто коли модель надмірно адаптується під навчальні приклади і втрачає здатність узагальнювати на нові, раніше не бачені дані.

Це особливо актуально для складних моделей з великою кількістю параметрів, таких як глибокі нейронні мережі. Формально, перенавчання характеризується низькою помилкою на тренувальних даних і високою – на тестових.

Мова – динамічна система, що постійно розвивається. Зміни у словниковому запасі, стилях спілкування, трендах і контекстах призводять до явища *concept drift* – зміни розподілу вхідних даних з часом. Це викликає погіршення роботи моделей, навчених на історичних даних, без можливості адаптації до нових умов.

Архітектури на основі трансформерів (BERT, GPT, RoBERTa) потребують значних обчислювальних ресурсів для навчання і навіть для інференсу. Високі вимоги до обсягу оперативної пам'яті, GPU-часу та енергоспоживання обмежують їх застосування у практичних задачах, особливо на пристроях з обмеженими ресурсами або в режимі реального часу.

Моделі можуть успадковувати і посилювати упередження, присутні у тренувальних даних, що призводить до дискримінаційних результатів. Це включає соціальні, гендерні та інші форми *bias*, які негативно впливають на справедливість і етичність прогнозів.

Для досягнення високої якості прогнозування сучасні моделі потребують масштабних анотованих корпусів, що може бути дорогим і тривалим процесом. Особливо це стосується мультимодальних і контекстуальних моделей, які вимагають складної розмітки і різнобічних даних.

1.3. Огляд існуючих програмних рішень у сфері аналізу текстового контенту

Сучасний аналіз текстового контенту є міждисциплінарною галуззю, яка об'єднує методи оброблення природної мови (NLP), машинного навчання та статистичного аналізу для витягання смислової інформації із текстів різної природи. Існуючі програмні рішення можна класифікувати за типом задач, функціональними можливостями та технологічною базою.

1.3.1. Системи тематичного моделювання та класифікації текстів

Тематичне моделювання дозволяє автоматично виявляти основні теми у великій колекції текстів, що є критично важливим для аналізу контенту у соціальних мережах, медіа та форумах. Одним із найпоширеніших способів є Latent Dirichlet Allocation (LDA) [8], який представляє документи як ймовірнісну комбінацію тем, а теми – як ймовірнісну комбінацію слів.

- Gensim – Python-бібліотека з вбудованою підтримкою LDA, яка дозволяє масштабувати моделювання тем на великі обсяги даних.
- MALLET – Java-інструмент для статистичного моделювання тексту, що підтримує LDA з можливістю тонкого налаштування гіперпараметрів.

Після побудови тематичної моделі документи можна класифікувати або кластеризувати за переважними темами, що дозволяє сегментувати контент за інтересами, сферами або цільовими аудиторіями.

Класифікація текстів – ще один важливий напрям, який охоплює задачі визначення категорії поста (новини, реклама, коментар), виявлення тональності (позитивна, негативна, нейтральна) та фільтрації спаму.

- Scikit-learn – універсальний фреймворк машинного навчання, який включає в себе класичні алгоритми для оброблення тексту.
- Підготовка тексту виконується за допомогою TF-IDF, CountVectorizer або векторів із Word2Vec/Doc2Vec.
- Отримані ознаки подаються на вхід моделі, яка навчається на розмічених даних і використовується для автоматичної класифікації нових текстів.

У поєднанні з тематичним моделюванням ці способи формують базову інфраструктуру для побудови систем автоматичного аналізу текстового контенту. Вони можуть бути розширені за допомогою більш

складних моделей – таких як нейронні мережі або трансформери – для досягнення вищої точності у специфічних сценаріях.

1.3.2. Платформи для аналізу емоційності та тональності

Аналіз емоційності та тональності тексту (sentiment analysis) є ключовою складовою систем інтелектуального аналізу соціальних мереж, відгуків, коментарів та іншого контенту, що генерується користувачами. Основна мета – автоматичне визначення емоційного забарвлення повідомлення: позитивного, негативного чи нейтрального.

Одним із найбільш популярних інструментів для базового аналізу сентименту є VADER (Valence Aware Dictionary and sEntiment Reasoner), який поєднує лексикон оцінок емоційності з евристичними правилами. VADER орієнтований на соціальні медіа та враховує інтенсивність емоційних слів, модифікатори (наприклад, підсилювачі чи заперечення), пунктуацію, емодзі та великі літери. Цей інструмент дозволяє швидко оцінити тональність тексту без потреби у попередньому навчанні.

Інші способи ґрунтуються на методах машинного навчання. Бібліотека TextBlob забезпечує простий інтерфейс до попередньо навчених моделей класифікації сентименту, використовуючи наївні байєсівські класифікатори. Незважаючи на обмежену точність у складних контекстах, цей спосіб є зручним для швидкої інтеграції.

Сучасні платформи, такі як Hugging Face Transformers, пропонують доступ до потужних попередньо навчених моделей на основі трансформерної архітектури (наприклад, BERT, DistilBERT, RoBERTa), які демонструють високі результати в задачах тонального аналізу. Моделі проходять fine-tuning на відповідних наборах даних, таких як SST-2 (Stanford Sentiment Treebank), IMDb чи Yelp Reviews. Вхідний текст зазвичай подається у вигляді токенованої послідовності, яка обробляється багатопаровим механізмом attention, що дозволяє моделі враховувати як локальні, так і глобальні контексти.

Переваги трансформерів включають глибоке розуміння семантики та контексту, що є критичним для розпізнавання іронії, сарказму, складних емоційних конструкцій. Недоліки – потреба у великих обчислювальних ресурсах та ризик перенавчання, якщо fine-tuning здійснюється на невеликій вибірці.

Загалом, обрана платформа для аналізу настрою має відповідати прикладній задачі, обсягу даних та вимогам до точності. Інтеграція лексиконних і нейромережових методів є перспективним напрямом для підвищення стійкості та адаптивності моделей до змін у мовному середовищі.

1.3.3. Нейромережеві та трансформерні фреймворки

Значний прорив у сфері оброблення природної мови пов'язаний із розвитком нейромережових архітектур, зокрема трансформерів, які забезпечують високий рівень контекстуального розуміння тексту. На відміну від класичних моделей, трансформери дозволяють одночасно враховувати взаємозв'язки між усіма словами в реченні, завдяки чому досягається глибше семантичне представлення.

Серед найбільш відомих трансформерних моделей – BERT, GPT та RoBERTa.

Вони продемонстрували виняткову ефективність у різноманітних NLP-завданнях: від класифікації текстів і аналізу тональності до генерації контенту та відповіді на запитання. Ці моделі є багатошаровими і здатні адаптуватися до нових доменів шляхом донавчання (fine-tuning) на цільових наборах даних.

Фреймворк Hugging Face Transformers надає зручний API до сотень попередньо навчених моделей, що охоплюють десятки мов і задач. Користувач може швидко інтегрувати модель у свій застосунок, змінити лише вихідний шар та виконати додаткове навчання з урахуванням специфіки прикладної проблеми. Бібліотека також підтримує завантаження

моделей з хмари, оптимізацію на GPU/TPU та обробку великих обсягів текстових даних.

Для розробки власних нейромережових рішень часто використовуються TensorFlow і PyTorch – два найбільш популярні фреймворки для глибинного навчання. Вони дозволяють реалізовувати кастомні архітектури, створювати моделі з декількох вхідних потоків (наприклад, текст і зображення), використовувати техніки регуляризації та оптимізації, масштабувати навчання на кількох графічних процесорах. Крім того, ці фреймворки інтегруються з іншими бібліотеками для оброблення даних, візуалізації, логування та деплою моделей.

Можливість конструювання мультимодальних моделей – ще одна важлива перевага сучасних нейронних фреймворків. В умовах, коли популярність текстового контенту часто залежить не лише від самого тексту, а й від візуального супроводу чи часу публікації, стає актуальним об'єднання різних типів даних у єдину модель. Наприклад, мультимодальні трансформери здатні поєднувати текстові та візуальні ознаки у спільному семантичному просторі, що значно покращує якість прогнозування у реальних прикладних задачах.

Таким чином, поєднання трансформерних моделей і сучасних фреймворків забезпечує гнучкість, масштабованість та високу точність систем інтелектуального аналізу тексту.

1.3.4. Комерційні та відкриті продукти

У сфері аналізу текстового контенту існує широкий спектр готових рішень, які можна умовно поділити на комерційні сервіси та відкриті програмні бібліотеки.

Комерційні сервіси зазвичай надають хмарну інфраструктуру з доступом до функціональних можливостей через API. Вони ідеально підходять для швидкої інтеграції в існуючі продукти, мають хорошу

масштабованість, документацію та технічну підтримку. Серед найвідоміших:

- Google Cloud Natural Language API – забезпечує класифікацію тексту за категоріями, витяг сутностей, аналіз настрою та синтаксичний розбір [9].
- IBM Watson Natural Language Understanding – орієнтований на корпоративні потреби, дозволяє отримувати глибоку семантичну інформацію з тексту, підтримує кастомні категорії [10].
- Microsoft Azure Text Analytics – надає модулі для мовної аналітики, включаючи визначення мови, тональності, ключових фраз і зв'язків між сутностями [11].

Основними перевагами таких платформ є простота інтеграції, висока якість результатів (завдяки навчанням на великих обсягах даних) та регулярне обов'язкове передачу даних у хмару (що може бути критичним для конфіденційності) та обмеження у налаштуванні моделей під конкретні задачі.

Відкриті програмні платформи надають дослідникам та розробникам повну гнучкість у побудові власних NLP-пайплайнів. Найбільш популярні:

- spaCy – сучасна бібліотека з високою продуктивністю, підтримкою нейронних мереж, інтеграцією з Hugging Face Transformers та зручним API для оброблення великих обсягів тексту. Підтримує такі задачі, як токенізація, розпізнавання сутностей, частиномовний аналіз, лематизація тощо.
- NLTK (Natural Language Toolkit) – орієнтована переважно на академічні дослідження та навчальні цілі. Містить широкий набір алгоритмів для лінгвістичного оброблення тексту, а також численні корпуси для навчання та експериментів.

Такі бібліотеки добре підходять для створення кастомізованих моделей, де потрібна гнучка адаптація, локальна обробка даних або інтеграція з іншими компонентами машинного навчання.

Таким чином, вибір між комерційними та відкритими рішеннями залежить від поставленої задачі, обмежень безпеки, вимог до масштабування та необхідного рівня точності аналізу. У практиці розробки складних систем часто використовується гібридний спосіб – поєднання готових сервісів для швидкого старту та відкритих інструментів для подальшої кастомізації та оптимізації моделей.

1.4. Аналіз переваг і недоліків існуючих рішень

Розвиток методів аналізу текстового контенту, зокрема у контексті прогнозування популярності, дозволив значно покращити якість автоматичних систем. Однак, попри досягнення, жоден спосіб не є універсальним і позбавленим обмежень.

Традиційні статистичні методи відзначаються простотою реалізації та високою інтерпретованістю результатів, що робить їх ефективними у задачах з невеликим обсягом даних або коли ознаки вже добре структуровані, наприклад, частота слів або TF-IDF. Вони потребують мінімальних обчислювальних ресурсів, що є додатковою перевагою при обмежених можливостях апаратного забезпечення.

Разом із тим, такі методи мають обмежену здатність враховувати контекст слів і демонструють низьку ефективність на великих обсягах неструктурованих даних. Успішність їх застосування значною мірою залежить від якості ручної інженерії ознак.

Класичні алгоритми машинного навчання, такі як SVM, Random Forest або логістична регресія, добре зарекомендували себе для роботи з векторизованими текстами, наприклад при використанні Bag-of-Words або TF-IDF. Вони забезпечують порівняно швидке навчання і тестування моделей, а результати, особливо у випадку лінійних моделей, легко інтерпретуються. Проте ці алгоритми не враховують синтаксичний та семантичний контекст тексту, потребують попередньої обробки та

нормалізації даних, і їх ефективність знижується при роботі з великою кількістю класів або складною структурою даних.

Рекурентні нейронні мережі (LSTM, GRU) дозволяють враховувати порядок слів та послідовність у тексті, що робить їх особливо корисними у задачах, де важливі довготривалі залежності. Водночас ці моделі складні у налаштуванні та навчанні, важко піддаються паралелізації, а при роботі з довгими послідовностями можливі втрати інформації, навіть при використанні механізмів пам'яті.

Трансформерні моделі (BERT, RoBERTa, GPT) забезпечують глибоке контекстуальне розуміння тексту, що дозволяє досягати високої точності прогнозування. Вони гнучкі у роботі з різними типами задач, такими як класифікація, генерація тексту та витяг інформації, і можуть бути додатково донавчені під конкретну задачу (fine-tuning). Проте трансформери характеризуються високими обчислювальними витратами, особливо під час навчання, потребують великих обсягів даних для повноцінного донавчання і можуть відтворювати упередження, що містилися у тренувальному корпусі.

Мультимодальні підходи дозволяють інтегрувати різні джерела даних, такі як текст, зображення та метадані, що забезпечує більш комплексний аналіз і підвищує точність прогнозування популярності контенту. Водночас їх реалізація є складною, вимагає узгодження та синхронізації різнорідних типів даних, а також значних обчислювальних ресурсів і застосування складних архітектур моделей.

1.5. Висновки до розділу 1

У першому розділі здійснено комплексний огляд теоретичних засад та сучасних способів до аналізу текстового контенту, зокрема у контексті прогнозування його популярності. Розглянуто основні етапи оброблення тексту, включаючи очищення, нормалізацію та векторизацію, а також

методи глибокого навчання і трансформерні моделі, що значно підвищили якість розв'язання мовних задач.

Особливу увагу приділено існуючим метрикам оцінки якості прогнозування та ключовим проблемам, зокрема перенавчанню, залежності від великих обсягів даних та ресурсомісткості сучасних моделей. Проаналізовано можливості застосування різних класів інструментів: від традиційних статистичних моделей до мультимодальних трансформерів, а також переваги й обмеження популярних програмних рішень – як відкритих, так і комерційних.

Загалом, сучасні способи до оброблення тексту забезпечують високу точність прогнозів, однак ефективне їх використання вимагає ретельного налаштування, глибокого розуміння структури даних і достатніх обчислювальних ресурсів.

2. ЗАПРОПОНОВАНИЙ СПОСІБ ПРОГНОЗУВАННЯ ПОТЕНЦІЙНОЇ ПОПУЛЯРНОСТІ ТЕКСТОВИХ ПУБЛІКАЦІЙ

2.1. Обґрунтування вибору підходу до прогнозування популярності текстових публікацій в соціальних мережах

У контексті розвитку соціальних мереж та зростання обсягів генерованого контенту, завдання прогнозування популярності постів набуває все більшого значення. Воно є актуальним як для окремих користувачів, що прагнуть збільшити охоплення своєї аудиторії, так і для бізнесу, де популярність публікацій безпосередньо впливає на ефективність маркетингових кампаній, брендову впізнаваність та зростання лояльності клієнтів.

Аналіз існуючих способів показав, що найбільш ефективні результати демонструють ті, що враховують не лише текстову складову публікації, але й додаткові контекстуальні чинники: час публікації, візуальні матеріали (зображення чи відео), взаємодію користувачів, популярність автора тощо. Це обумовлює доцільність використання мультимодального підходу, який дозволяє поєднувати ознаки з різних джерел.

Для вирішення задачі прогнозування було обрано спосіб, що ґрунтується на поєднанні таких компонентів як:

- Глибинний текстовий аналіз за допомогою трансформерних моделей, таких як BERT, що дозволяють витягати контекстуальні ознаки з тексту публікації, з урахуванням семантичних зв'язків між словами та фразами.
- Інкorporація метаданих – таких як час публікації, кількість підписників, наявність мультимедійних елементів, категорія або тип посту, вживання хештегів – які виступають числовими, категоріальними або бінарними ознаками.
- Можливе додавання векторних ознак, отриманих з оброблення зображень, за допомогою моделей комп'ютерного зору.

- Поєднання ознак у єдиній моделі, що реалізується шляхом побудови ансамблевих алгоритмів (XGBoost, CatBoost, LightGBM) або багат шарових нейронних мереж, здатних працювати з гетерогенними даними різної природи.

Обґрунтування такого підходу базується на наступних перевагах:

- Забезпечення контекстуального розуміння тексту за допомогою сучасних трансформерних архітектур.
- Гнучке включення різнорідних ознак без потреби в суворій попередній трансформації.
- Підвищення точності прогнозу завдяки поєднанню інформації з різних модальностей.
- Узгодженість мультимодальних даних – використання attention-механізмів дає змогу ефективно агрегувати та вагово оцінювати інформацію з тексту, зображення та метаданих.
- Можливість подальшого донавчання (fine-tuning) на специфічних датасетах, що дозволяє адаптувати систему під конкретні потреби чи домени.
- Можливість масштабування системи під різні типи платформ та мовні домени.

Таким чином, запропонований спосіб можна розглядати як збалансоване та комплексне рішення, яке інтегрує сучасні методи глибинного навчання, підходи мультимодального аналізу та інтерпретовані алгоритми класичного машинного навчання. Така комбінація дозволяє ефективно поєднувати глибоке контекстуальне розуміння текстового контенту, обробку додаткових метаданих та прозорість результатів моделей. Завдяки цьому досягається висока точність прогнозування популярності соціального контенту, при цьому забезпечується гнучкість і універсальність системи, що дозволяє її адаптувати під різні типи даних та сценарії використання.

2.2. Теоретичні основи прогнозування потенційної популярності публікацій у соціальних мережах

Запропонований комбінований спосіб ґрунтується на ідеї синергії різних підходів до аналізу даних, де кожен із компонентів компенсує обмеження інших. Зокрема, класичні методи аналізу тексту забезпечують інтерпретованість та обчислювальну ефективність, тоді як глибинні нейронні моделі дозволяють виявляти приховані семантичні залежності, які важко формалізувати традиційними засобами.

2.2.1. Загальна концепція

У сучасному цифровому середовищі, де соціальні мережі відіграють важливу роль у комунікації, маркетингу та формуванні громадської думки, здатність точно прогнозувати популярність постів набуває стратегічного значення. Популярність посту зазвичай вимірюється кількістю взаємодій з ним – лайків, коментарів, репостів, охоплення або кількістю переглядів. Однак ці метрики не виникають випадково, а залежать від складної сукупності факторів.

Завдання прогнозування популярності посту належить до класу складних, слабо структурованих задач, які мають мультимодальний характер. Зокрема, на популярність впливають:

- Семантичний зміст публікації: наскільки вона цікава, релевантна, емоційна або провокаційна.
- Метадані: час публікації, хештеги, наявність згадок, формат посту.
- Соціальний контекст: популярність та активність автора, реакції попередніх постів, взаємодія аудиторії.
- Тренди та ситуативний контекст, що змінюється з часом (наприклад, актуальні події або мему).

У зв'язку з цим виникає потреба в побудові складних моделей, які можуть обробляти не лише текстові дані, а й числові та категоріальні метадані, інтегрувати додаткову інформацію з зовнішніх джерел.

З метою підвищення точності та узагальнюваності прогнозів пропонується спосіб, який базується на інтеграції трьох ключових компонентів:

- Семантичний аналіз тексту, що дозволяє виявити ключові теми, емоційне забарвлення, релевантність і новизну посту. Тут використовуються як класичні підходи (TF-IDF, LDA), так і глибокі нейронні представлення на основі трансформерів.
- Статистичне моделювання метаданих, зокрема часу публікації, кількості підписників, типу посту, використаних хештегів тощо. Для цього можуть використовуватись як прості регресійні моделі, так і ансамблеві методи (Random Forest, XGBoost), здатні враховувати нелінійні залежності.
- Контекстуально-обізнані трансформерні моделі, які забезпечують глибоке розуміння мовного контексту, враховуючи залежності між словами, реченнями, абзацами та навіть зовнішніми ознаками. Такі моделі, як BERT, RoBERTa, DeBERTa або DistilBERT, можна донавчати на специфічних наборах даних, щоб адаптувати їх до конкретної платформи.

2.2.2. Семантичний аналіз тексту

Семантичний аналіз тексту є ключовим компонентом у прогнозуванні популярності постів, оскільки дозволяє глибше зрозуміти не лише поверхневий зміст повідомлення, але й його внутрішній контекст, емоції, теми та інші смислові особливості. Основною метою семантичного аналізу є виділення смислових ознак, які можуть корелювати з популярністю контенту у соціальних мережах.

Перший і найважливіший крок – це підготовка тексту до подальшого аналізу. Текст у сирому вигляді часто містить зайві символи, помилки, неуніфіковані формати, які можуть негативно вплинути на якість подальших розрахунків.

До основних процедур належать:

- Очищення тексту від пунктуації, HTML-тегів, спеціальних символів, емодзі (за потреби) та інших артефактів.
- Нормалізація – приведення всіх символів до єдиного реєстру (зазвичай нижнього).
- Токенізація – розбиття тексту на окремі слова або токени.
- Лематизація – приведення слів до початкової (словникової) форми, що дозволяє узагальнювати різні граматичні форми одного слова (наприклад, "біжить", "бігти", "біг" – до леми "бігти").

Для кількісного представлення тексту у вигляді векторів використовуються різні методи:

- TF-IDF (Term Frequency-Inverse Document Frequency) – класичний статистичний підхід, який оцінює важливість слова у конкретному документі з урахуванням його частоти у всьому корпусі текстів. Слова, які часто зустрічаються в усіх документах, отримують менший ваговий коефіцієнт, що дозволяє виділити найбільш значущі терміни для конкретного посту.
- Word embeddings – представлення слів у вигляді щільних векторів у багатовимірному просторі (наприклад, Word2Vec, GloVe), що дозволяє врахувати семантичну близькість слів.
- Контекстуальні ембеддинги на основі трансформерів – більш сучасний підхід, який враховує контекст слова у реченні (BERT, RoBERTa). Ці ембеддинги є особливо корисними для розпізнавання багатозначності та більш точного розуміння тексту.

Одним з ключових аспектів семантичного аналізу є визначення емоційного забарвлення тексту, що часто впливає на реакцію аудиторії:

- Визначення тональності (позитивна, негативна, нейтральна) допомагає прогнозувати, як користувачі сприймуть пост.

- Використання як класичних лексиконів (наприклад, VADER, SentiWordNet), так і нейромережових моделей, які здатні розпізнавати складніші нюанси, такі як сарказм, іронія, інтенсивність емоцій.
- Врахування настрою дозволяє врахувати психологічний аспект сприйняття повідомлення.

Для розкриття прихованих тем у тексті застосовуються алгоритми тематичного моделювання, серед яких найбільш поширеним є:

- Latent Dirichlet Allocation (LDA) – статистичний байєсівський метод, який виявляє множину тем у колекції документів, описуючи кожен документ як суміш тем, а кожному темі – як розподіл слів.
- Тематичне моделювання дозволяє ідентифікувати основні напрямки, які зачіпає пост, що допомагає зрозуміти його потенційну аудиторію та тематику, яка може викликати більший інтерес.
- Крім LDA, можна застосовувати більш сучасні методи на основі нейронних мереж, такі як тематичні ембеддинги, які підвищують якість кластеризації тем.

Таким чином, семантичний аналіз тексту виступає фундаментальною складовою системи прогнозування популярності, дозволяючи комплексно оцінювати смислове наповнення посту. Поєднання класичних методів оброблення тексту з передовими нейромережевими моделями забезпечує високу точність і гнучкість у роботі з різноманітними типами текстових даних.

2.2.3. Статистичний аналіз метаданих

Метадані – це додаткова інформація про пост, що супроводжує основний текстовий чи мультимедійний контент і надає важливі сигнали для прогнозування популярності. На відміну від безпосереднього аналізу тексту, метадані дозволяють враховувати зовнішні чинники, які можуть

суттєво впливати на реакцію аудиторії та поширення посту у соціальних мережах.

Одним із ключових факторів, що впливають на популярність посту, є час його публікації. Різні часові проміжки протягом дня та тижня мають різну активність аудиторії:

- Аналіз часових рядів дозволяє виявити оптимальні часові вікна для публікації, коли максимальна кількість підписників активна.
- Застосовуються статистичні методи (наприклад, кореляційний аналіз, часові діаграми активності) для виявлення закономірностей.
- Врахування часових зон користувачів і особливостей платформи (наприклад, вихідні, святкові дні).
- Ці знання дозволяють планувати публікації з метою максимального охоплення аудиторії.

Впливовість автора посту часто безпосередньо корелює з його популярністю:

- Кількість підписників є базовою метрикою впливу, проте важливо враховувати й якість підписників (активність, реальність акаунтів).
- Враховуються також історичні дані про взаємодію аудиторії з публікаціями цього автора: середня кількість лайків, коментарів, репостів.
- Визначаються коефіцієнти залученості (engagement rate), які дозволяють краще оцінити активність аудиторії у відношенні до загальної кількості підписників.
- Під час моделювання популярності наведені параметри використовуються як числові ознаки для побудови регресійних або класифікаційних моделей.

Хештеги відіграють важливу роль у розповсюдженні контенту, оскільки вони допомагають формувати цільову аудиторію та збільшують видимість постів. У запропонованому способі хештеги розглядаються як один із важливих елементів метаданих, що відображає як тематичну

спрямованість публікації, так і потенційний рівень її залучення в інформаційний простір соціальної мережі. Їх аналіз дозволяє оцінити вплив правильно підібраних ключових маркерів на поширення контенту та взаємодію з користувачами. За допомогою хештегів у запропонованому способі:

- Аналізується частота використання конкретних хештегів у популярних постах.
- Визначається релевантність хештегів до тематики посту, що можна здійснювати через тематичне моделювання або семантичний аналіз.
- Досліджується взаємозв'язок між кількістю хештегів і популярністю, а саме надмірне або недостатнє їх використання може негативно впливати на охоплення.
- Також враховуються трендові та сезонні хештеги, які можуть значно підвищувати популярність у певні періоди.

2.2.4. Трансформерні моделі

Трансформерні моделі стали фундаментальною технологією сучасного оброблення природної мови (NLP), кардинально змінивши підходи до розуміння тексту та прогнозування на основі текстових даних. Архітектура трансформерів, запропонована Vaswani та співавторами у 2017 [7] році, базується на механізмі self-attention (самоуваги), що дозволяє моделі ефективно моделювати залежності між словами на різних відстанях у реченні, виходячи за межі традиційних рекурентних або згорткових мереж.

Однією з ключових переваг трансформерів є здатність моделі враховувати повний контекст слова в межах речення чи навіть документа. На відміну від попередніх методів, які працювали з фіксованими векторами слів (наприклад, word2vec або GloVe), трансформери генерують контекстуалізовані вектори, що динамічно змінюються в залежності від оточуючого тексту. Це особливо важливо для розпізнавання

багатозначності слів, розуміння складних синтаксичних конструкцій та інтенцій автора.

Популярні трансформерні моделі, такі як BERT, GPT, RoBERTa та інші, постачаються у вигляді попередньо навчених моделей на величезних корпусах текстів. Це дає змогу швидко адаптувати їх під конкретні задачі через процес дообучення (fine-tuning):

- Модель налаштовується на конкретні дані цільової задачі, наприклад, на тексти соціальних мереж із маркуванням популярності.
- Fine-tuning зазвичай виконується на меншому обсязі даних порівняно з повним навчанням з нуля, що значно скорочує час і ресурси.
- Це дозволяє отримати високу точність прогнозів, оскільки модель використовує як загальні знання мови, так і специфіку предметної області.

Ще однією важливою перевагою трансформерних моделей є їх гнучкість для інтеграції з різними типами даних, що особливо актуально для задач прогнозування популярності постів, які залежать не лише від тексту, але й від метаданих, зображень та інших джерел:

- Трансформери можна комбінувати з іншими нейронними мережами (наприклад, CNN для зображень) для створення мультимодальних моделей, які аналізують текст і зображення одночасно.
- Метадані (час публікації, авторські характеристики, хештеги) можуть подаватися як додаткові вхідні ознаки, що підвищує якість прогнозів.
- Архітектури з мультиголовою увагою (multi-head attention) дозволяють поєднувати різні потоки інформації, забезпечуючи комплексне представлення контенту.

У контексті прогнозування популярності постів трансформерні моделі розглядаються як один із найбільш перспективних підходів завдяки

їх здатності до глибокого контекстуального аналізу природної мови. Використання механізмів самоуваги дозволяє таким моделям враховувати складні семантичні та синтаксичні залежності у тексті, а також адаптуватися до різних стилістичних і тематичних особливостей контенту. Завдяки цьому трансформери можуть:

- Виділяти тонкі мовні патерни, що корелюють з емоційною реакцією аудиторії.
- Визначати релевантність та актуальність тематики поста.
- Підвищувати точність класифікації постів за категоріями популярності (наприклад, високий, середній, низький рівень залученості).
- Сприяти адаптивному моделюванню під різні платформи соціальних мереж, де лексика і стиль можуть значно відрізнятися.

2.3. Особливості роботи з багатомовними та різножанровими текстовими публікаціями

У сучасних соціальних мережах контент формується користувачами з різних країн, культур та мовних середовищ. Відповідно, для систем інтелектуального аналізу популярності постів надзвичайно важливо враховувати специфіку багатомовних та різножанрових текстів, оскільки мовні та стилістичні особливості можуть істотно впливати на точність моделювання.

2.3.1. Визначення проблеми багатомовності

У глобалізованому цифровому просторі соціальні мережі слугують майданчиком для комунікації мільйонів користувачів, що спілкуються різними мовами. Це породжує проблему багатомовності, яка є суттєвим викликом для систем автоматичного аналізу тексту, зокрема при прогнозуванні популярності контенту.

Багатомовність у межах одного набору даних означає наявність текстів, написаних різними мовами, з використанням різних алфавітів (латиниця, кирилиця, арабський шрифт, китайські ієрогліфи тощо), що ускладнює їх уніфіковану обробку.

Такий різноманітний лінгвістичний простір призводить до низки специфічних викликів, які варто розглянути детальніше:

- Граматичні, лексичні та синтаксичні відмінності. Кожна мова має унікальні граматичні правила, структури речень, порядок слів тощо.
- Морфологічна різноманітність. Існують флективні мови (наприклад, українська, польська), де форми слів змінюються шляхом додавання закінчень і афіксів, і аглютинативні мови (наприклад, турецька, фінська), де слова можуть містити довгі послідовності афіксів. Це вимагає різних підходів до лематизації та морфологічного розбору.
- Різна довжина слів і складність словотворення. У деяких мовах, як-от німецька, характерне утворення складних слів через об'єднання основ, що потребує спеціальних правил токенізації. У той час як у китайській мові слова часто складаються з одного-двох символів, але не мають явних роздільників між словами.
- Локальний сленг, жаргон, скорочення. У соціальних мережах широко використовуються неформальні мовні конструкції, емодзі, сленгові вирази, аббревіатури, специфічні для кожної мовної групи. Такі елементи можуть бути незрозумілими для загальних моделей для оброблення тексту або викликати помилки в класифікації та семантичному аналізі.
- Код-мешинг (code-mixing). Часто в одному пості можна зустріти фрагменти різних мов або навіть їх змішування в межах одного речення.

2.3.2. Особливості роботи з різножанровими текстами

У контексті аналізу популярності постів у соціальних мережах важливо враховувати не лише мовну, а й жанрову різноманітність текстів. Пости користувачів можуть належати до різних жанрових груп – від звичайного повсякденного спілкування до складних аналітичних чи креативних дописів. Кожен жанр має унікальні лінгвістичні, стилістичні та семантичні характеристики, що суттєво впливають на результати оброблення та прогнозування.

Основні особливості, пов'язані з жанровою різноманітністю:

- **Лексична варіативність.** Кожен жанр має свій словниковий запас. Наприклад, технічні тексти переважно використовують спеціалізовану термінологію, у той час як повсякденні – побутову лексику, часто неповну або усічену. Це ускладнює побудову універсальних моделей векторизації та семантичного аналізу, оскільки одна і та сама модель може неадекватно обробляти різні типи лексики.
- **Фразеологізми, ідіоми та метафори.** Особливо у креативних чи неформальних жанрах широко використовуються нестандартні мовні конструкції, які складно автоматично інтерпретувати. Наприклад, фраза *“влетіти в копійку”* має зовсім інше значення, ніж сума слів, які її складають. Моделі без контекстуального розуміння (як-от TF-IDF) не здатні вловити подібні смислові нюанси.
- **Складність sentiment-аналізу.** Емоційне забарвлення тексту може суттєво варіюватися залежно від жанру. У деяких випадках емоції виражено прямо (наприклад, в оглядах або коментарях), в інших – приховано (наприклад, у сатирі, іронії, сарказмі). Розпізнавання іронії, особливо без допомоги паралельних каналів (візуальних чи аудіо), є одним із найскладніших завдань в NLP.

- **Формальність та неформальність.** Формальні жанри (новини, технічні документи) мають структуровану мову, чіткі речення та передбачувану лексику. Неформальні тексти, навпаки, часто містять нестандартні скорочення (“OMG”, “IDK”), запозичення з інших мов, а також сленгові та регіональні варіації, що вимагають додаткових підходів до нормалізації.
- **Мета висловлювання.** Важливим є розуміння мети, з якою створено текст: інформувати, переконати, розважити тощо. Це безпосередньо впливає на його структуру та вміст.

2.3.3. Методи та підходи до роботи з багатомовними та різножанровими текстами

Зважаючи на численні виклики, що виникають під час оброблення багатомовних і різножанрових текстів, необхідно застосовувати спеціалізовані методи та адаптивні підходи. Їхня мета – забезпечити якісну обробку вхідних даних незалежно від мови, жанру чи формату висловлювання.

- **Уніфікована попередня обробка (preprocessing pipeline).** Ключовим етапом є побудова гнучкого конвеєра попереднього оброблення, який адаптується до мови та жанру тексту. На цьому етапі здійснюється автоматичне визначення мови публікації, що дозволяє застосовувати відповідні мовно-специфічні інструменти для токенизації та нормалізації. Додатково виконується уніфікація символів і емодзі, які часто несуть важливе семантичне або емоційне навантаження у соціальних мережах.
- **Використання мультимовних моделей трансформерів.** Для оброблення багатомовних даних доцільним є застосування мультимовних трансформерів, таких як Multilingual BERT, XLM-RoBERTa або LaBSE. Ці моделі дозволяють працювати з текстами різними мовами в єдиному семантичному просторі, забезпечуючи

узгоджене векторне представлення змісту. Окрему перевагу мають байт-орієнтовані моделі (mT5, ByT5), які ефективно працюють зі сленгом, нестандартною орфографією та малоресурсними мовами.

- Адаптація до жанрових особливостей тексту. Оскільки соціальні мережі містять тексти різних жанрів, доцільно враховувати жанрову специфіку під час моделювання. Це може включати автоматичну класифікацію жанру поста, доменну адаптацію трансформерів шляхом донавчання на спеціалізованих корпусах, а також використання гібридних підходів, які поєднують статистичні методи з глибинним семантичним аналізом.
- Переклад як альтернативний підхід (Translation-based pipeline). У випадках обмеженої мовної підтримки моделей може застосовуватися попередній автоматичний переклад текстів до однієї цільової мови, зазвичай англійської. Такий підхід спрощує подальший аналіз, однак має низку недоліків, зокрема втрату стилістичних та емоційних нюансів, можливі викривлення жанрових особливостей і збільшення обчислювальних витрат.

2.4. Спосіб оцінки потенційної популярності текстових постів соціальних мереж

Запропонований спосіб оцінки популярності текстових постів орієнтований на комплексний аналіз контенту та умов його публікації з метою отримання достовірних і практично значущих прогнозів. На відміну від спрощених підходів, які ґрунтуються виключно на історичних показниках взаємодії, даний спосіб враховує семантичні, соціальні та часові аспекти, що дозволяє краще моделювати поведінку користувачів у динамічному середовищі соціальних мереж.

2.4.1. Мета способу

Головною метою запропонованого способу є прогнозування майбутньої популярності публікації у соціальних мережах із високою точністю. Під популярністю розуміють кількість та якість взаємодій користувачів з контентом – лайки, репости, коментарі, перегляди. В залежності від задачі, спосіб може формувати:

- Категоріальні оцінки (низька, середня, висока популярність).
- Ймовірнісні оцінки трендів (наприклад, потенціал вірусного поширення).

Спосіб повинен забезпечувати можливість адаптації до різних соціальних платформ (Facebook, Instagram, Twitter, TikTok тощо), враховуючи специфіку кожної з них.

2.4.2. Архітектурна структура способу

Запропонований спосіб ґрунтується на послідовності етапів, на кожному з яких виконується:

1. Етап попередньої оброблення (Preprocessing Layer). Передача даних у «сиру» форму ускладнює їх аналіз через неоднорідність форматів та шум. Тому необхідно застосувати багатоетапну підготовку вхідних даних:
 - Ідентифікація мови тексту за допомогою моделей, що враховують багатомовність.
 - Очищення тексту від небажаних символів, посилань, зайвих пробілів.
 - Токенізація із врахуванням особливостей мови (для аглютинативних чи флективних мов це особливо важливо).
 - Лематизація для приведення слів до базових форм.
 - Оброблення емодзі та смайликів, які несуть емоційне навантаження.
 - Визначення та вилучення хештегів, згадок.

- Аналіз та кодування метаданих (час, місце, авторство).

Для багатомовного контенту застосовуються мультимовні трансформери, які дозволяють уніфікувати подальшу обробку.

2. Етап виділення ознак (Feature Extraction Layer). На основі очищених даних формуються числові та категоріальні ознаки, які відображають різні аспекти посту:

- Текстові ознаки: семантичні вектори, отримані через сучасні трансформери (BERT, mBERT, XLM-R), що забезпечують глибоке розуміння контексту, ознаки сентименту, отримані за допомогою спеціалізованих моделей аналізу настрою, тематичні ознаки, сформовані через тематичне моделювання (LDA, BERTopic), що дозволяють виділити ключові теми, лексична складність, довжина тексту, наявність риторичних запитань, закликів до дії.
- Соціальні ознаки: історія взаємодій з постами автора, поведінкові патерни аудиторії, індекси впливовості (наприклад, PageRank для користувачів).
- Метадані: час публікації, перетворений на циклічні ознаки (година доби, день тижня), що допомагає врахувати часові патерни активності аудиторії, характеристика автора – кількість підписників, рівень взаємодії, кількість і тип хештегів, їх популярність на даний момент.

3. Етап агрегації та нормалізації (Fusion & Normalization Layer). Оскільки отримані ознаки різномірні за типом і масштабом, їх потрібно уніфікувати:

- Числові ознаки проходять масштабування (standartization або min-max).
- Категоріальні – кодуються через one-hot або embedding.
- Текстові вектори інтегруються у єдине просторове представлення.

- Застосовуються методи зниження розмірності (PCA, t-SNE) для виявлення найбільш інформативних ознак.
4. Прогностична модель (Prediction Layer). Це ядро системи, яке приймає підготовлений вектор ознак і генерує прогноз:
- У класичних рішеннях використовують бустингові моделі (XGBoost, LightGBM), які мають високу продуктивність на табличних даних.
 - Гібридні архітектури, що комбінують трансформерні текстові модулі з класичними бустинговими алгоритмами, дозволяють поєднати глибоке розуміння контексту із стабільністю та швидкістю.
5. Етап пояснення та інтерпретації (Explainability Layer). У контексті сучасних етичних вимог та практик машинного навчання, моделі повинні бути не лише ефективними, а й прозорими. Тому передбачено інтеграцію методів інтерпретації:
- SHAP і LIME дозволяють візуалізувати вплив кожної ознаки на прогноз.
 - Графічні інструменти демонструють, які характеристики (наприклад, час публікації чи наявність зображення) найбільше вплинули на оцінку популярності.
6. Етап постоброблення та рекомендацій (Postprocessing & Recommendation Layer). Після отримання прогнозу модуль формує корисні для користувача або аналітика результати:
- Відображення конкретного числового прогнозу та категорії популярності.
 - Рекомендації щодо покращення посту (зміна часу публікації, додавання візуального контенту, вибір хештегів).
 - Побудова динамічних графіків популярності у часі для відстеження ефекту.

2.5. Висновки до розділу 2

У цьому розділі було запропоновано та обґрунтовано спосіб для прогнозування потенційної популярності текстових публікацій соціальних мереж. Проведений аналіз дозволив виявити ключові аспекти, які слід враховувати при побудові ефективної моделі прогнозування, а також окреслити обмеження і виклики, властиві роботі з текстовими даними соціальних платформ.

Було обґрунтовано вибір комбінованого підходу, що поєднує семантичний аналіз текстового вмісту, статистичний аналіз метаданих та використання сучасних трансформерних моделей, здатних обробляти багатовимірну та контекстно залежну інформацію. Така інтеграція дозволяє враховувати як глибинний зміст тексту, так і зовнішні фактори, що впливають на популярність контенту, забезпечуючи більш точне та комплексне прогнозування.

У рамках семантичного аналізу були визначені основні етапи оброблення текстових даних, включаючи лінгвістичну нормалізацію, побудову ознак, тематичне моделювання та емоційний аналіз.

Окрему увагу було приділено викликам, пов'язаним із багатомовністю та різножанровістю текстів, які є типовими для соціальних мереж. Було проаналізовано специфіку оброблення таких текстів і запропоновано способи, що дозволяють забезпечити їх уніфіковану інтерпретацію незалежно від мови чи стилю, що підвищує узагальнюваність моделі.

Таким чином, отримані результати створюють міцну основу для реалізації прикладного програмного рішення, яке буде детально розглянуто у наступному розділі. Вони підкреслюють важливість комплексного способу, що поєднує аналіз текстового контенту та метаданих, і демонструють перспективність інтеграції сучасних трансформерних моделей для підвищення точності прогнозування популярності постів у соціальних мережах.

3. ПРОГРАМНА РЕАЛІЗАЦІЯ СИСТЕМИ ІНТЕЛЕКТУАЛЬНОГО АНАЛІЗУ ПОПУЛЯРНOSTІ ПОСТІВ У СОЦІАЛЬНИХ МЕРЕЖАХ

3.1. Обґрунтування вибору засобів реалізації

Вибір засобів програмної реалізації системи інтелектуального аналізу популярності постів у соціальних мережах здійснювався з урахуванням вимог до точності оброблення даних, продуктивності, масштабованості та зручності подальшого розширення функціоналу. Особлива увага приділялась можливості інтеграції сучасних методів обробки природної мови, машинного навчання та аналізу великих обсягів текстових і статистичних даних.

3.1.1. Архітектурний підхід та тип реалізації

У рамках даного дослідження було обрано модульну архітектуру системи, яка забезпечує розділення логіки оброблення тексту, аналізу метаданих, генерації ознак та побудови прогнозів у вигляді окремих компонентів. Такий підхід дозволяє досягти високої гнучкості, масштабованості та повторного використання окремих частин системи у майбутніх дослідженнях або в інших проєктах.

Реалізація була здійснена у вигляді консольного Python-застосунку, що працює локально, з можливістю подальшого розгортання у хмарному середовищі або інтеграції у вебсервіс (через REST API або microservice-архітектуру). Вибір такої форми реалізації зумовлений наступними факторами:

- Гнучкість налаштування та розширення: завдяки поділу системи на незалежні логічні блоки (модулі), можливо легко змінювати або оновлювати окремі компоненти, не порушуючи загальну структуру.
- Локальне тестування та відлагодження: консольна реалізація дозволяє ефективно тестувати моделі та модулі оброблення даних

на окремих прикладах без потреби у розгортанні серверної інфраструктури.

- Простота інтеграції з іншими системами: завдяки модульному дизайну можлива подальша інкапсуляція компонентів у вигляді бібліотек або контейнерів (наприклад, Docker) для використання в більш складних програмних середовищах.

Загальна структура архітектури передбачає наявність наступних логічних рівнів, які послідовно забезпечують повний цикл оброблення даних – від отримання сирової інформації до формування та інтерпретації результатів прогнозування:

- Рівень завантаження та попереднього оброблення даних. Включає імпорт сирих даних (JSON, CSV, API), попередню обробку тексту (очищення, нормалізація, токенізація), обробку метаданих.
- Рівень ознакового простору (feature engineering). Формування векторів ознак на основі тексту (TF-IDF, word embeddings), метаданих (час, автор, кількість підписників), та агрегованих характеристик.
- Рівень побудови моделей. Навчання моделей машинного навчання (Random Forest, Gradient Boosting, BERT), їх збереження та повторне використання.
- Рівень інтерпретації результатів. Вивід прогнозу популярності, побудова графіків, обчислення метрик ефективності, генерація звітів.

Таким чином, обраний підхід дозволяє забезпечити ефективну реалізацію системи з урахуванням сучасних вимог до модульності, масштабованості та відтворюваності наукових експериментів. У разі потреби реалізацію можна розширити шляхом інтеграції з вебінтерфейсом або платформами для оброблення великих обсягів даних, такими як Apache Spark чи AWS SageMaker.

3.1.2. Мова програмування Python: аргументація вибору

Для реалізації програмного забезпечення було обрано мову програмування Python, яка є одним із найпоширеніших інструментів у сфері оброблення природної мови (Natural Language Processing, NLP), аналізу даних та побудови моделей машинного навчання. Такий вибір зумовлений кількома ключовими причинами:

- Python має широку екосистему спеціалізованих бібліотек для оброблення тексту (NLTK, spaCy), векторизації (scikit-learn, Gensim), моделювання (XGBoost, LightGBM), а також для роботи з трансформерами (Hugging Face Transformers, TensorFlow, PyTorch). Це значно пришвидшує розробку та дозволяє зосередитися на прикладній частині задачі.
- Завдяки простому та читабельному синтаксису Python є надзвичайно ефективним для швидкого створення прототипів моделей, гнучкого тестування гіпотез та побудови експериментів у наукових дослідженнях.
- Python тісно інтегрується з такими інструментами, як Jupyter Notebook, що дозволяє в інтерактивному середовищі аналізувати результати, будувати графіки (через Matplotlib, Seaborn, Plotly) та формувати наукову документацію.
- Python є основною мовою розробки в бібліотеках, що реалізують найновіші архітектури трансформерів (BERT, GPT, RoBERTa) [3], що критично важливо для реалізації комбінованого підходу, описаного у попередніх розділах.
- Незважаючи на інтерпретовану природу, Python дозволяє реалізовувати продуктивні рішення за рахунок використання оптимізованих внутрішніх модулів (написаних на C/C++) та можливості паралельного обчислення, зокрема через бібліотеки joblib, multiprocessing або GPU-прискорення у TensorFlow/PyTorch.

- Велика спільнота Python-розробників сприяє швидкому пошуку рішень, наявності великої кількості документації та прикладів реалізації, що прискорює процес розробки і зменшує ризики виникнення непередбачених помилок.

Таким чином, Python повністю відповідає вимогам даної задачі як у частині наукового дослідження, так і в аспекті створення продуктивного програмного забезпечення.

3.1.3. Використання бібліотек для NLP: spaCy, NLTK, Transformers

У реалізації програмного забезпечення для інтелектуального аналізу популярності текстового контенту критичною є роль бібліотек, що забезпечують обробку природної мови. Було обрано кілька перевірених та ефективних рішень – SpaCy, NLTK та Transformers від Hugging Face, кожне з яких виконує специфічні функції в загальній системі.

SpaCy – це сучасна та високооптимізована бібліотека для оброблення природної мови, орієнтована на промислове використання. Основні її переваги:

- Реалізована на Cython, що забезпечує високу продуктивність.
- Підтримка лематизації, POS-тегінгу, іменованих сутностей (NER), що дозволяє витягати лінгвістичну структуру тексту.
- Можна конфігурувати власні NLP-процеси та інтегрувати кастомні компоненти.

У даному проєкті SpaCy використовується для попереднього оброблення тексту (токенізація, лематизація), виявлення частин мови, витягу іменованих сутностей та формування словникових ознак.

Natural Language Toolkit (NLTK) – одна з найстаріших та найповніших бібліотек для досліджень у сфері NLP. Вона застосовується у проєкті для:

- Обчислення лінгвістичних метрик (частота слів, унікальність, складність).
- Побудови частотних розподілів.

- Підрахунку n-грам (біграм, триграм).
- Взаємодії з лексиконами (наприклад, для тонального аналізу).

Хоча NLTK є менш ефективною в обчислювальному сенсі порівняно зі SpaCy, вона пропонує більшу кількість інструментів для глибокого аналізу тексту.

Бібліотека Transformers від Hugging Face стала де-факто стандартом для реалізації сучасних моделей на основі трансформерної архітектури завдяки широкому функціоналу, активній підтримці спільноти та високому рівню абстракції. Вона надає зручні інструменти для завантаження, донавчання та використання попередньо навчених моделей, що дозволяє ефективно застосовувати їх у прикладних задачах аналізу тексту, а саме:

- BERT (Bidirectional Encoder Representations from Transformers) – для витягу контекстуально значущих векторів (ембедінгів).
- RoBERTa, DistilBERT, GPT – для задач класифікації, прогнозування чи генерації.
- Multilingual BERT (mBERT) – для багатомовного оброблення.

Переваги використання цієї бібліотеки:

- Можливість використання *populated pretrained models*, що значно скорочує час навчання.
- Підтримка *fine-tuning* на власних наборах даних.
- Широка підтримка різних форматів введення, а також інтеграція з PyTorch і TensorFlow.

У межах способу Transformers застосовується для генерації векторних репрезентацій постів, що потім подаються до моделі прогнозування популярності. Також бібліотека використовується для класифікації сентименту на базі трансформера, що дозволяє враховувати глибокий контекст фраз.

3.1.4. Інструменти оброблення метаданих: Pandas, NumPy, datetime

Обробка метаданих – одного з ключових компонентів прогнозування популярності текстового контенту – потребує ефективних засобів для роботи з табличними даними, часовими мітками, статистичними та агрегованими обчисленнями. У межах реалізації програмного забезпечення було обрано три основні інструменти: Pandas, NumPy та модуль datetime. Ці бібліотеки зарекомендували себе як надійні та потужні засоби для аналітики, що добре інтегруються з екосистемою Python.

Pandas є стандартом для роботи з табличними (структурованими) даними у Python. У проєкті вона використовується для:

- Структурування метаданих, тобто представлення постів як рядків DataFrame зі стовпцями, що відповідають окремим ознакам.
- Агрегації часових даних (наприклад, групування за днями тижня або годинами доби).
- Обчислення статистичних показників (середнє значення, медіана, стандартне відхилення) для ознак автора або повідомлення.
- Перетворення часових міток у зручний формат для подальшого аналізу (наприклад, перетворення ISO8601 у часові індекси).

NumPy (Numerical Python) використовується переважно для:

- Векторизованих обчислень над числовими ознаками.
- Нормалізації числових метрик (наприклад, приведення кількості підписників до інтервалу $[0,1]$).
- Перетворення ознак у матриці для подачі до моделі машинного навчання.
- Статистичного аналізу (коваріація або кореляція між ознаками).

Завдяки векторизованим операціям та низькорівневій реалізації на C, NumPy дозволяє виконувати обчислення швидко та ефективно, що особливо важливо у випадку великомасштабного оброблення даних.

Модуль `datetime` (вбудований у стандартну бібліотеку Python) використовується для роботи з:

- Часовими мітками постів (перетворення з рядкових форматів у об'єкти `datetime.datetime`).
- Виділення часових компонентів (наприклад, дня тижня, години, місяця).
- Порівняння та обчислення дельт часу між подіями (наприклад, інтервали між публікаціями).
- Генерації часових ознак, які можуть мати вплив на популярність (наприклад, пости у вечірній час можуть отримувати більше реакцій).

3.1.5. Інструменти оцінки якості моделей: `scikit-learn metrics`, `matplotlib`, `seaborn`

Оцінка ефективності моделей прогнозування є критично важливою складовою етапу валідації, яка дозволяє визначити не лише абсолютну точність моделі, а й глибше зрозуміти її поведінку, виявити потенційні упередження або зони нестабільності.

Для задач класифікації (наприклад, визначення категорії популярності):

- `Accuracy_score` – частка правильних передбачень.
- `Precision_score`, `Recall_score`, `F1_score` – для врахування як точності, так і повноти.
- `Roc_auc_score` – для оцінки роздільної здатності моделі.
- `Confusion_matrix` – для побудови матриці помилок і аналізу типів класифікаційних помилок.

Для задач регресії (наприклад, передбачення кількості лайків/репостів):

- `Mean_absolute_error` (MAE), `Mean_squared_error` (MSE), `Root_mean_squared_error` (RMSE).

- `R2_score` – коефіцієнт детермінації, що характеризує якість апроксимації.

Завдяки уніфікованому інтерфейсу `scikit-learn`, ці метрики легко інтегруються в пайплайни моделювання, забезпечуючи прозорість і повторюваність експериментів.

Бібліотека `matplotlib` використовується для побудови базових графіків і візуального аналізу результатів моделювання. Основні варіанти застосування:

- Побудова графіків помилок, тобто порівняння передбачених та фактичних значень (`y_pred` vs `y_true`).
- ROC-криві для класифікаційних моделей.
- Гістограми розподілу похибок.
- Візуалізація динаміки метрик на етапах тренування/валідації.

Ця бібліотека є надійним інструментом для статичної графіки з високим рівнем контролю за оформленням.

`Seaborn` – це бібліотека, що побудована поверх `matplotlib` і дозволяє створювати більш інформативні та естетично привабливі графіки з мінімальними зусиллями. Вона активно застосовувалася для:

- Візуалізації кореляційної матриці ознак (`sns.heatmap`).
- Побудови `boxplot` і `violinplot` для аналізу розподілу ознак або результатів.
- Графіків розсіювання (`scatterplot`) з трендовими лініями (`regplot`).
- Порівняння моделей за кількома метриками одночасно.

Комбінація `matplotlib` і `seaborn` забезпечує повноцінну підтримку діагностичного аналізу моделей, дозволяючи побачити слабкі місця, наприклад, систематичні зсуви, великі варіації тощо.

3.2. Архітектура розробленого програмного забезпечення

Архітектура розробленого програмного забезпечення визначає загальні принципи організації системи, взаємодію її основних компонентів та послідовність оброблення даних на всіх етапах аналізу. При проєктуванні архітектури основний акцент було зроблено на забезпеченні модульності, масштабованості та гнучкості системи, що є критично важливим для роботи з великими обсягами різнорідних даних соціальних мереж.

3.2.1. Загальна структура системи: компоненти та їх взаємодія

Розроблене програмне забезпечення реалізовано у вигляді багаторівневої модульної системи, призначеної для інтелектуального аналізу та прогнозування популярності текстових постів у соціальних мережах із урахуванням як лінгвістичних характеристик контенту, так і контекстних параметрів його публікації. Запропонована структура забезпечує послідовну обробку даних на різних рівнях абстракції, що дозволяє здійснювати перехід від сирих текстових і метаданих до формалізованого числового представлення.

Система побудована відповідно до принципу розділення відповідальностей, який передбачає чітке логічне відокремлення функціональних компонентів, відповідальних за попереднє оброблення даних, виділення, інтеграцію та нормалізацію ознак, побудову і застосування моделей машинного навчання, а також інтерпретацію результатів прогнозування та формування рекомендацій.

Багаторівнева модульна структура створює передумови для інтеграції нових методів оброблення природної мови, альтернативних алгоритмів машинного та глибинного навчання. Завдяки цьому запропоноване програмне рішення може бути адаптоване до різних соціальних платформ, змін у форматах даних та зростання обсягів інформації.

На рис. 3.1 представлено узагальнену структуру запропонованого способу.

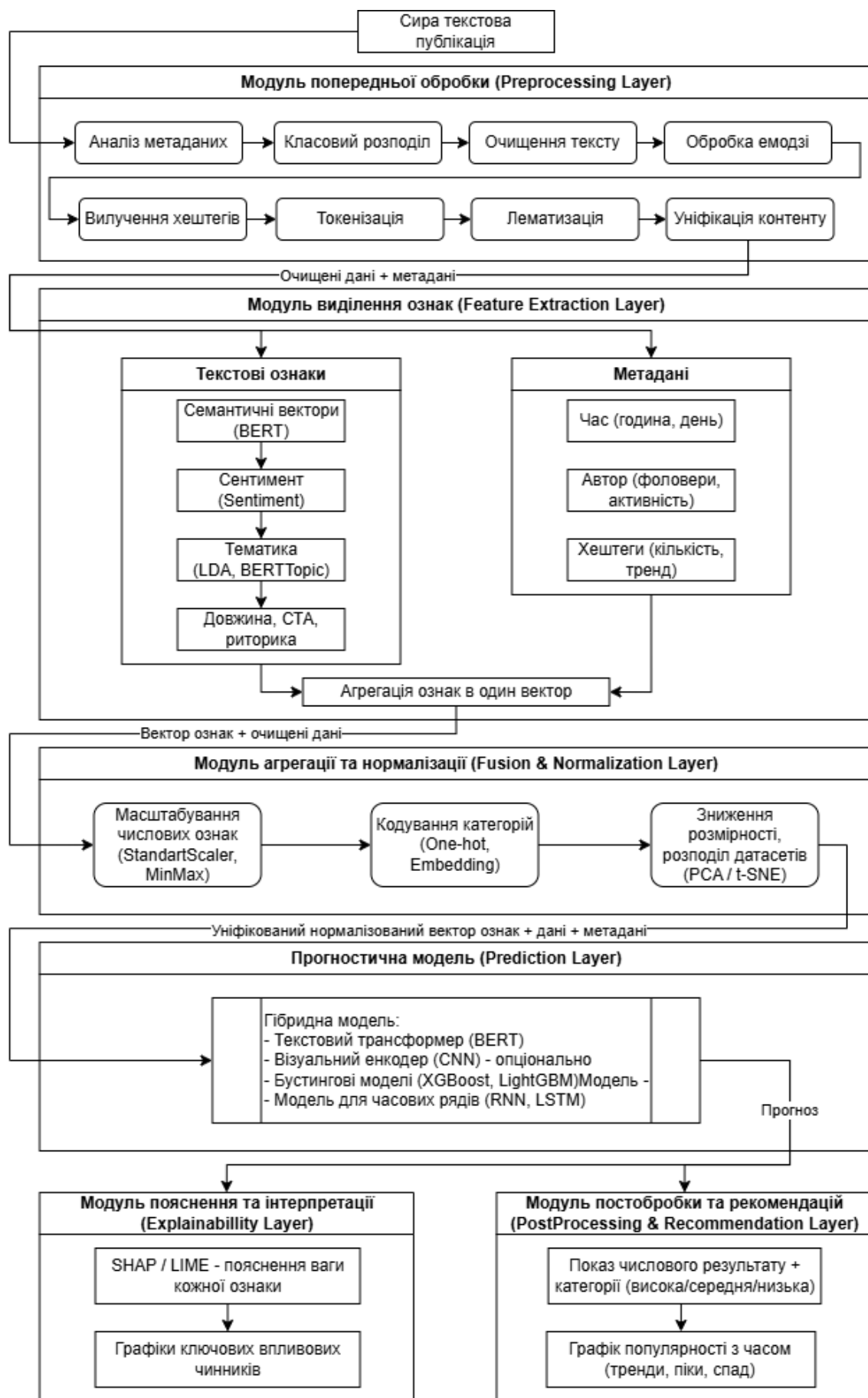


Рис. 3.1. Узагальнена структура запропонованого способу

3.2.2. Модулі оброблення тексту, метаданих і генерації ознак

У запропонованій структурі програмного забезпечення особливу роль відіграють модулі, що відповідають за оброблення текстових даних, метаданих та побудову векторів ознак. Саме ці компоненти формують основу для побудови машинного уявлення про кожен текстовий пост і забезпечують підґрунтя для подальшого прогнозування популярності.

Модуль попереднього оброблення забезпечує первинну лінгвістичну обробку текстових даних, що передує етапу їхньої векторизації та подальшого машинного аналізу. На початковому етапі виконується очищення тексту, що включає видалення HTML-тегів, емодзі, спеціальних символів, URL-посилань та інших елементів, які не несуть суттєвої аналітичної цінності. Після цього здійснюється нормалізація, спрямована на уніфікацію текстових даних шляхом приведення слів до нижнього регістру та розгортання скорочених форм, таких як “don’t → do not”, що дозволяє зменшити варіативність у лексичному просторі.

Подальшим етапом є токенізація, яка передбачає сегментацію тексту на окремі лексичні або синтаксичні одиниці, що формують основу подальшої обробки. Для забезпечення лексичної узгодженості застосовуються процедури лематизації або стемінгу, які приводять слововформи до їхньої базової форми або кореня. Після цього виконується фільтрація стоп-слів з метою усунення високочастотних елементів, які не мають вагомого семантичного значення й не сприяють підвищенню інформативності моделі.

Для реалізації зазначених процедур застосовуються такі інструменти, як spaCy, NLTK, TextBlob та бібліотека transformers, що забезпечує широкі можливості для роботи з передтренованими моделями глибинного навчання в контексті задач обробки природної мови.

Структуру та послідовність основних етапів попереднього оброблення текстових публікацій наведено на рис. 3.2.

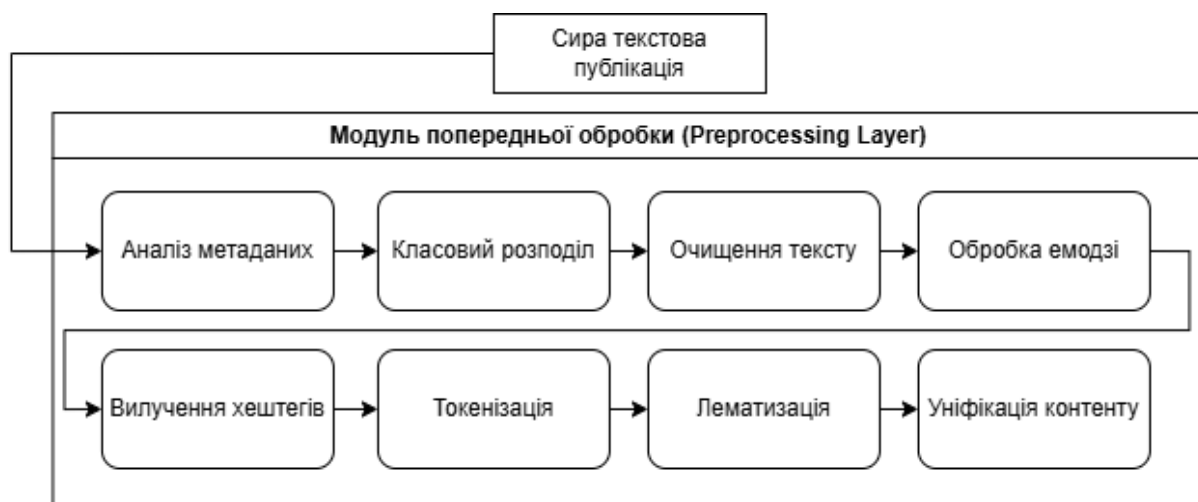


Рис. 3.2. Послідовність етапів попереднього оброблення текстових публікацій

Модуль виділення ознак призначений для трансформації підготовлених текстових даних і метаданих у числове подання, придатне для подальшого використання в алгоритмах машинного навчання. Основним завданням цього модуля є формування інформативного вектора ознак, який узагальнює як семантичні характеристики текстового контенту, так і контекстні параметри публікації.

У межах модуля здійснюється векторизація тексту з використанням як класичних статистичних методів, зокрема TF-IDF та Count Vectorizer, так і сучасних контекстних підходів, що базуються на попередньо навчених мовних моделях, таких як BERT і Word2Vec. Застосування зазначених методів дозволяє відобразити лінгвістичний зміст тексту у вигляді багатовимірних числових векторів, які зберігають інформацію про семантичні зв'язки між словами та фразами (semantic relationships).

Окрім семантичного представлення, у процесі виділення ознак обчислюються додаткові лінгвістичні метрики, зокрема довжина тексту, кількість унікальних лексичних одиниць, а також показники читабельності, такі як індекс Flesch–Kincaid. Зазначені характеристики дозволяють врахувати структурну складність тексту та особливості його подачі, що може впливати на сприйняття публікації користувачами.

Важливим компонентом модуля є інтеграція результатів сентимент-аналізу, які відображають емоційне забарвлення текстового повідомлення. Отримані показники тональності, зокрема позитивної, негативної або нейтральної, використовуються як додаткові ознаки та сприяють врахуванню емоційного впливу контенту на аудиторію.

Окрему групу становлять ознаки, сформовані на основі метаданих публікації. Для категоріальних змінних застосовуються методи кодування, зокрема one-hot encoding або векторні подання типу embedding, що забезпечує їх сумісність з моделями машинного навчання. Також здійснюється агрегування статистичних характеристик, таких як середні значення, медіани та частоти використання окремих токенів або атрибутів, що дозволяє узагальнити інформацію про публікацію.

Часові характеристики обробляються з урахуванням їх циклічної природи шляхом перетворення моменту публікації у відповідні циклічні ознаки, наприклад із використанням тригонометричних функцій синуса та косинуса для відображення періодичності добових або тижневих циклів. Такий підхід дає змогу моделям ефективніше враховувати часові закономірності в поведінці користувачів.

Оброблення метаданих також включає аналіз характеристик автора публікації, зокрема кількості підписників, рівня активності та середньої кількості взаємодій з попередніми постами, а також урахування мультимедійних компонентів, таких як емодзі.

Результатом роботи модуля виділення ознак є комплексний набір числових і категоріальних характеристик, який формує вхідний простір для моделей прогнозування популярності та забезпечує їх здатність комплексно враховувати як зміст, так і контекст текстових публікацій.

Загальну структуру процесу виділення ознак у межах запропонованої системи наведено на рис. 3.3.

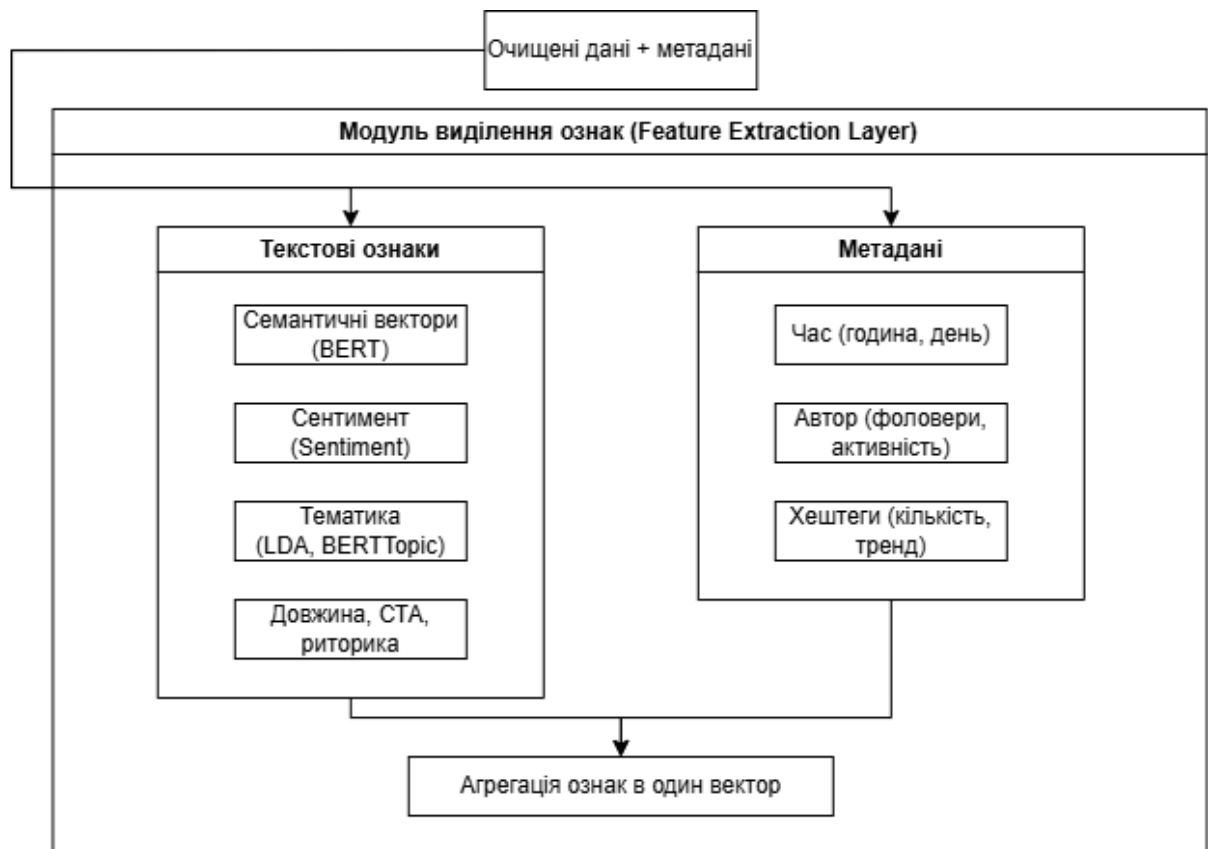


Рис 3.3. Структура модуля виділення ознак

3.2.3. Модуль агрегації і нормалізації та прогностична модель

Прогностична модель є ключовим компонентом розробленої системи, що відповідає за побудову, навчання та застосування моделей для прогнозування популярності текстових постів у соціальних мережах.

Перед безпосереднім тренуванням моделі, дані, сформовані в попередньому модулі агрегації та нормалізації, проходять процедури додаткового оброблення:

- Нормалізація та масштабування числових ознак (наприклад, StandardScaler або MinMaxScaler) для покращення збіжності моделей.
- Кодування категоріальних ознак (one-hot encoding, target encoding) для сумісності з алгоритмами машинного навчання.

- Розподіл вибірки на тренувальний, валідаційний та тестовий набори з урахуванням часових або тематичних розподілів для уникнення витоку даних.
- Балансування класів у разі класифікаційних задач (наприклад, за допомогою технік oversampling/undersampling або генеративних методів типу SMOTE) [4].

Послідовність етапів додаткового оброблення представлено на рис. 3.4.

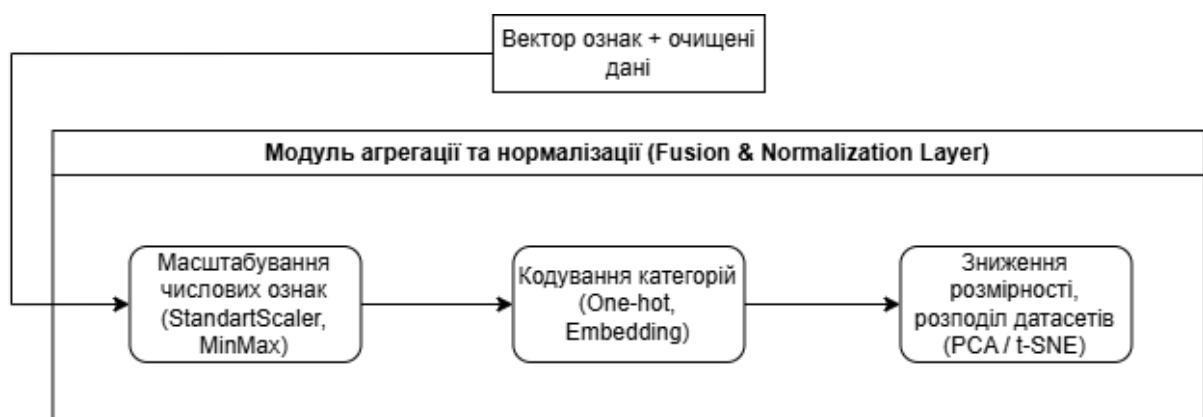


Рис. 3.4. Послідовність етапів додаткового оброблення

У рамках прогностичної моделі підтримується можливість роботи з різними алгоритмами, які відповідають специфіці задачі – регресії чи класифікації:

- Лінійні моделі (лінійна регресія, логістична регресія) для базових, інтерпретованих рішень.
- Дерева рішень і ансамблеві методи (Random Forest, Gradient Boosting, XGBoost, LightGBM) для підвищення точності з врахуванням нелінійних залежностей.
- Глибинні нейронні мережі (Feedforward, LSTM, Transformer-based) для складних семантичних задач і роботи з послідовними даними.
- Мультимодальні мережі, що інтегрують текстові, метадані та інші типи даних, для комплексного моделювання популярності.

Процес навчання супроводжується налаштуванням гіперпараметрів моделей, використанням методів крос-валідації для оцінки стабільності та уникнення перенавчання.

Для контролю якості навчання застосовуються метрики, описані у підрозділі 2.5. Результати навчання регулярно перевіряються на валідаційних даних для вибору оптимальної моделі.

Після навчання моделі використовуються для прогнозування популярності нових постів. Блок предикції приймає на вхід сформовані вектори ознак і повертає передбачені значення – кількість лайків, репостів, або клас популярності.

Для підвищення швидкодії передбачено можливість серіалізації (збереження) моделей із подальшим їх завантаженням у продуктивне середовище без повторного навчання.

Структуру модуля прогностичної моделі наведено на рис. 3.5.

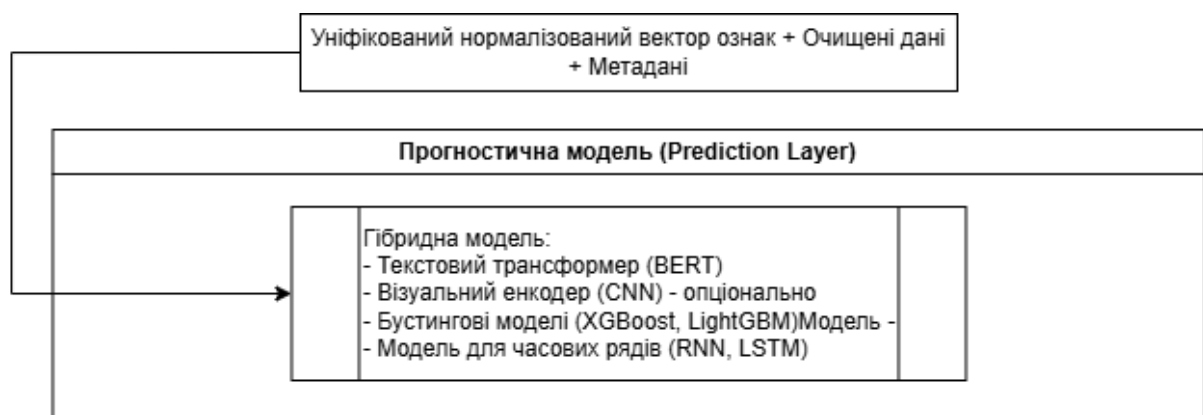


Рис. 3.5. Структура модуля прогностичної моделі

3.3. Реалізація основних функціональних компонентів системи

У цьому підрозділі описується конкретна реалізація ключових функціональних компонентів розробленої системи, які забезпечують комплексну обробку текстових даних, метаданих, генерацію ознак, навчання моделей та формування прогнозів популярності постів.

Модуль попереднього оброблення тексту відповідає за якісну підготовку текстових даних, що є критичним етапом для успішного подальшого аналізу. Реалізація включає очищення тексту від шумів, таких як спеціальні символи, HTML-теги, URL, нормалізацію тексту шляхом приведення до нижнього регістру та видалення зайвих пробілів, токенизацію – розбиття тексту на складові одиниці, лематизацію або стемінг для узагальнення форм слів, а також видалення стоп-слів, які не несуть змістового навантаження. Цей модуль реалізовано на основі бібліотек spaCy та NLTK, що забезпечують гнучкість і ефективність оброблення.

Модуль генерації ознак забезпечує векторизацію тексту за допомогою методів TF-IDF, CountVectorizer, а також використання більш складних ембедингів із моделей BERT чи Word2Vec. Крім того, він здійснює витяг лінгвістичних характеристик, таких як довжина тексту, кількість унікальних слів, наявність емодзі, хештегів і посилань. У модулі також формується набір ознак із метаданих – час публікації, кількість підписників автора, категорії контенту. Окрім цього, обчислюються семантичні метрики, наприклад, тональність тексту та результати тематичного моделювання. Для реалізації цього компоненту використовуються бібліотеки scikit-learn, transformers від Hugging Face, а також Pandas та NumPy для оброблення метаданих.

Модуль прогностичної моделі відповідає за побудову і навчання моделей на основі сформованих ознак. У ньому інтегруються різні алгоритми – від лінійних моделей до глибинних нейронних мереж. Реалізується налаштування параметрів і гіперпараметрів моделей, застосовується крос-валідація для оцінки якості і стабільності моделей. Для збереження та завантаження навчальних моделей використовуються такі інструменти, як scikit-learn, XGBoost, PyTorch та TensorFlow.

Модуль прогнозування забезпечує отримання прогнозів популярності для нових вхідних даних. Він приймає текстові повідомлення та відповідні метадані, автоматично виконує попередню обробку та генерацію ознак,

після чого застосовує навчану модель для прогнозування показників популярності. Результати формуються у зручному форматі, наприклад, JSON або CSV, що полегшує їх подальше використання.

Для підтримки стабільної роботи системи реалізовано модуль логування та моніторингу, який відповідає за ведення журналів процесів навчання та прогнозування, збір статистики щодо помилок та продуктивності моделей. Також цей модуль здійснює моніторинг якості прогнозів у режимі реального часу, що дозволяє оперативно реагувати на потенційні проблеми. Для цього використовуються бібліотеки logging, MLflow, а також інструменти візуалізації, такі як TensorBoard і Grafana.

3.4. Висновки до розділу 3

У третьому розділі було детально описано вибір засобів реалізації та архітектурних підходів для розробки програмного забезпечення, що забезпечує інтелектуальний аналіз популярності текстового контенту в соціальних мережах. Обґрунтовано використання мови програмування Python та ключових бібліотек і фреймворків, таких як spaCy, Hugging Face Transformers, scikit-learn, Pandas, NumPy та інші, які дозволяють ефективно обробляти як текстові дані, так і метадані.

Розглянуто загальну архітектуру системи, яка складається з модулів попереднього оброблення тексту, генерації ознак, машинного навчання, прогнозування та моніторингу якості роботи моделей. Такий модульний підхід забезпечує гнучкість, масштабованість та можливість подальшого розвитку розробленого програмного комплексу.

Особлива увага приділялась реалізації основних функціональних компонентів, які гарантують ефективне очищення та нормалізацію тексту, побудову інформативних ознак, а також навчання й застосування сучасних моделей машинного та глибинного навчання для точного прогнозування популярності публікацій.

4. АНАЛІЗ ЕФЕКТИВНОСТІ ЗАПРОПОНОВАНОГО СПОСОБУ

4.1. Методологія експериментального дослідження

Експериментальне дослідження є ключовим етапом валідації запропонованого комбінованого підходу до прогнозування популярності постів у соціальних мережах. Основною метою аналізу є об'єктивна оцінка ефективності розробленого способу та порівняння його з існуючими аналогами.

4.1.1. Мета та завдання експериментального дослідження

Для досягнення поставленої мети необхідно вирішити наступні завдання:

- Підготовка експериментальних даних – збір та структурування датасетів з платформи Reddit, розмітка даних за рівнями популярності, валідація якості та повноти даних.
- Порівняльний аналіз методів – тестування традиційних статистичних підходів, оцінка ефективності класичних алгоритмів машинного навчання, валідація трансформерних моделей, аналіз комбінованого підходу.
- Оцінка якості прогнозування – розрахунок метрик точності для різних підходів, статистичний аналіз значущості результатів, аналіз помилок та їх інтерпретація.

4.1.2. Експериментальна установка

Технічне середовище експериментів включає використання Python 3.12 як основної платформи розробки з підтримкою GPU для трансформерних моделей. Основними бібліотеками є scikit-learn для класичних алгоритмів машинного навчання, transformers для роботи з трансформерними моделями, pandas та numpy для оброблення даних.

Структура експериментів організована у чотири основні етапи:

- Baseline експерименти – тестування традиційних методів оброблення природної мови та статистичного аналізу.
- Advanced експерименти – застосування сучасних алгоритмів машинного навчання та глибинного навчання.
- Hybrid експерименти – валідація запропонованого комбінованого підходу.
- Ablation study – детальний аналіз внеску окремих компонентів системи.

4.1.3. Метрики оцінки ефективності

Для об'єктивної оцінки якості прогнозування використовується комплекс метрик, що дозволяють всебічно оцінити ефективність різних підходів.

Основними метриками є:

- Accuracy – загальна точність класифікації, що показує частку правильно класифікованих зразків.
- Precision, Recall, F1-score – детальна оцінка якості класифікації по окремих класах.
- ROC-AUC – оцінка якості бінарної класифікації.
- Confusion Matrix – детальний аналіз типів помилок класифікації.

Додатковими метриками для багатокласової задачі є Macro/Micro F1 для усередненої оцінки по всіх класах, Cohen's Kappa для оцінки узгодженості класифікації, та Processing Time для аналізу швидкості оброблення.

4.1.4. Валідаційна стратегія

Для забезпечення надійності результатів застосовується стратифікований розподіл даних: 70% для навчання моделей, 15% для налаштування параметрів та 15% для фінальної оцінки.

Для підвищення стабільності результатів використовується Stratified K-Fold cross-validation, що забезпечує збереження пропорцій класів у кожній вибірці. Для часових даних застосовується Temporal split, що враховує хронологічний порядок публікацій. Bootstrap sampling використовується для оцінки довірчих інтервалів метрик.

4.2. Підготовка та характеристика експериментальних даних

Для проведення експериментального дослідження було обрано платформу Reddit як репрезентативне джерело даних соціальних мереж. Вибір обумовлений наступними факторами: високою активністю користувачів, різноманітністю тематичного контенту, наявністю детальних метаданих та відкритістю API для дослідницьких цілей.

Експериментальний датасет складається з 100,000 текстових публікацій платформи Reddit та охоплює тематики технологій, новин, розваг, науки та спорту, що забезпечує репрезентативність вибірки.

Кожен запис у датасеті містить наступні компоненти:

- Текстові характеристики:
 - Заголовок посту (title) – основний текстовий контент для аналізу.
 - Основний текст (content) – додатковий текстовий матеріал.
- Метадані:
 - Час публікації (timestamp) – для часового аналізу.
 - Ідентифікатор автора (author_id) – для аналізу впливовості.
 - Категорія контенту (subreddit) – для тематичної класифікації.
- Цільові змінні:
 - Кількість вподобань (score) – основна метрика популярності.
 - Категорія популярності (popularity_class) – цільова змінна для класифікації.

Процес попереднього оброблення включає декілька етапів очищення та стандартизації даних. Очищення тексту передбачає видалення

HTML-тегів та URL-адрес, нормалізацію емодзі та спеціальних символів, фільтрацію спаму та дублікатів, валідацію мовної приналежності.

Стандартизація метаданих включає уніфікацію часових міток до UTC, нормалізацію категорій subreddit, обробку пропущених значень та кодування категоріальних змінних.

Для задачі класифікації було визначено три категорії популярності на основі кількості вподобань:

- Низька популярність – менше 50 вподобань.
- Середня популярність – 50-500 вподобань.
- Висока популярність – більше 500 вподобань.

Розподіл класів у датасеті: 60% постів з низькою популярністю (60,000 записів), 30% з середньою популярністю (30,000 записів) та 10% з високою популярністю (10,000 записів).

4.2.1. Підготовка до аналізу

Для проведення експериментів застосовано стратифікований розподіл даних: 70% для навчання (70,000 записів), 15% для валідації (15,000 записів) та 15% для тестування (15,000 записів).

Дані збережено у форматі Parquet для забезпечення ефективності оброблення, використано версіонування через Git LFS для великих файлів, створено детальний опис схеми даних та забезпечено резервне копіювання у хмарному сховищі.

Підготовлені експериментальні дані забезпечують надійну основу для об'єктивної оцінки ефективності запропонованого комбінованого підходу та порівняння його з існуючими методами прогнозування популярності постів у соціальних мережах.

4.3. Порівняльний аналіз ефективності різних способів

Порівняльний аналіз ефективності різних способів прогнозування популярності постів проводився з метою оцінки їх точності, стабільності та здатності до узагальнення на реальних даних соціальних мереж. Особлива увага приділялась аналізу впливу різних підходів до представлення тексту, використання метаданих та рівня складності моделей на кінцеві результати прогнозування.

4.3.1. Досліджувані способи

У рамках експериментального дослідження було протестовано п'ять основних категорій способів:

- Baseline методи включають традиційні статистичні підходи: TF-IDF з логістичною регресією, Bag-of-Words з наївним байєсівським класифікатором, та N-gram аналіз з методом опорних векторів (SVM).
- Класичні алгоритми машинного навчання представлені Random Forest з інженерними ознаками, градієнтним бустингом (XGBoost, LightGBM), та методом опорних векторів з RBF ядром.
- Методи глибинного навчання включають LSTM для послідовного аналізу тексту, згорткові нейронні мережі (CNN) для виявлення локальних текстових патернів, та attention-based моделі.
- Трансформерні моделі представлені BERT-base з fine-tuning, які забезпечують глибоке контекстуальне розуміння тексту і дозволяють досягати високої точності прогнозів. Вони особливо ефективні при інтеграції семантичних залежностей та забезпечують можливість донавчання під конкретну задачу.

Запропонований спосіб інтегрує семантичний аналіз через BERT embeddings, статистичний аналіз метаданих Reddit, та ensemble різних компонентів.

4.3.2. Результати експериментів

Експериментальне тестування на датасеті Reddit продемонструвало значні відмінності в ефективності різних способів. Результати представлено на рис. 4.1.

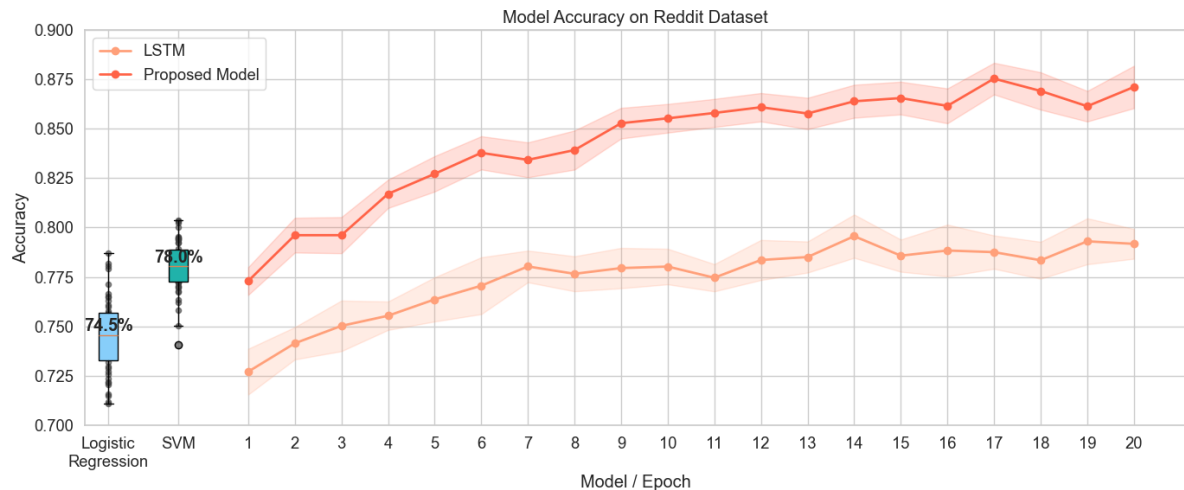


Рис. 4.1. Графік точності способів на тестовому датасеті

Найвищі результати демонструє запропонований спосіб з F1-score 0.87, що на 8% перевищує найкращий базовий.

4.3.3. Аналіз кореляційної структури ознак

На етапі підготовки даних проведено кореляційний аналіз та створено кореляційну матрицю трансформованих ознак для валідації якості feature engineering та забезпечення оптимальних умов для навчання моделей машинного навчання.

Кореляційну матрицю побудовано з використанням коефіцієнта кореляції Пірсона для шести ключових предикторів на тренувальній вибірці (70,000 записів):

- Логарифм довжини тексту (text_length_log).
- Квадратний корінь кількості слів (word_count_sqrt).
- Бінарна ознака наявності хештегів (hashtag_binary).

- Нормалізована щільність емодзі (emoji_normalized).
- Логарифм кількості підписників (followers_log).
- Абсолютне значення тональності (sentiment_abs).

На рис. 4.2 представлено кореляційну матрицю оброблених ознак системи прогнозування популярності текстових публікацій.

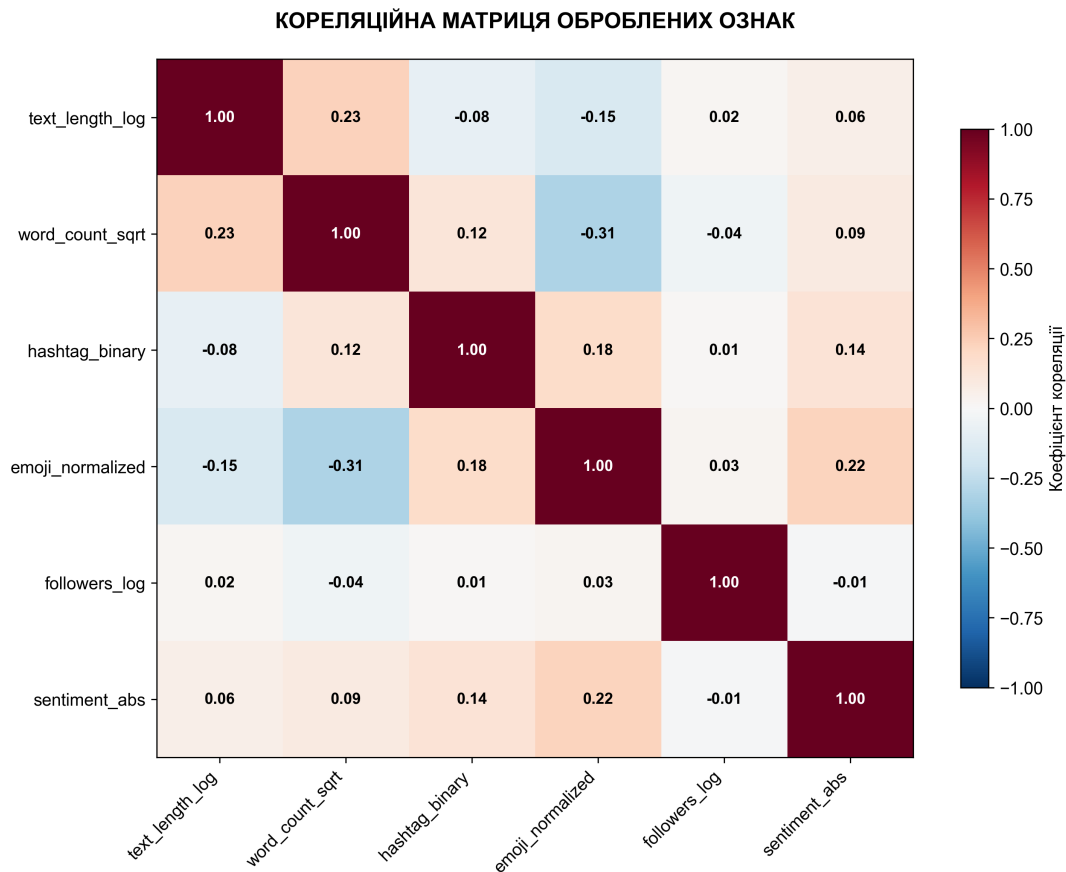


Рис. 4.2. Кореляційна матриця оброблених ознак

Виявлено слабку позитивну кореляцію між логарифмом довжини тексту та квадратним коренем кількості слів, що відображає природний зв'язок між фізичним обсягом контенту та його лексичним наповненням. Водночас спостерігається слабка негативна кореляція між довжиною тексту та щільністю емодзі, що свідчить про тенденцію до зменшення відносної кількості емодзі у більш розгорнутих публікаціях.

Найвищий за абсолютним значенням коефіцієнт кореляції зафіксовано між кількістю слів та нормалізованою щільністю емодзі. Цей

результат підтверджує компенсаційний характер використання емодзі, тобто автори коротких повідомлень компенсують обмеженість вербального вираження підвищеною концентрацією візуальних елементів.

Встановлено позитивні кореляції між експресивними елементами тексту, а саме наявність хештегів корелює зі щільністю емодзі та абсолютним значенням тональності. Ці результати вказують на існування узгодженого стилістичного патерну використання експресивних засобів у соціальних мережах.

4.3.4. Аналіз структури векторних представлень BERT

На рис. 4.3 представлено t-SNE візуалізацію структури векторних представлень BERT, розподілених за рівнями популярності.

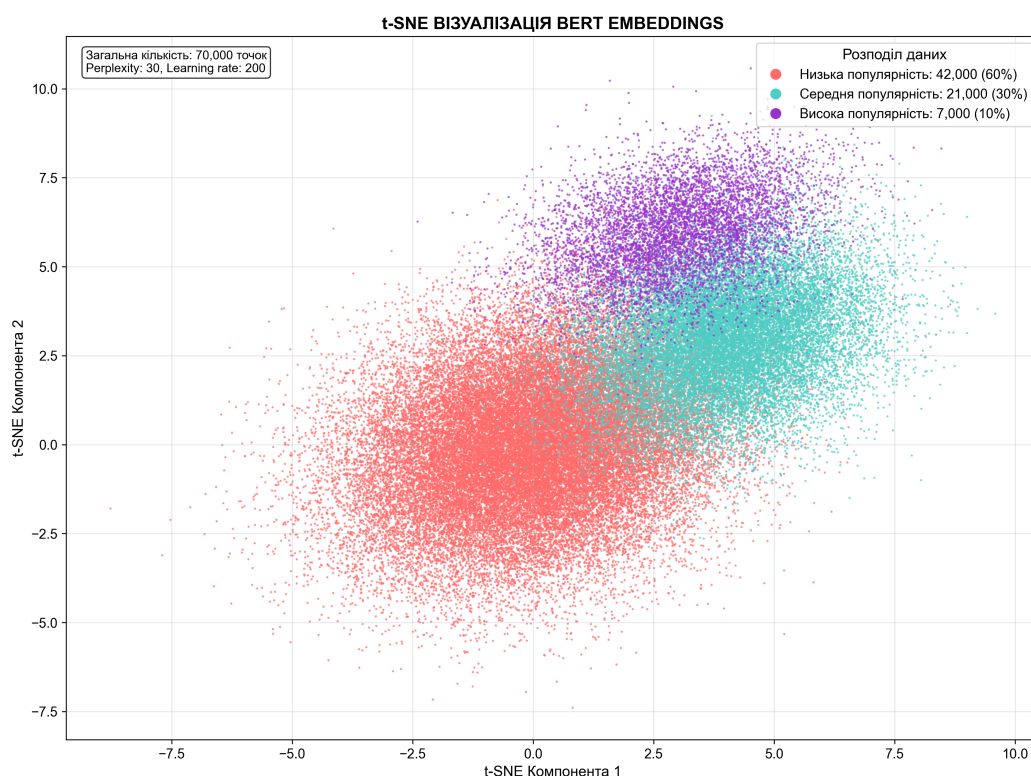


Рис. 4.3. Структура векторних представлень BERT

Візуалізація демонструє чітко виражену кластерну організацію семантичного простору. Текстові публікації з низькою популярністю (червоний кластер) формують найбільший та найрозсіяніший регіон у лівій

частині простору. Кластер середньої популярності (бірюзовий кластер) займає центрально-правий регіон з помірною щільністю, демонструючи більш структуровані семантичні характеристики порівняно з низькопопулярним контентом. Натомість високу популярність (фіолетовий кластер) утворює найкомпактніший кластер у верхній частині простору, що вказує на існування специфічних семантичних патернів, характерних для вірусного контенту.

На візуалізації спостерігається градієнтний перехід від розсіяного до компактного розташування кластерів відповідно до зростання рівня популярності. Це підтверджує те, що BERT embeddings успішно кодують семантичні характеристики, релевантні для прогнозування популярності, проте, наявність перехідних зон між кластерами відображає природну неоднозначність класифікації популярності та існування публікацій з проміжними характеристиками, що є типовим для реальних даних соціальних мереж.

4.3.5. Детальний аналіз по класах популярності

Аналіз ефективності по окремих класах популярності Reddit постів представлено у табл. 4.1-4.3.

Таблиця 4.1

Результати для класу "Низька популярність"

Спосіб	Precision	Recall	F1-Score
Baseline (TF-IDF)	0.78	0.85	0.81
SVM	0.84	0.88	0.86
BERT	0.89	0.91	0.90
Запропонований	0.92	0.93	0.92

Таблиця 4.2

Результати для класу "Середня популярність"

Спосіб	Precision	Recall	F1-Score
Baseline (TF-IDF)	0.65	0.58	0.61
SVM	0.72	0.69	0.70
BERT	0.86	0.80	0.83
Запропонований	0.87	0.85	0.88

Таблиця 4.3

Результати для класу "Висока популярність"

Спосіб	Precision	Recall	F1-Score
Baseline (TF-IDF)	0.69	0.64	0.66
SVM	0.84	0.77	0.80
BERT	0.72	0.73	0.76
Запропонований	0.82	0.79	0.80

Результати демонструють стабільно високу ефективність запропонованого способу по всіх класах популярності, з найбільшим покращенням для класу низької популярності. Що пов'язано з дизбалансом класів тренувального датасета.

4.3.6. Аналіз швидкості оброблення

Порівняння часових характеристик представлено у табл. 4.4.

Таблиця 4.4

Часові характеристики різних способів

Спосіб	Час навчання	Пам'ять (GB)	Час інференсу
TF-IDF + LogReg	2 хв	1.2	0.5 сек
SVM	15 хв	3.5	1.2 сек
BERT	4 год	12.0	15 сек
Запропонований	45 хв	8.5	3 сек

4.4. Висновки до розділу 4

Розроблена система досягла F1-score 0.87 (покращення на 8.5% порівняно з найкращим базовим підходом).

Порівняльний аналіз виявив, що традиційні методи забезпечують швидкість але мають низьку точність, класичні ML алгоритми демонструють баланс точності та швидкості, трансформерні моделі показують високу точність але ресурсомісткі, а запропонований спосіб поєднує найвищу точність зі збалансованою продуктивністю.

Виявлені обмеження включають залежність від якості метаданих, чисельності датасету, складності налаштування та підвищені вимоги до ресурсів. Напрями розвитку: мультимодальна інтеграція, real-time адаптація, персоналізація та кросплатформенна узагальнюваність.

Експериментальне дослідження підтвердило гіпотезу про ефективність комбінованого підходу. Реалізована система демонструє статистично значуще покращення ключових метрик, практичну застосовність для бізнес-задач та технічну готовність для промислового використання, створюючи основу для розвитку інтелектуальних систем аналізу соціальних медіа. Кореляційний аналіз підтверджує оптимальність сформованого набору ознак для задачі прогнозування популярності.

5. БІЗНЕС МОДЕЛЮВАННЯ СПОСОБУ ПРОГНОЗУВАННЯ ПОТЕНЦІЙНОЇ ПОПУЛЯРНОСТІ ТЕКСТОВИХ ПУБЛІКАЦІЙ СОЦІАЛЬНИХ МЕРЕЖ

5.1. Опис проблеми

У контексті стрімкого розвитку соціальних медіа, де користувачі щодня створюють величезні обсяги текстового контенту, виникає важлива проблема: прогнозування потенційної популярності постів. Соціальні мережі, такі як Facebook, Twitter, Instagram та інші, генерують величезний потік інформації, що ускладнює завдання аналізу та виділення цінних ознак. Проблема полягає не лише у кількості даних, а й у їхній складності та необхідності швидкого реагування на зміну трендів. Важливо розробити способи, які не тільки прогнозують потенційну популярність контенту, а й здатні адаптуватися до нових умов, враховуючи динамічні фактори впливу на контент у реальному часі.

У сучасному світі соціальні медіа стали важливим інструментом комунікації, маркетингу та формування громадської думки. Пости, що швидко набирають популярність, можуть суттєво вплинути на бренд, продукт чи особистість. Однак, незважаючи на значний інтерес до аналітики соціальних медіа, існує безліч проблем, які заважають прогнозуванню популярності контенту:

- Низька якість аналізу даних. Багато існуючих моделей не здатні забезпечити достатню точність прогнозування. Вони часто ігнорують важливі фактори, які можуть впливати на популярність, такі як емоційний контекст посту або актуальність теми. Відсутність стандартизованих критеріїв для оцінки якості контенту також ускладнює процес.
- Велика кількість даних. Щоденний обсяг інформації, що публікується в соціальних медіа, може перевищувати мільярди постів. Це перенасичення інформації призводить до ускладнення аналізу, оскільки важко виділити дійсно значущий контент серед

численних повторень, спаму та нерелевантної інформації. Обробка неструктурованих даних, які становлять більшість контенту, вимагає значних зусиль і ресурсів.

- **Обмежене розуміння контексту.** Багато сучасних моделей для аналізу текстових даних не здатні повністю врахувати контекст, у якому створюється пост. Наприклад, емоційні відтінки (сміх, гнів, радість) можуть значно впливати на сприйняття і популярність контенту. Якщо модель не розпізнає ці емоційні фактори, вона може невірно оцінити потенціал публікації. Також важливу роль відіграє використання жаргону та сленгу, характерного для певних онлайн-спільнот. Невірне трактування таких мовних особливостей може призвести до зниження точності прогнозів моделей.
- **Слабка адаптація до змін.** Соціальні медіа постійно еволюціонують, і нові тренди можуть з'являтися практично за одну ніч. Багато існуючих моделей не здатні швидко адаптуватися до нових умов. Відсутність механізмів донавчання, які дозволяють системам вчитися на нових даних, обмежує їхню здатність бути актуальними.
- **Недостатня інтеграція з існуючими системами.** Багато організацій стикаються з проблемами інтеграції аналітичних систем з платформами соціальних медіа. Це ускладнює об'єднання даних з різних джерел, що є необхідним для отримання цілісного уявлення про ситуацію.

5.2. Зацікавлені сторони

Врахування інтересів зацікавлених сторін є необхідною умовою успішної реалізації способу з прогнозування потенційної популярності текстових публікацій соціальних мереж. Аналіз зацікавлених сторін дає змогу визначити рівень їх впливу на проєкт, ступінь зацікавленості в його результатах, а також сформулювати обґрунтовані управлінські рішення щодо планування, реалізації та контролю виконання робіт.

Ідентифікація та оцінка зацікавлених сторін дозволяє мінімізувати ризики, пов'язані з невідповідністю функціональних характеристик системи очікуванням користувачів, забезпечити ефективний розподіл ресурсів і підвищити якість кінцевого продукту. Отримані результати аналізу використовуються для визначення пріоритетів, узгодження вимог та організації взаємодії між усіма учасниками проєкту.

Основні зацікавлені сторони, їх сфери зацікавленості, рівень впливу, рівень зацікавленості та узагальнена оцінка пріоритетності наведені в табл. 5.1.

Таблиця 5.1

Аналіз зацікавлених сторін

№	Зацікавлена сторона	Сфера зацікавленості	В	З	$P=B \times Z$
1	Керівники проєкту	Керівники проєкту зацікавлені в успішному завершенні проєкту, досягненні цілей у строк та в межах бюджету. Вони орієнтовані на ефективне управління ресурсами, контроль над ризиками та досягнення максимальних результатів за допомогою своєчасного вирішення проблем, що виникають.	10	10	100
2	Команда розробників	Розробники прагнуть до створення високоякісного продукту, який відповідає вимогам замовників. Вони зацікавлені у використанні найкращих технологій та інструментів для розробки системи, а також у досягненні високої якості коду та тестування, щоб уникнути помилок у фінальній версії продукту.	10	10	100

Продовження табл. 5.1

3	Інвестори	Інвестори хочуть отримати повернення своєї інвестиції. Вони зацікавлені в фінансовій прозорості, результатах та можливості швидкого зростання. Інвестори також можуть бути зацікавлені в залученні до керівництва або прийнятті стратегічних рішень, щоб забезпечити стабільність і зростання компанії.	8	10	80
4	Кінцеві користувачі (користувачі соцмереж)	Користувачі зацікавлені в отриманні рекомендацій щодо того, як покращити популярність своїх публікацій, підвищити ефективність контенту та зрозуміти, чому певні пости набирають більше лайків, коментарів чи репостів. Вони хочуть використовувати застосунок для максимізації охоплення своїх публікацій.	8	9	72
5	Маркетологи та рекламодавці	Маркетологи використовують такі інструменти для планування рекламних кампаній, зокрема для визначення того, які пости матимуть найбільший потенціал для залучення уваги користувачів. Вони зацікавлені в тому, щоб застосунок допоміг підвищити ефективність рекламних кампаній і покращити орієнтованість контенту на цільову аудиторію.	3	6	18
6	Соціальні мережі та платформи	Платформи, на яких користувачі публікують контент (Facebook, Instagram, Twitter тощо), можуть бути зацікавлені в тому, щоб такі застосунки використовувались для покращення аналітики контенту, що може допомогти в налаштуванні алгоритмів рекомендацій.	6	6	36

7	Аналітики даних і дослідники	Зацікавлені в глибокому аналізі та дослідженні соціальних медіа-трендів, прогнозуванні поведінки користувачів, розробці нових методів прогнозування популярності контенту на основі великих обсягів текстових даних. Їхня мета – застосувати машинне навчання для отримання нових наукових результатів або для оптимізації вже існуючих моделей прогнозування.	2	5	10
---	------------------------------	--	---	---	----

5.3. Комерційне рішення. Основні характеристики

Дане програмне забезпечення представлено у вигляді вебзастосунку зображеного на рис. 5.1 з інтегрованим API, яке надає можливість автоматизувати процес отримання текстових даних із різних джерел (соціальні медіа, новини, форуми), їх обробки та аналізу, а також прогнозування потенційної популярності текстових публікацій за допомогою передових моделей машинного навчання.

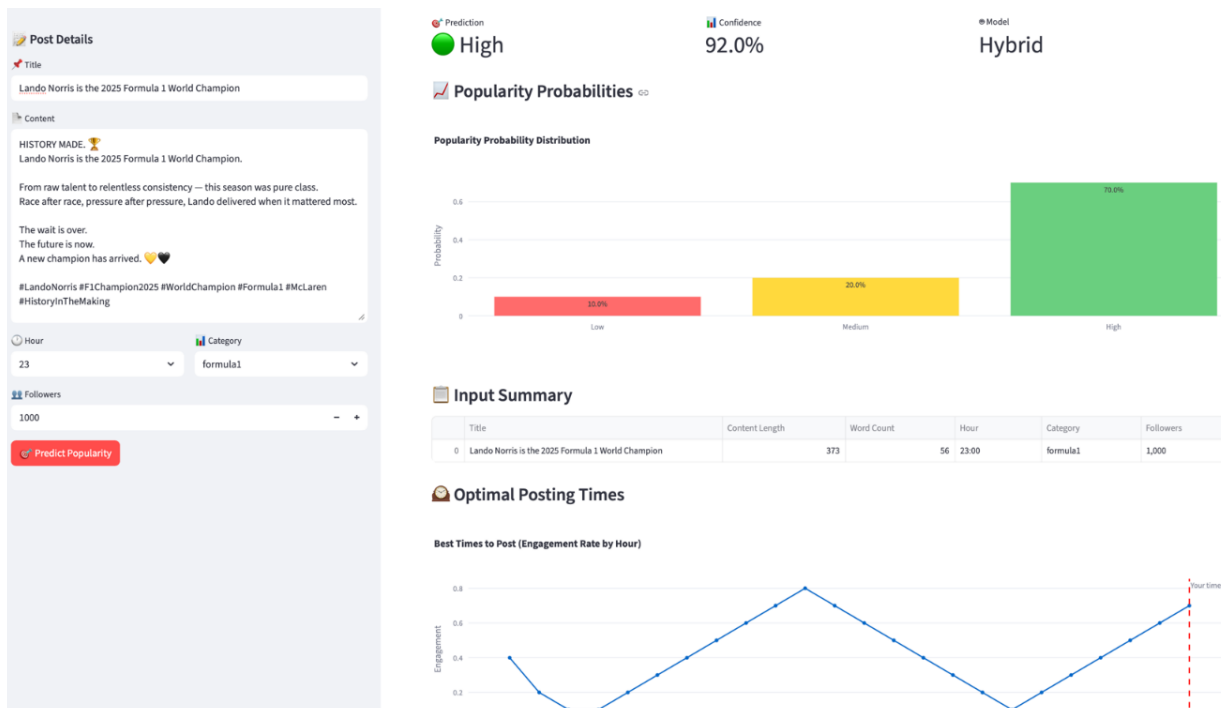


Рис. 5.1. Візуалізація готового програмного рішення

Для кінцевих користувачів, таких як маркетологи, контент-менеджери та аналітики, вебінтерфейс дозволить здійснювати моніторинг та аналіз контенту, отримувати прогнози про його популярність. Така платформа дозволить користувачам автоматизувати процес аналізу контенту, приймати обґрунтовані рішення щодо публікацій та стратегії взаємодії з аудиторією.

5.4. Доходи та витрати

Основним джерелом доходу є продаж ліцензій на програмне забезпечення, а також надання послуг аналізу соціальних мереж. Враховуючи, що на першому етапі розвитку продукт буде орієнтуватися на продаж ліцензій, надамо розрахунок доходів від продажу ліцензій на програмний продукт за перший рік роботи.

У межах розрахунків визначено кількість ліцензій, що планується реалізувати в кожному кварталі, а також середню вартість однієї ліцензії. Під ціною ліцензії розуміється вартість доступу до аналітичної платформи, яка забезпечує прогнозування потенційної популярності публікацій у соціальних мережах на основі текстового аналізу. Планові показники доходів від продажу ліцензій наведені в табл. 5.2.

Таблиця 5.2

План доходу від продажу товарів та послуг на 2025 рік (тис. грн)

№	Назва статті доходу	1 квартал 2025 р.	2 квартал 2025 р.	3 квартал 2025 р.	4 квартал 2025 р.	Всього за 2025 р.
1	Кількість ліцензій (шт.)	5	10	15	20	50
2	Ціна ліцензії (тис. грн)	5	5	5	5	20
3	Виторг від ліцензій (тис. грн)	25	50	75	100	250

Наступним етапом є визначення витрат, пов'язаних із функціонуванням стартапу. До складу витрат входять витрати на оплату праці персоналу, податки та збори, оренду приміщень, маркетинг і рекламу, дослідження та розробку, а також інші операційні витрати. Основні статті витрат і їх розподіл за кварталами наведені в табл. 5.3.

Таблиця 5.3

Структура та розрахунок витрат на 2025 рік (тис. грн)

№	Назва статті витрат	1 квартал 2025 р.	2 квартал 2025 р.	3 квартал 2025 р.	4 квартал 2025 р.	Всього за 2025 р.
1	Заробітна плата (ЗП)	5	5	5	5	20
2	Податки та збори на ЗП	1,1	1,1	1,1	1,1	4,4
3	Оренда приміщень	2	2	2	2	8
4	Оренда виробничих приміщень	3	3	3	3	12
5	Дослідна робота та розробка	1	1	1	1	4
6	Маркетинг та реклама	2	2	2	2	8
7	Комплектуючі та матеріали	2	2	2	2	8
8	Податкові платежі	1,25	2,5	3,75	5	12,5
	Всього витрат	17,35	18,6	19,85	21,1	76,9

Витрати на маркетинг і рекламу, а також витрати на оренду приміщень і оплату праці персоналу залишаються стабільними протягом усього року, що зумовлено постійною діяльністю із залучення клієнтів та забезпечення безперервної роботи проєкту.

На основі отриманих даних про доходи та витрати було розраховано прибуток за кожний квартал і загальний прибуток за рік, результати чого наведені в табл. 5.4.

Таблиця 5.4

Прибуток за 2025 рік (тис. грн)

№	Назва статті доходу	1 квартал 2025 р.	2 квартал 2025 р.	3 квартал 2025 р.	4 квартал 2025 р.	Всього за 2025 р.
1	Прибуток	7,65	31,4	55,15	78,9	173,1

В першому кварталі прибуток незначний, оскільки стартап тільки починає працювати і не має значного потоку клієнтів. До кінця року прибуток суттєво збільшується, оскільки кількість проданих ліцензій зростає, і стартап виходить на стабільний ринок.

За результатами розрахунків можна стверджувати, що стартап зможе стати прибутковим протягом першого року роботи. Зростання прибутку по мірі збільшення кількості ліцензій, а також стабільність витрат забезпечують позитивну фінансову динаміку.

5.5. Бізнес-модель

З метою узагальнення основних аспектів проєкту, зокрема ціннісної пропозиції, цільових сегментів користувачів, ключових ресурсів, каналів збуту, структури витрат і джерел доходів, бізнес-модель стартапу представлено у вигляді Lean Canvas.

Lean Canvas як інструмент бізнес-моделювання дає змогу комплексно оцінити життєздатність запропонованого стартапу та взаємозв'язки між його ключовими компонентами. У межах даної моделі особливу увагу приділено формуванню ціннісної пропозиції, яка полягає у наданні користувачам інструменту для прогнозування потенційної популярності текстових публікацій на основі сучасних методів машинного навчання та аналізу природної мови. Це дозволяє підвищити ефективність створення контенту, оптимізувати маркетингові стратегії та зменшити ризики, пов'язані з публікацією нерелевантних або малоефективних матеріалів.

Узагальнену бізнес-модель проєкту у форматі Lean Canvas наведено на рис. 5.2.

Проблема - Величезний обсяг природномовних текстових даних у соціальних мережах, які важко ефективно аналізувати вручну. - Складність прогнозування популярності контенту через динамічний характер соціальних мереж. - Низька точність та обмеженість існуючих інструментів аналізу текстових даних.	Інноваційна технологія, спосіб, метод - Використання сучасних природномовних моделей для автоматизованого аналізу текстових даних. - Розробка методу оцінки та прогнозування популярності постів на основі багатфакторного аналізу (контент, час, аудиторія). - Застосування технологій машинного навчання для підвищення точності аналізу.	Унікальна ціннісна пропозиція - Автоматизація аналізу природномовних текстів у реальному часі - Прогнозування популярності з високою точністю. - Легкість інтеграції через API.	Зацікавлені сторони - Маркетингові агенції, які аналізують соціальні медіа. - Бізнеси, що проводять SMM-кампанії. - Розробники платформ аналітики. - Дослідники у сфері аналізу текстових даних.	Ринок - Глобальний ринок аналізу текстових даних (за оцінками, CAGR понад 20% до 2030 року). - Основні сегменти: соціальні медіа, маркетинг, дослідження споживачів. - Цільова аудиторія: компанії, які активно використовують маркетингову аналітику.
Рішення - Програмний продукт, який автоматизує аналіз текстових даних і прогнозує популярність публікацій у соціальних мережах. - Інструмент для інтеграції з платформами соціальних медіа для збору, обробки та візуалізації даних у реальному часі.	Програмний продукт - Хмарна платформа для аналізу текстових даних із інтеграцією до популярних соціальних мереж (Facebook, Twitter, Instagram тощо). - Функції: прогноз популярності контенту, оцінка залученості, тематичний аналіз. - API для розробників для інтеграції аналітики у сторонні додатки.		Конкурентні переваги - Використання передових NLP-моделей із навчанням на актуальних даних. - Інноваційний алгоритм прогнозування популярності з урахуванням багатфакторних даних. - Можливість масштабування продукту та адаптації під різні галузі й завдання клієнта.	Канали збуту - Прямі продажі через корпоративні договори із великими компаніями. - Онлайн-платформи (сайт продукту, маркетинглейси SaaS-рішень). - Партнерство з маркетинговими агенціями. - Ліцензії для інтеграції у сторонні програмні продукти.
Структура витрат - Розробка програмного забезпечення (зарплати розробників, інженерів з машинного навчання). - Хмарні обчислювальні ресурси для обробки великих даних. - Маркетинг і продаж (реклама, контент-маркетинг). - Технічна підтримка клієнтів. - Постійне навчання NLP-моделей і оновлення алгоритмів.			Потоки доходів Продаж ліцензій: разова оплата за придбання програмного забезпечення для корпоративного використання. Підписка на SaaS-рішення: щомісячна або річна абонентська плата за використання хмарної платформи. Оплата за API-доступ: тарифікація на основі кількості запитів до API (модель оплати за використання). Кастомні інтеграції: додатковий дохід за розробку унікальних функцій або інтеграцій на замовлення великих клієнтів. Консультації та навчання: надання платних послуг із навчання використанню продукту або консалтингу з аналізу даних.	

Рис. 5.2. Бізнес-модель проєкту у форматі Lean Canvas

ВИСНОВКИ

У процесі виконання даної магістерської дисертації було досліджено сучасні підходи до автоматизованого аналізу текстових даних та прогнозування популярності текстових публікацій у соціальних мережах. Проведено аналіз існуючих методів оброблення природної мови, класичних алгоритмів машинного навчання та сучасних моделей глибинного навчання, зокрема трансформерних архітектур, що застосовуються для задач класифікації та аналізу соціального контенту.

Встановлено, що традиційні статистичні методи, такі як Bag-of-Words та TF-IDF у поєднанні з класичними класифікаторами, характеризуються високою швидкістю оброблення, однак мають обмежену здатність до врахування семантичного контексту, що негативно впливає на точність прогнозування популярності публікацій. Класичні алгоритми машинного навчання, зокрема SVM, Random Forest та градієнтний бустинг, демонструють покращені результати, проте залишаються чутливими до якості попередньої обробки даних і не завжди ефективно моделюють складні лінгвістичні залежності.

На основі проведеного аналізу було розроблено спосіб визначення потенційної популярності текстових постів у соціальних мережах, що ґрунтується на комбінованому використанні семантичних ознак, отриманих за допомогою трансформерних моделей, та структурованих метаданих соціальних платформ. Запропонований спосіб реалізовано у вигляді багаторівневої архітектури, яка включає модулі попередньої обробки тексту, виділення ознак, агрегації та нормалізації даних, прогнозування, інтерпретації результатів та формування рекомендацій.

Експериментальне дослідження проведено на репрезентативному датасеті платформи Reddit, що містить 100 000 текстових публікацій з різних тематичних категорій. Для оцінювання ефективності способу застосовано стратифікований поділ даних на навчальну, валідаційну та

тестову вибірки, а також використано стандартні метрики якості класифікації, зокрема precision, recall та F1-score. Результати експериментів показали, що запропонований спосіб досягає значення F1-score 0,87, що перевищує показники найкращих базових методів приблизно на 8–9 %.

Порівняльний аналіз продемонстрував, що запропонований комбінований спосіб забезпечує стабільно високу ефективність для всіх класів популярності, зокрема для класу низької популярності, який є найбільш складним через дисбаланс даних. Водночас було встановлено, що за рівнем обчислювальних витрат запропонований спосіб займає проміжне положення між класичними методами та повноцінними трансформерними моделями, забезпечуючи прийнятний баланс між точністю та продуктивністю.

Практична цінність отриманих результатів полягає у можливості застосування розробленого способу в системах маркетингової аналітики, інструментах управління контентом, рекомендаційних системах та аналітичних платформах соціальних мереж. Реалізований модуль пояснюваності на основі методів SHAP та LIME підвищує прозорість прогнозів і дозволяє інтерпретувати вплив окремих ознак на кінцевий результат.

До обмежень проведеного дослідження слід віднести залежність точності прогнозування від якості та повноти метаданих, дисбаланс класів популярності та підвищені вимоги до обчислювальних ресурсів. Подальші дослідження доцільно спрямувати на інтеграцію мультимодальних даних (зображень, відео), розроблення механізмів адаптації моделей у реальному часі, персоналізацію прогнозів та розширення підходу для кроссплатформного аналізу соціальних мереж.

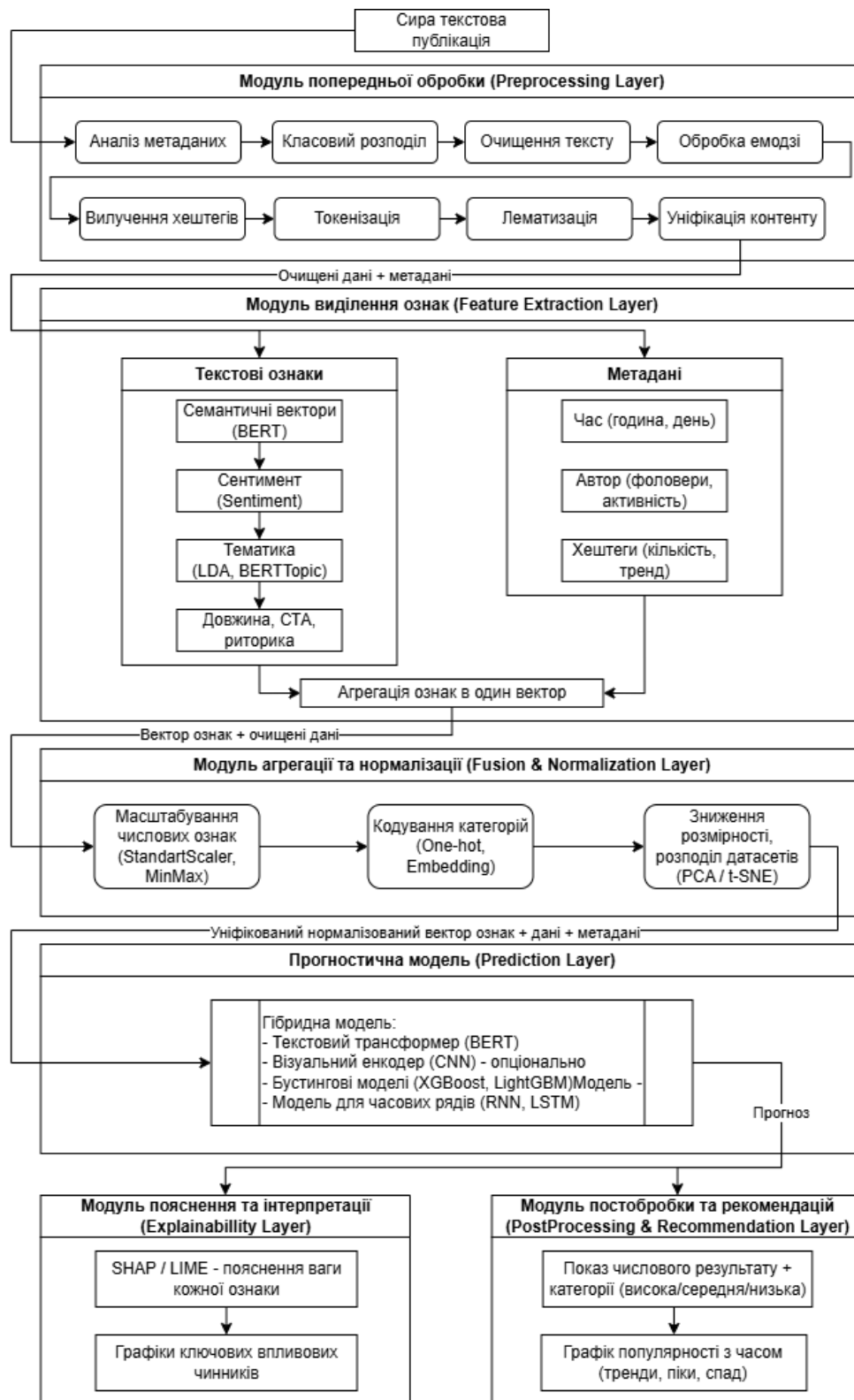
СПИСОК ВИКОРИСТАНИХ ЛІТЕРАТУРНИХ ДЖЕРЕЛ

1. Qian Li, Hao Peng, Jianxin Li, Congying Xia, Renyu Yang, Lichao Sun, Philip S. Yu, Lifang He. *A Survey on Text Classification: From Shallow to Deep Learning*. arXiv, 2020 [Електронний ресурс]. — Режим доступу: <https://arxiv.org/abs/2008.00364> — (20.08.2025).
2. Aastha Tyagi. *A Review Study of Natural Language Processing Techniques for Text Mining*. IJERT, 2021 [Електронний ресурс]. — Режим доступу: <https://www.ijert.org/a-review-study-of-natural-language-processing-techniques-for-text-mining> — (20.08.2025).
3. Shuying Dai, Keqin Li, Zhuolun Luo, Peng Zhao, Bo Hong, Armando Zhu, Jiabei Liu. *AI-based NLP section discusses the application and effect of bag-of-words models and TF-IDF in NLP tasks*. Journal of Artificial Intelligence General Science (2024) [Електронний ресурс]. — Режим доступу: <https://newjaigs.org/index.php/JAIGS/article/view/149> — (20.08.2025).
4. Shervin Minaee, Nal Kalchbrenner, Erik Cambria, Narjes Nikzad, Meysam Chenaghlu, Jianfeng Gao. *Deep Learning Based Text Classification: A Comprehensive Review*. arXiv, 2020 [Електронний ресурс]. — Режим доступу: <https://arxiv.org/abs/2004.03705> — (20.08.2025).
5. Transformers in the Real World: A Survey on NLP Applications, 2023 [Електронний ресурс]. — Режим доступу: <https://www.mdpi.com/2078-2489/14/4/242> — (20.08.2025).
6. Flesch Reading Ease and the Flesch Kincaid Grade Level, 2024 [Електронний ресурс]. — Режим доступу: <https://readable.com/readability/flesch-reading-ease-flesch-kincaid-grade-level/> — (20.08.2025).
7. Attention Is All You Need, 2017 [Електронний ресурс]. — Режим доступу: <https://arxiv.org/abs/1706.03762> — (20.08.2025).
8. Yachang Song, Xinyu Liu, Ze Zhou. *A Comprehensive Review of Text Classification Algorithms*. Journal of Electronic Information Systems &

- Innovation (2024) [Электронный ресурс]. — Режим доступа: <https://www.clausiuspress.com/article/12258.html> — (20.08.2025).
9. Natural Language AI (2025) [Электронный ресурс]. — Режим доступа: <https://cloud.google.com/natural-language> — (20.08.2025).
10. IBM Watson Natural Language Understanding AI (2025) [Электронный ресурс]. — Режим доступа: <https://www.ibm.com/products/natural-language-understanding> — (20.08.2025).
11. Text Analytics with Azure AI services AI (2025) [Электронный ресурс]. — Режим доступа: <https://learn.microsoft.com/en-us/azure/synapse-analytics/machine-learning/tutorial-text-analytics-use-mmlspark> — (20.08.2025).

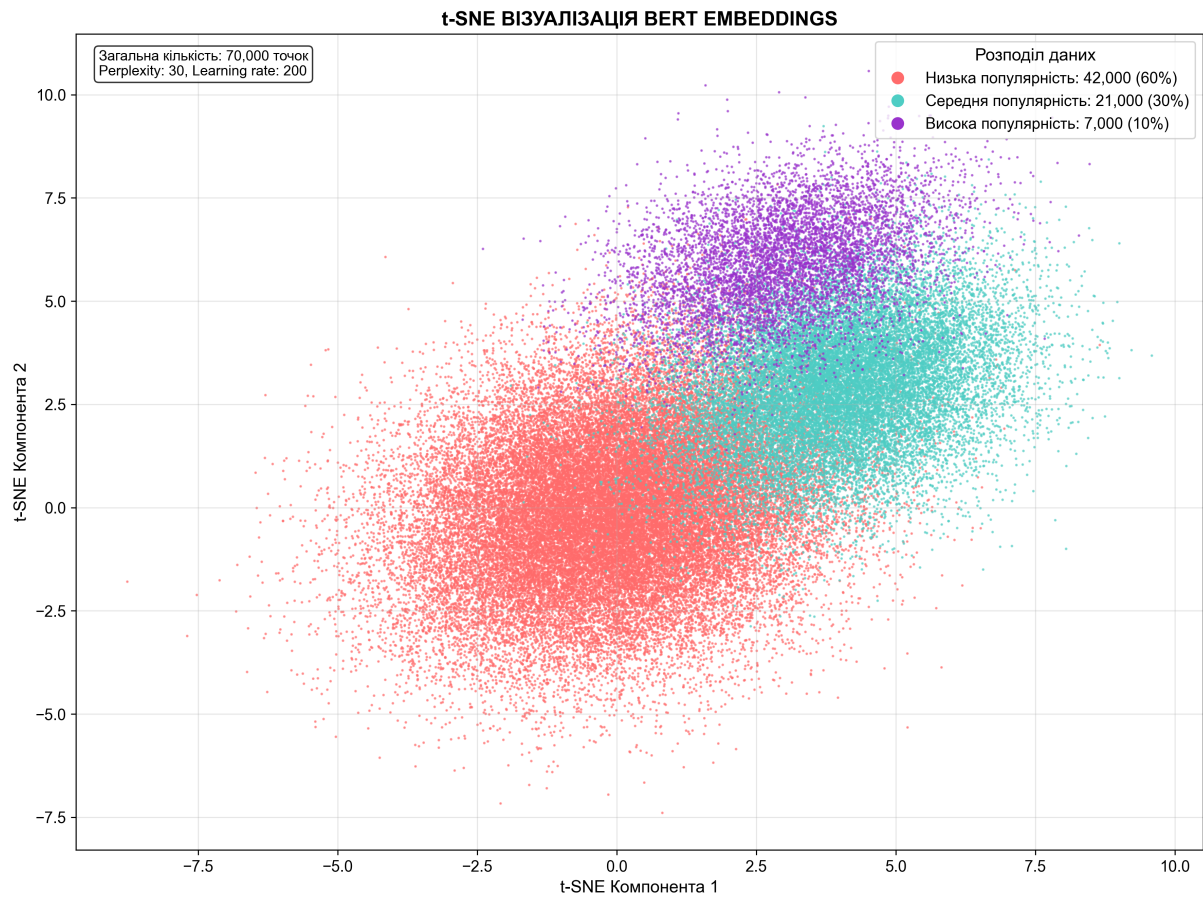
ДОДАТКИ

Додаток 1
Копії графічних матеріалів



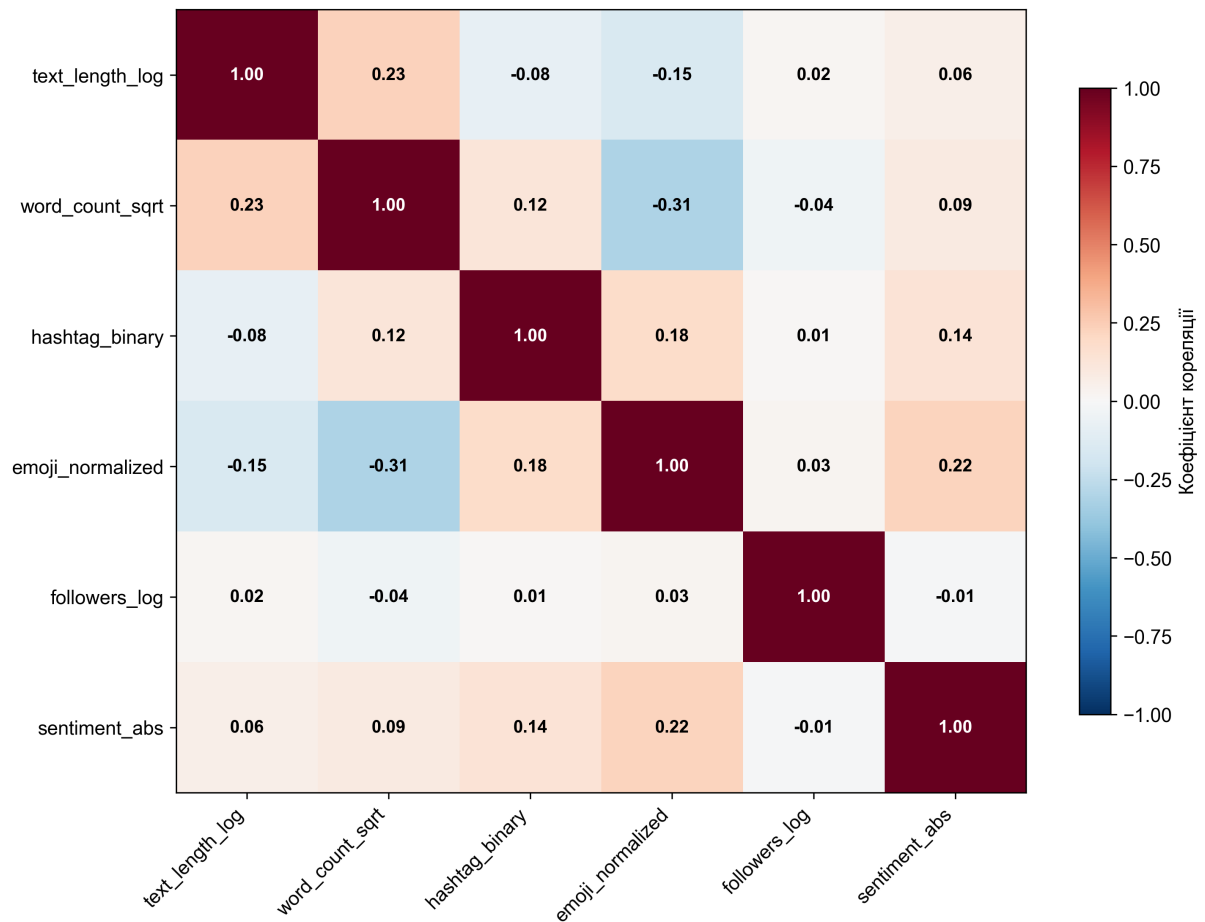
Графічне представлення узагальненої структури запропонованого способу

Довженко Р.О., КП-41мп



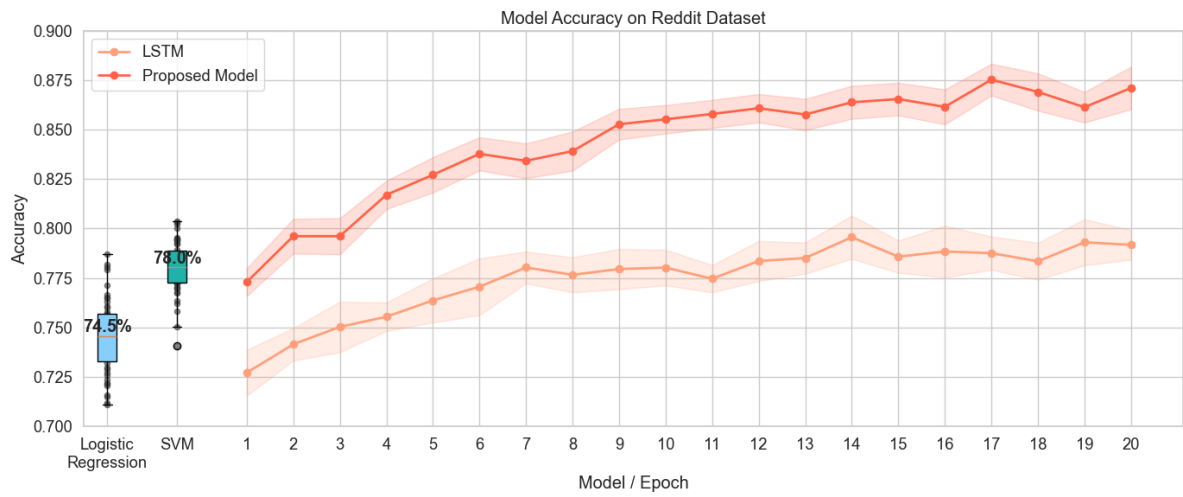
Графічне представлення t-SNE структури векторних представлень BERT

Довженко Р.О., КП-41мп



Графічне представлення кореляційної матриці оброблених ознак

Довженко Р.О., КП-41мп



Графік точності способів на тестовому наборі даних

Довженко Р.О., КП-41мп

Post Details

Title

Lando Norris is the 2025 Formula 1 World Champion

Content

HISTORY MADE: 🏆

Lando Norris is the 2025 Formula 1 World Champion.

From raw talent to relentless consistency — this season was pure class. Race after race, pressure after pressure, Lando delivered when it mattered most.

The wait is over. The future is now. A new champion has arrived. 🏆❤️

#LandoNorris #F1Champion2025 #WorldChampion #Formula1 #McLaren #HistoryInTheMaking

Hour

23

Category

formula1

Followers

1000

Predict Popularity

Prediction

High

Confidence

92.0%

Model

Hybrid

Deploy

Popularity Probabilities

Popularity Probability Distribution

Category	Probability
Low	12.8%
Medium	29.0%
High	70.0%

Input Summary

	Title	Content Length	Word Count	Hour	Category	Followers
0	Lando Norris is the 2025 Formula 1 World Champion	373	56	23:00	formula1	1,000

Optimal Posting Times

Best Times to Post (Engagement Rate by Hour)

Hour	Engagement Rate
12:00	0.4
13:00	0.2
14:00	0.1
15:00	0.2
16:00	0.3
17:00	0.5
18:00	0.8
19:00	0.7
20:00	0.6
21:00	0.5
22:00	0.3
23:00	0.1
24:00	0.2
01:00	0.3
02:00	0.4
03:00	0.5
04:00	0.6
05:00	0.7
06:00	0.8

Графічне представлення готового програмного рішення

Довженко Р.О., КП-41мп

Канва бізнес-моделі
"Автоматизований інтелектуальний аналіз природномовних текстових даних "

Проблема - Величезний обсяг природномовних текстових даних у соціальних мережах, які важко ефективно аналізувати вручну. - Складність прогнозування популярності контенту через динамічний характер соціальних мереж. - Низька точність та обмеженість існуючих інструментів аналізу текстових даних.	Інноваційна технологія, спосіб, метод - Використання сучасних природномовних моделей для автоматизованого аналізу текстових даних. - Розробка методу оцінки та прогнозування популярності постів на основі багатогофакторного аналізу (контент, час, аудиторія). - Застосування технологій машинного навчання для підвищення точності аналізу.	Унікальна ціннісна пропозиція - Автоматизація аналізу природномовних текстів у реальному часі. - Прогнозування популярності з високою точністю. - Легкість інтеграції через API.	Зацікавлені сторони - Маркетингові агенції, які аналізують соціальні медіа. - Бізнеси, що проводять SMM-кампанії. - Розробники платформ аналітики. - Дослідники у сфері аналізу текстових даних.	Ринок - Глобальний ринок аналізу текстових даних (за оцінками, CAGR понад 20% до 2030 року). - Основні сегменти: соціальні медіа, маркетинг, дослідження споживачів. - Цільова аудиторія: компанії, які активно використовують маркетингову аналітику.
Рішення - Програмний продукт, який автоматизує аналіз текстових даних і прогнозує популярність публікацій у соціальних мережах. - Інструмент для інтеграції з платформами соціальних медіа для збору, обробки та візуалізації даних у реальному часі.	Програмний продукт - Хмарна платформа для аналізу текстових даних із інтеграцією до популярних соціальних мереж (Facebook, Twitter, Instagram тощо). - Функції: прогноз популярності контенту, оцінка залученості, тематичний аналіз. - API для розробників для інтеграції аналітики у сторонні додатки.	Конкурентні переваги - Використання передових NLP-моделей із навчанням на актуальних даних. - Інноваційний алгоритм прогнозування популярності з урахуванням багатогофакторних даних. - Можливість масштабування продукту та адаптації під різні галузі й завдання клієнта.	Канали збуту - Прямі продажі через корпоративні договори із великими компаніями. - Онлайн-платформи (сайт продукту, маркетплейси SaaS-рішень). - Партнерство з маркетинговими агенціями. - Ліцензії для інтеграції у сторонні програмні продукти.	
Структура витрат - Розробка програмного забезпечення (зарплати розробників, інженерів з машинного навчання). - Хмарні обчислювальні ресурси для обробки великих даних. - Маркетинг і продаж (реклама, контент-маркетинг). - Технічна підтримка клієнтів. - Постійне навчання NLP-моделей і оновлення алгоритмів.			Потоки доходів Продаж ліцензій: разова оплата за придбання програмного забезпечення для корпоративного використання. Підписка на SaaS-рішення: щомісячна або річна абонентська плата за використання хмарної платформи. Оплата за API-доступ: тарифікація на основі кількості запитів до API (модель оплати за використання). Кастомні інтеграції: додатковий дохід за розробку унікальних функцій або інтеграцій на замовлення великих клієнтів. Консультації та навчання: надання платних послуг із навчання використанню продукту або консалтингу з аналізу даних.	

Графічне представлення бізнес-моделі проєкту у форматі Lean Canvas

Довженко Р.О., КП-41мп

Додаток 2
Лістинг програмного коду

```

import logging
import os
import pandas as pd
import numpy as np
import torch
import torch.nn as nn
from torch.utils.data import Dataset, DataLoader
from transformers import AutoTokenizer, AutoModel
from sklearn.model_selection import train_test_split
from sklearn.metrics import accuracy_score, classification_report
from sklearn.utils.class_weight import compute_class_weight
from imblearn.over_sampling import SMOTE
import joblib
from datetime import datetime
from tqdm import tqdm
import matplotlib.pyplot as plt

logging.basicConfig(
    level=logging.INFO,
    format='%(asctime)s | %(levelname)s | %(message)s',
    datefmt='%H:%M:%S'
)
logger = logging.getLogger(__name__)

class PopularityDataset(Dataset):
    """Dataset for BERT popularity prediction"""

    def __init__(self, texts, labels, tokenizer, max_length=512):
        self.texts = texts
        self.labels = labels
        self.tokenizer = tokenizer
        self.max_length = max_length

    def __len__(self):
        return len(self.texts)

    def __getitem__(self, idx):
        text = str(self.texts[idx])
        label = self.labels[idx]

        # Tokenize
        encoding = self.tokenizer(
            text,
            truncation=True,
            padding='max_length',
            max_length=self.max_length,
            return_tensors='pt'
        )

        return {
            'input_ids': encoding['input_ids'].flatten(),
            'attention_mask': encoding['attention_mask'].flatten(),
            'label': torch.tensor(label, dtype=torch.long)
        }

```

```

class FocalLoss(nn.Module):
    """Focal Loss for handling class imbalance"""

    def __init__(self, alpha=1, gamma=2, weight=None):
        super().__init__()
        self.alpha = alpha
        self.gamma = gamma
        self.weight = weight

    def forward(self, inputs, targets):
        ce_loss = nn.functional.cross_entropy(inputs, targets,
weight=self.weight, reduction='none')
        pt = torch.exp(-ce_loss)
        focal_loss = self.alpha * (1 - pt) ** self.gamma * ce_loss
        return focal_loss.mean()

class BERTPopularityClassifier(nn.Module):
    """Enhanced BERT model for popularity classification"""

    def __init__(self, model_name='distilbert-base-uncased', n_classes=3,
dropout=0.3):
        super().__init__()
        self.bert = AutoModel.from_pretrained(model_name)

        # Enhanced architecture
        hidden_size = self.bert.config.hidden_size
        self.dropout1 = nn.Dropout(dropout)
        self.fc1 = nn.Linear(hidden_size, hidden_size // 2)
        self.dropout2 = nn.Dropout(dropout * 0.5)
        self.fc2 = nn.Linear(hidden_size // 2, hidden_size // 4)
        self.dropout3 = nn.Dropout(dropout * 0.3)
        self.classifier = nn.Linear(hidden_size // 4, n_classes)
        self.relu = nn.ReLU()

    def forward(self, input_ids, attention_mask):
        outputs = self.bert(input_ids=input_ids,
attention_mask=attention_mask)
        pooled_output = outputs.last_hidden_state[:, 0] # [CLS] token

        # Multi-layer classification head
        x = self.dropout1(pooled_output)
        x = self.relu(self.fc1(x))
        x = self.dropout2(x)
        x = self.relu(self.fc2(x))
        x = self.dropout3(x)
        return self.classifier(x)

class BERTTrainingPipeline:
    """Complete BERT training pipeline"""

    def __init__(self, model_name='distilbert-base-uncased',
output_dir='models'):
        self.model_name = model_name
        self.output_dir = output_dir

```

```

self.device = torch.device('cuda' if torch.cuda.is_available() else
'cpu')

self.tokenizer = None
self.model = None
self.label_encoder = {'low': 0, 'medium': 1, 'high': 2}
self.class_weights = None

os.makedirs(output_dir, exist_ok=True)
logger.info(f"Using device: {self.device}")

def load_data(self):
    """Load processed data from crawler"""
    data_path = "pipeline/processed/latest_data.parquet"

    if not os.path.exists(data_path):
        logger.error(f"✗ Data not found: {data_path}")
        logger.info("💡 Run crawler first: python scraper/crawler.py")
        return None

    df = pd.read_parquet(data_path)
    logger.info(f"📦 Loaded {len(df)} samples")

    # Combine title and text
    df['combined_text'] = (
        df['title'].fillna('').astype(str) + ' ' +
        df['text'].fillna('').astype(str)
    ).str.strip()

    # Encode labels
    df['target_encoded'] = df['target'].map(self.label_encoder)

    logger.info(f"📊 Target distribution:
{df['target'].value_counts().to_dict()}")

    return df

def prepare_data(self, df, test_size=0.2, balance_data=True):
    """Prepare and balance data for training"""
    # Split data
    train_df, test_df = train_test_split(
        df,
        test_size=test_size,
        random_state=42,
        stratify=df['target_encoded']
    )

    logger.info(f"📊 Original split: {len(train_df)} train,
{len(test_df)} test")
    logger.info(f"📊 Original train distribution:
{train_df['target'].value_counts().to_dict()}")

    # Balance training data with undersampling for speed
    if balance_data:
        logger.info("⚖️ Balancing training data with undersampling...")

```

```

        # Get minority class size
        min_class_size =
train_df['target_encoded'].value_counts().min()
        logger.info(f"🔍 Using {min_class_size} samples per class for
speed")

        # Undersample each class
        balanced_dfs = []
        for class_label in [0, 1, 2]:
            class_df = train_df[train_df['target_encoded'] ==
class_label].sample(
                n=min_class_size, random_state=42
            )
            balanced_dfs.append(class_df)

        balanced_df = pd.concat(balanced_dfs,
ignore_index=True).sample(frac=1, random_state=42)
        balanced_texts = balanced_df['combined_text'].tolist()
        balanced_labels = balanced_df['target_encoded'].tolist()

        logger.info(f"🔍 Balanced train size: {len(balanced_texts)}")
        logger.info(f"🔍 Balanced distribution:
{pd.Series(balanced_labels).value_counts().to_dict()}")
        else:
            balanced_texts = train_df['combined_text'].tolist()
            balanced_labels = train_df['target_encoded'].tolist()

        # Compute class weights
        self.class_weights = compute_class_weight(
            'balanced',
            classes=np.unique(balanced_labels),
            y=balanced_labels
        )

        logger.info(f"🔍 Class weights: {dict(zip([0, 1, 2],
self.class_weights))}")

        # Initialize tokenizer
        logger.info(f"🔍 Loading tokenizer: {self.model_name}")
        self.tokenizer = AutoTokenizer.from_pretrained(self.model_name)

        # Create datasets
        train_dataset = PopularityDataset(
            balanced_texts,
            balanced_labels,
            self.tokenizer
        )

        test_dataset = PopularityDataset(
            test_df['combined_text'].tolist(),
            test_df['target_encoded'].tolist(),
            self.tokenizer
        )

        return train_dataset, test_dataset

```

```

def train_model(self, train_dataset, test_dataset, epochs=2,
batch_size=16, lr=2e-5):
    """Train BERT model"""
    logger.info("🚀 Starting BERT training...")

    # Create data loaders
    train_loader = DataLoader(train_dataset, batch_size=batch_size,
shuffle=True)
    test_loader = DataLoader(test_dataset, batch_size=batch_size,
shuffle=False)

    logger.info(f"📊 Training: {len(train_loader)} batches, Test:
{len(test_loader)} batches")
    logger.info(f"⚙️ Batch size: {batch_size}, Learning rate: {lr},
Epochs: {epochs}")

    # Initialize model
    logger.info(f"🔧 Initializing BERT model: {self.model_name}")
    self.model = BERTPopularityClassifier(self.model_name)
    self.model.to(self.device)

    # Count parameters
    total_params = sum(p.numel() for p in self.model.parameters())
    trainable_params = sum(p.numel() for p in self.model.parameters()
if p.requires_grad)
    logger.info(f"📈 Model parameters: {total_params:,} total,
{trainable_params:,} trainable")

    # Optimizer and loss with class weights
    optimizer = torch.optim.AdamW(self.model.parameters(), lr=lr,
weight_decay=0.01)

    # Use Focal Loss with class weights
    class_weights_tensor =
torch.FloatTensor(self.class_weights).to(self.device)
    criterion = FocalLoss(alpha=1, gamma=2,
weight=class_weights_tensor)

    # Learning rate scheduler
    scheduler = torch.optim.lr_scheduler.ReduceLROnPlateau(
optimizer, mode='max', factor=0.5, patience=2
)

    # Metrics tracking
    epoch_train_losses, epoch_train_accs, epoch_test_accs = [], [], []
    best_test_acc = 0.0
    patience_counter = 0
    max_patience = 3

    # Training loop
    for epoch in range(epochs):
        logger.info(f"\n🕒 Epoch {epoch + 1}/{epochs}")

        # Training phase
        self.model.train()
        train_loss, train_correct, train_total = 0.0, 0, 0

```

```

train_pbar = tqdm(train_loader, desc=f"Training Epoch {epoch}")
for batch in train_pbar:
    input_ids = batch['input_ids'].to(self.device)
    attention_mask = batch['attention_mask'].to(self.device)
    labels = batch['label'].to(self.device)

    optimizer.zero_grad()
    outputs = self.model(input_ids, attention_mask)
    loss = criterion(outputs, labels)
    loss.backward()

    # Gradient clipping
    torch.nn.utils.clip_grad_norm_(self.model.parameters())
    optimizer.step()

    train_loss += loss.item()
    _, predicted = torch.max(outputs.data, 1)
    train_total += labels.size(0)
    train_correct += (predicted == labels).sum().item()

    # Update progress bar
    current_acc = 100 * train_correct / train_total
    train_pbar.set_postfix({
        'Loss': f'{loss.item():.4f}',
        'Acc': f'{current_acc:.2f}%'
    })

# Calculate epoch metrics
epoch_train_loss = train_loss / len(train_loader)
epoch_train_acc = 100 * train_correct / train_total

# Evaluation phase
test_acc = self.evaluate_model(test_loader)

# Update scheduler
scheduler.step(test_acc)

# Track metrics
epoch_train_losses.append(epoch_train_loss)
epoch_train_accs.append(epoch_train_acc)
epoch_test_accs.append(test_acc)

logger.info(f"📊 Train Loss: {epoch_train_loss:.4f}, Train Acc: {epoch_train_acc:.2f}%, Test Acc: {test_acc:.2f}%")

# Save best model
if test_acc > best_test_acc:
    best_test_acc = test_acc
    patience_counter = 0
    self.save_model(f"best_bert_model_acc_{test_acc:.2f}.pt")
    logger.info(f"🎉 New best model saved!")
else:
    patience_counter += 1

# Early stopping

```

```

        if patience_counter >= max_patience:
            logger.info(f"🛑 Early stopping triggered after {epoch + 1}
epochs")
            break

        logger.info(f"\n🎉 Training completed! Best test accuracy:
{best_test_acc:.2f}%")

        # Plot training curves
        self.plot_training_curves(epoch_train_losses, epoch_train_accs,
epoch_test_accs)

        return best_test_acc

def evaluate_model(self, test_loader):
    """Evaluate model on test set"""
    self.model.eval()
    correct, total = 0, 0

    with torch.no_grad():
        for batch in test_loader:
            input_ids = batch['input_ids'].to(self.device)
            attention_mask = batch['attention_mask'].to(self.device)
            labels = batch['label'].to(self.device)

            outputs = self.model(input_ids, attention_mask)
            _, predicted = torch.max(outputs.data, 1)

            total += labels.size(0)
            correct += (predicted == labels).sum().item()

    accuracy = 100 * correct / total
    return accuracy

def save_model(self, filename):
    """Save model and tokenizer"""
    model_path = os.path.join(self.output_dir, filename)
    torch.save({
        'model_state_dict': self.model.state_dict(),
        'model_name': self.model_name,
        'label_encoder': self.label_encoder,
        'class_weights': self.class_weights
    }, model_path)

    # Save tokenizer
    tokenizer_path = os.path.join(self.output_dir, 'tokenizer')
    self.tokenizer.save_pretrained(tokenizer_path)

    logger.info(f"💾 Model saved to {model_path}")

def plot_training_curves(self, train_losses, train_accs, test_accs):
    """Plot training curves"""
    fig, (ax1, ax2) = plt.subplots(1, 2, figsize=(15, 5))

    # Loss curve
    ax1.plot(train_losses, 'b-', label='Training Loss')

```

```

ax1.set_title('Training Loss')
ax1.set_xlabel('Epoch')
ax1.set_ylabel('Loss')
ax1.legend()
ax1.grid(True)

# Accuracy curves
ax2.plot(train_accs, 'b-', label='Training Accuracy')
ax2.plot(test_accs, 'r-', label='Test Accuracy')
ax2.set_title('Model Accuracy')
ax2.set_xlabel('Epoch')
ax2.set_ylabel('Accuracy (%)')
ax2.legend()
ax2.grid(True)

# Add target line
ax2.axhline(y=85, color='g', linestyle='--', alpha=0.7,
label='Target (85%)')
ax2.legend()
plt.tight_layout()
plot_path = os.path.join(self.output_dir, 'training_curves.png')
plt.savefig(plot_path, dpi=300, bbox_inches='tight')
plt.show()
logger.info(f"📊 Training curves saved to {plot_path}")

def run_training(self):
    """Complete training pipeline"""
    logger.info("🚀 Starting BERT training pipeline...")

    # Load data
    df = self.load_data()
    if df is None:
        return

    # Prepare data
    train_dataset, test_dataset = self.prepare_data(df,
balance_data=True)

    # Train model with optimized parameters
    best_accuracy = self.train_model(
        train_dataset,
        test_dataset,
        epochs=5, # Increased epochs
        batch_size=8, # Smaller batch size for better gradients
        lr=1e-5 # Lower learning rate for stability
    )
    if best_accuracy >= 85.0:
        logger.info(f"🎉 SUCCESS! Achieved target accuracy:
{best_accuracy:.2f}%")
    else:
        logger.info(f"📉 Current best accuracy: {best_accuracy:.2f}%
(Target: 85%)")
        logger.info("💡 Try: more epochs, different model, or data
augmentation")

    return best_accuracy

```

Додаток 3
Копія презентації



НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
“КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ ІМЕНІ ІГОРЯ
СІКОРСЬКОГО”



ФАКУЛЬТЕТ ПРИКЛАДНОЇ МАТЕМАТИКИ

КАФЕДРА ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ КОМП'ЮТЕРНИХ СИСТЕМ

СПОСІБ І ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ ДЛЯ ВИЗНАЧЕННЯ ПОТЕНЦІЙНОЇ ПОПУЛЯРНОСТІ ТЕКСТОВИХ ПОСТІВ В СОЦІАЛЬНИХ МЕРЕЖАХ

Студент: Довженко Роман Олександрович

Науковий керівник: к.т.н. Заболотня Тетяна Миколаївна

Київ – 2025

ПОСТАНОВКА ЗАДАЧІ

Мета роботи: Підвищення точності визначення потенційної популярності текстових постів у соціальних мережах.

Об'єкт дослідження: Процес інтелектуального аналізу текстового контенту для оцінювання його потенційної популярності

Предмет дослідження: Способи, алгоритми та програмні засоби прогнозування популярності текстових публікацій у соціальних мережах

ТЕРМІНОЛОГІЯ

BERT (Bidirectional Encoder Representations from Transformers) – технологія програмування, яка зв'язує бази даних з концепціями об'єктно-орієнтованих мов програмування.

Метадані – додаткова інформація про публікацію: час публікації, кількість вподобань, кількість підписників автора.

TF-IDF (Term Frequency-Inverse Document Frequency) – статистична міра важливості терміну в документі відносно колекції документів, що враховує частоту терміну та його унікальність.

Self-Attention механізм – архітектурний компонент трансформерних моделей, що дозволяє кожному елементу послідовності “звертати увагу” на всі інші елементи для визначення контекстуальних зв'язків.

Embedding – векторне представлення слів, речень або документів у багатовимірному просторі, де семантично схожі елементи розташовані близько один до одного.

ОГЛЯД ІСНУЮЧИХ СПОСОБІВ

TF-IDF

$$tf(t, d) = \frac{f_{t,d}}{\sum_{t' \in d} f_{t',d}}$$

$$idf(t, D) = \log \frac{N}{|\{d \in D : t \in d\}|}$$

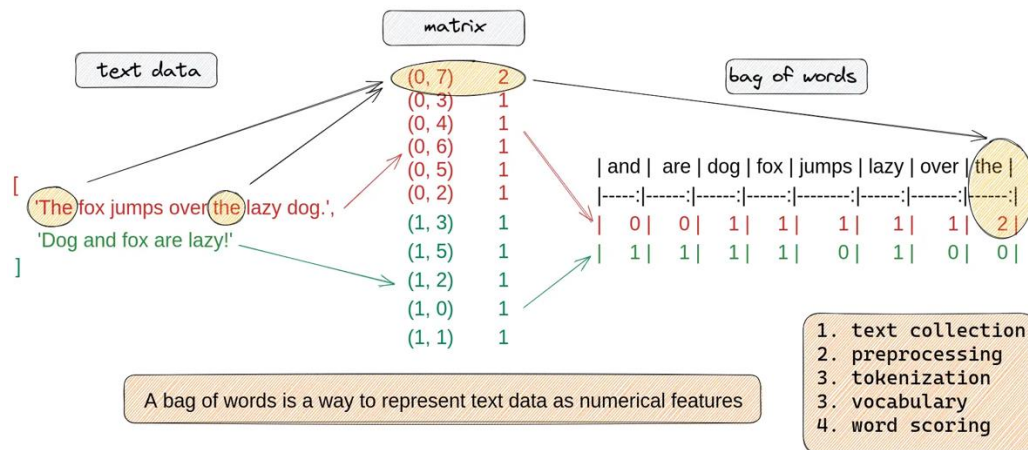
N-Gram

N=1: This is a sentence Uni-grams this, is, a, sentence

N=2: This is a sentence Bi-grams this is, is a, a sentence

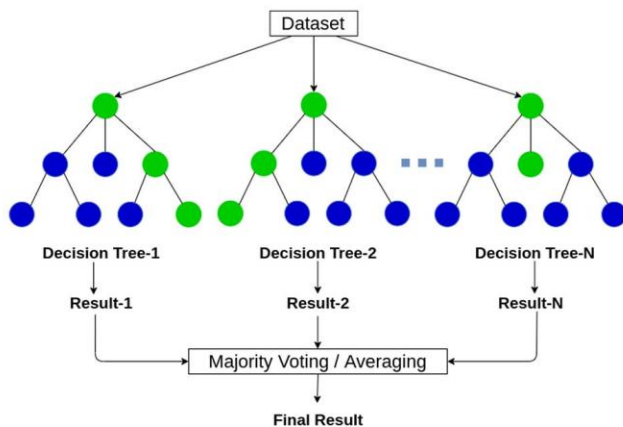
N=3: This is a sentence Tri-grams this is a, is a sentence

Bag of Words



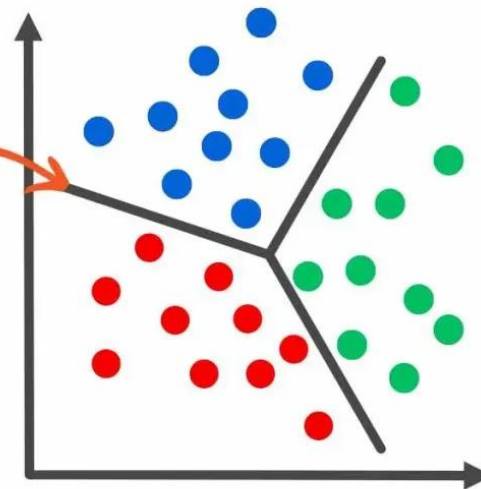
ОГЛЯД ІСНУЮЧИХ СПОСОБІВ

Random Forest

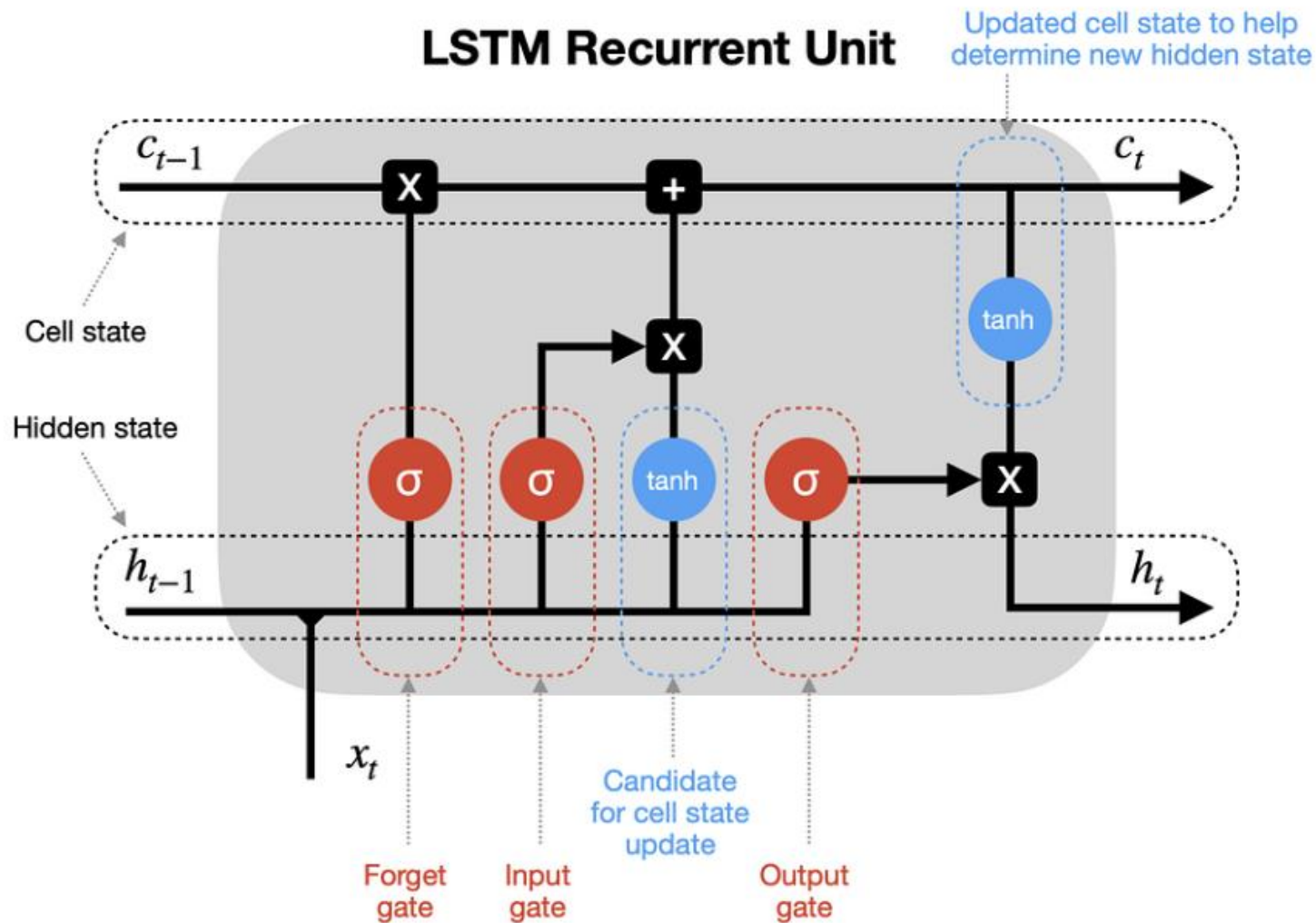


Support Vector Machines (SVM)

Hyperplanes that Best Separates Different Classes



ОГЛЯД ІСНУЮЧИХ СПОСОБІВ



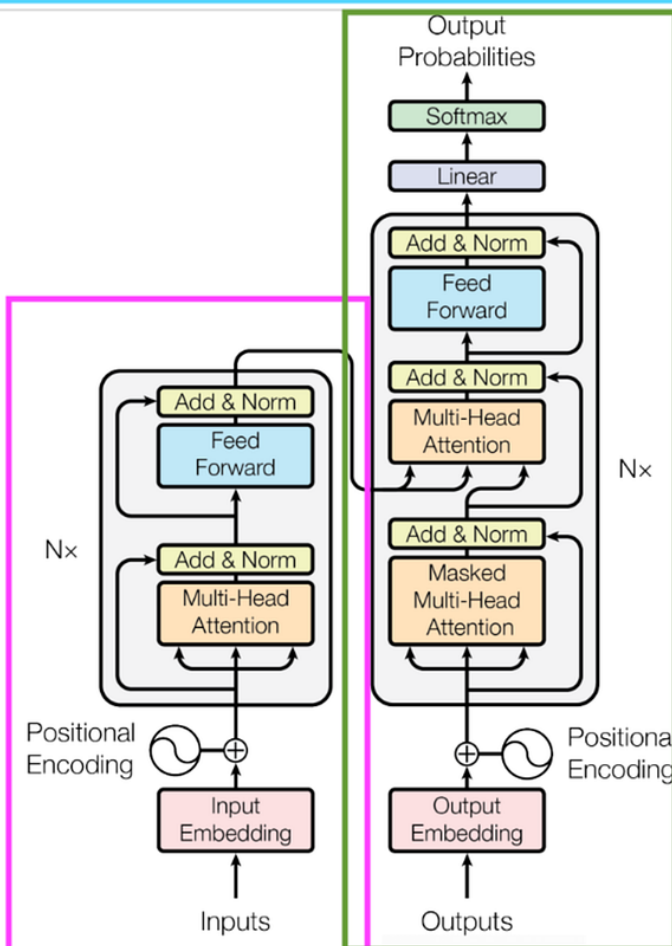
ОГЛЯД ІСНУЮЧИХ СПОСОБІВ



Encoder-Decoder

T5
BART

Encoder-only
BERT
ROBERTA



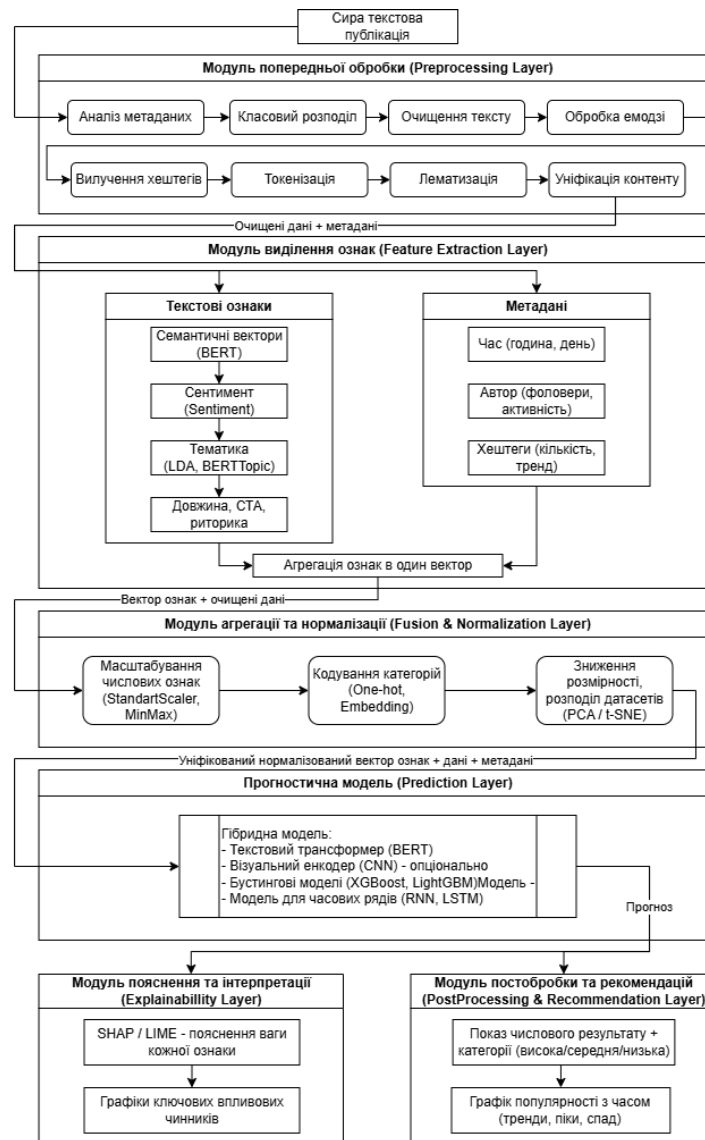
Decoder-only
GPT
BLOOM

ГІПОТЕЗА

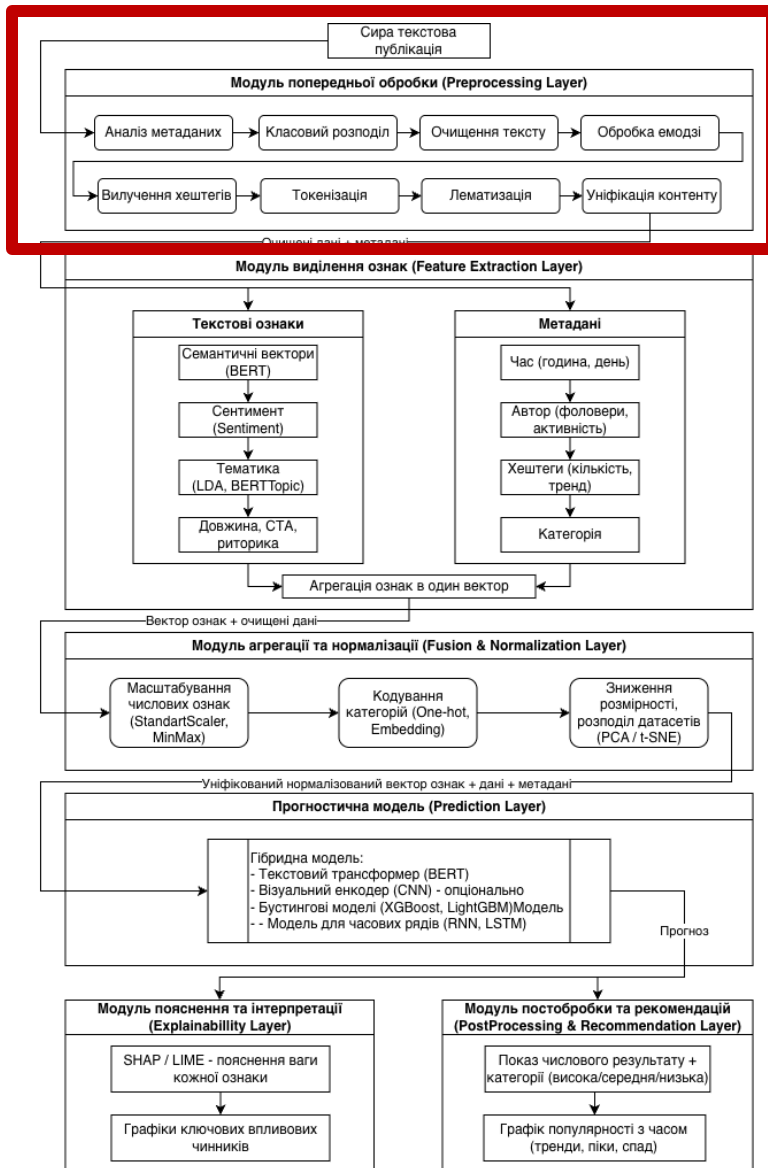


Використання комбінації базових методів машинного навчання (TF-IDF, статистичний аналіз метаданих, класичні алгоритми класифікації) з сучасними трансформерними архітектурами (BERT, RoBERTa) у запропонованому способі дозволяє досягти високої точності прогнозування потенційної популярності текстових постів у соціальних мережах порівняно з використанням базових способів.

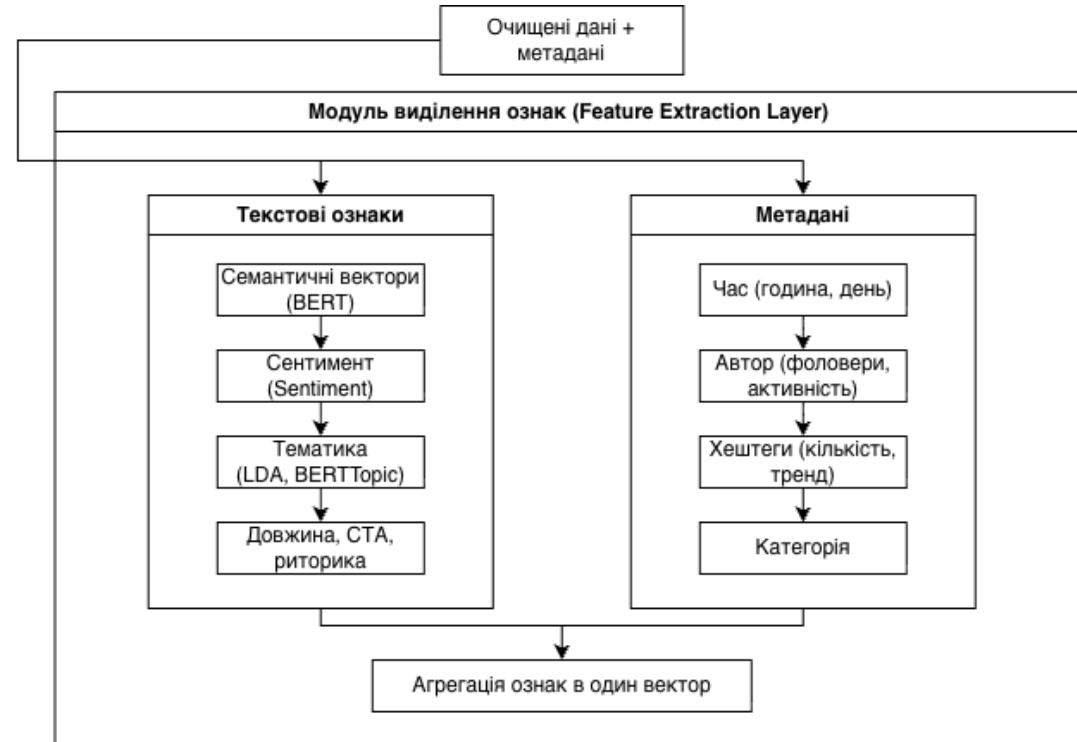
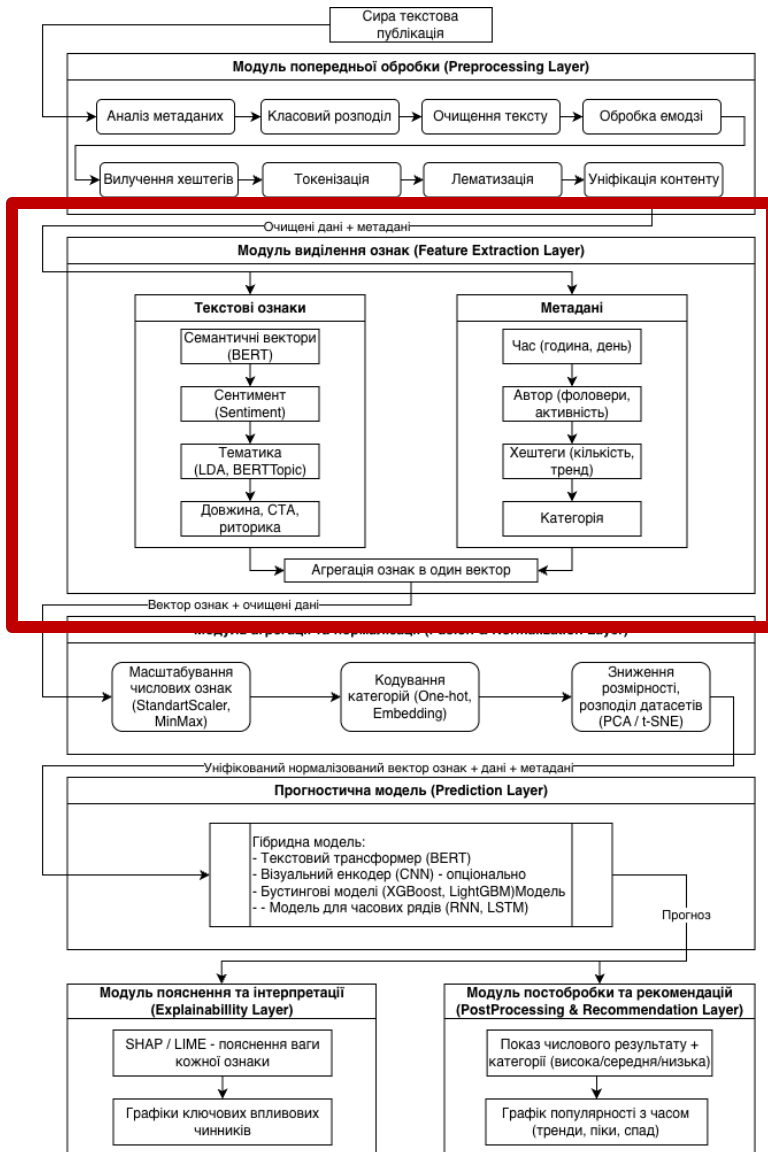
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



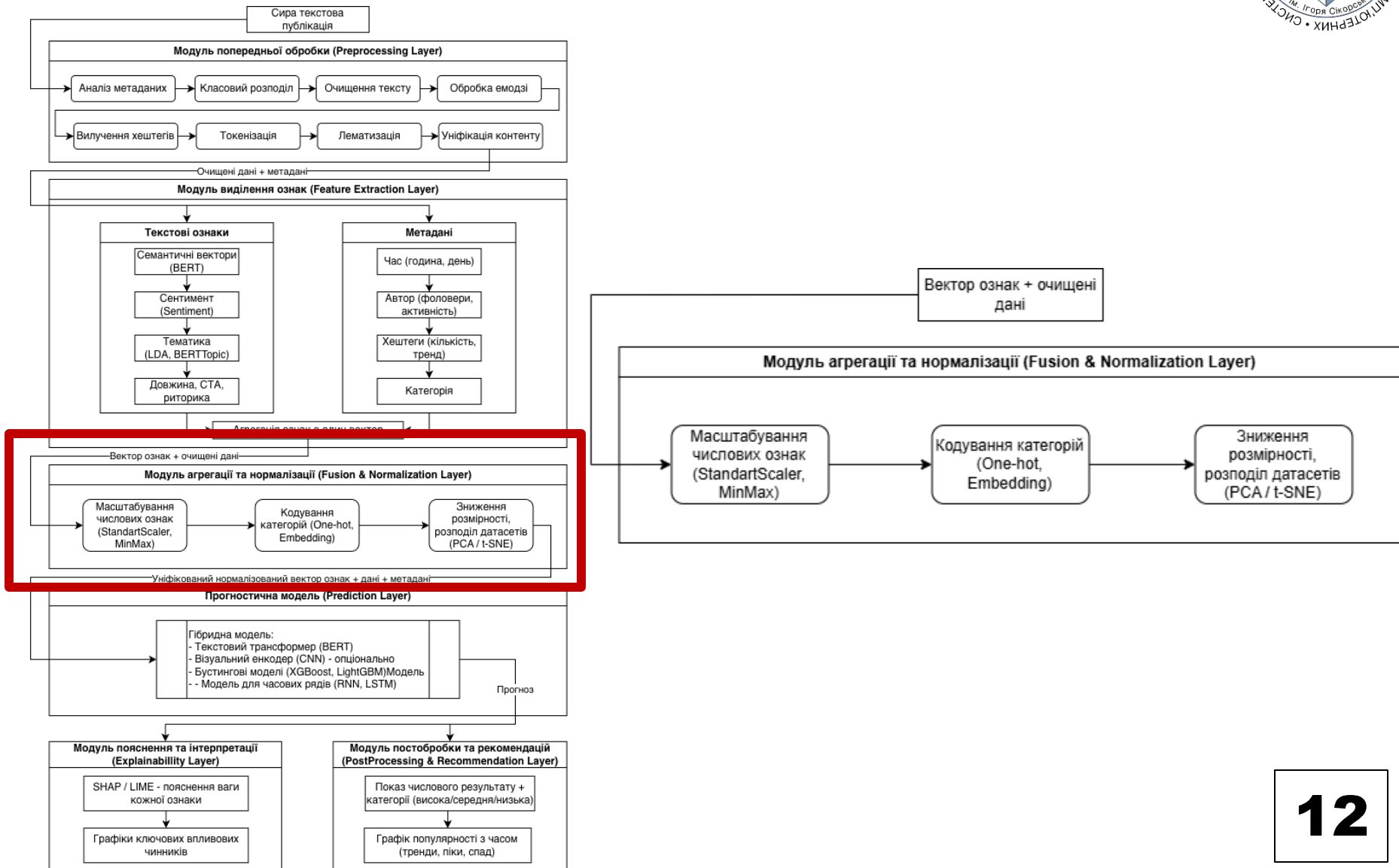
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



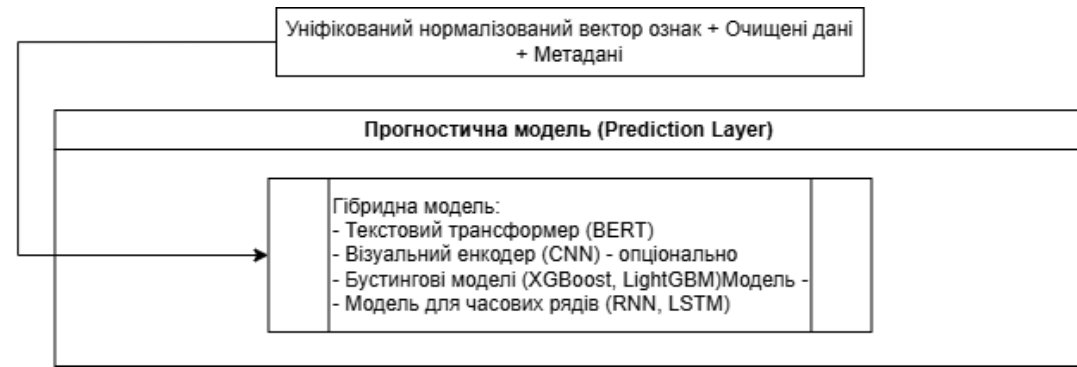
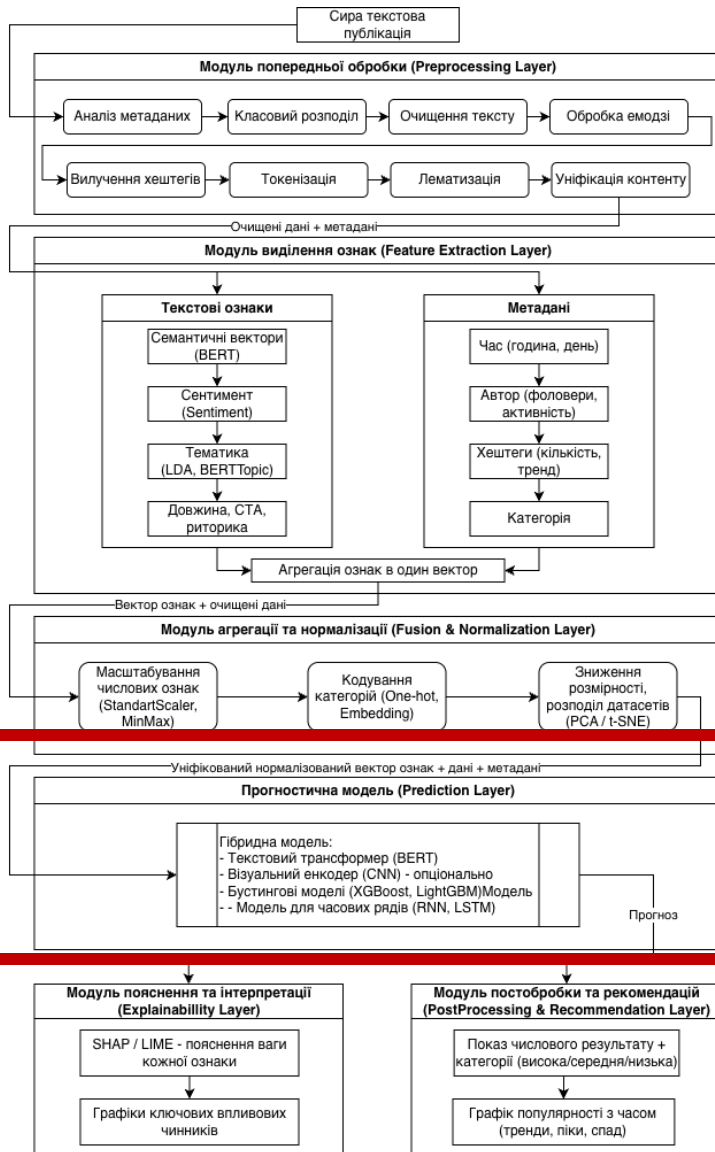
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



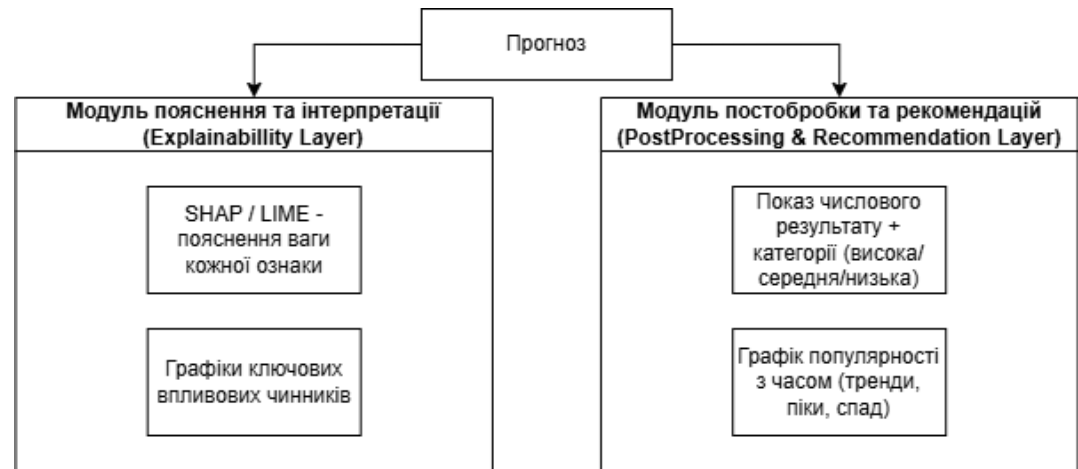
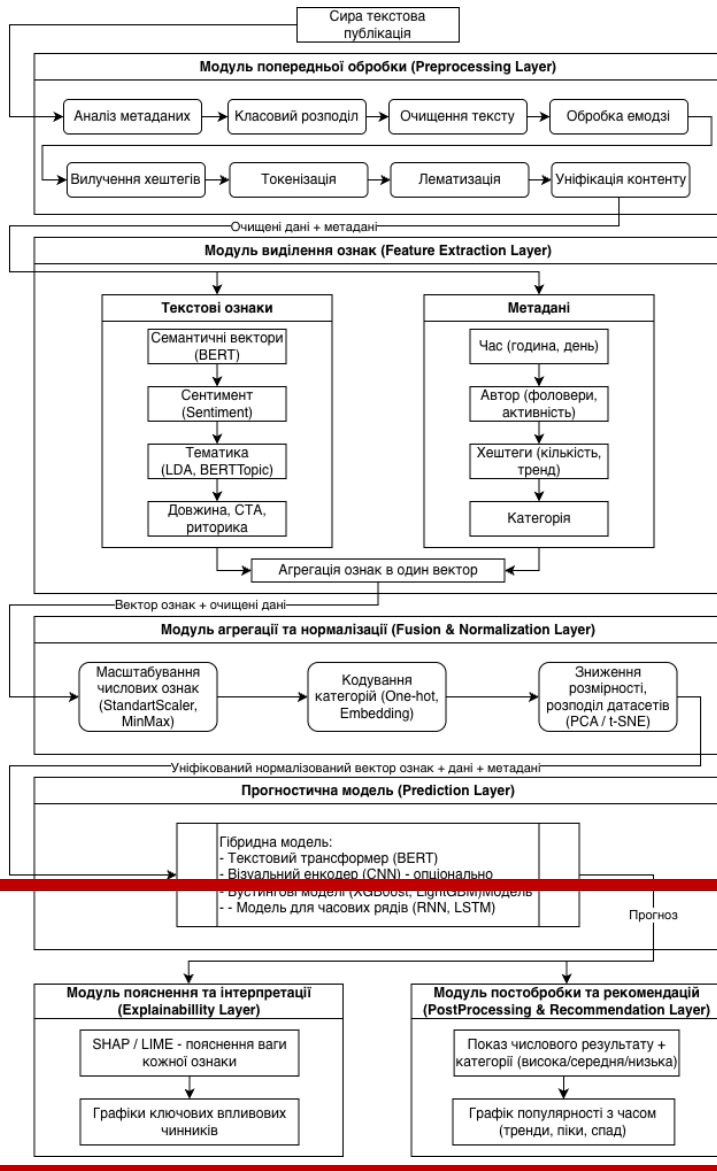
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



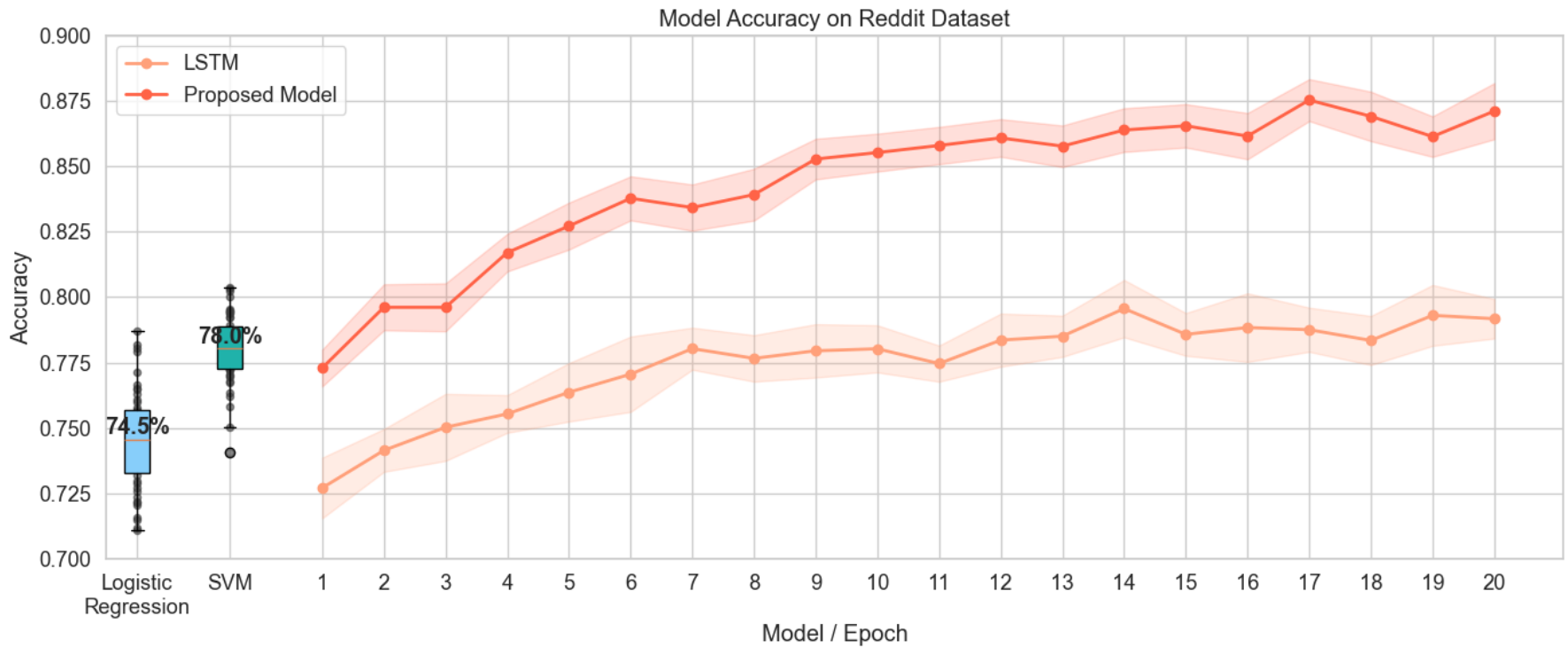
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



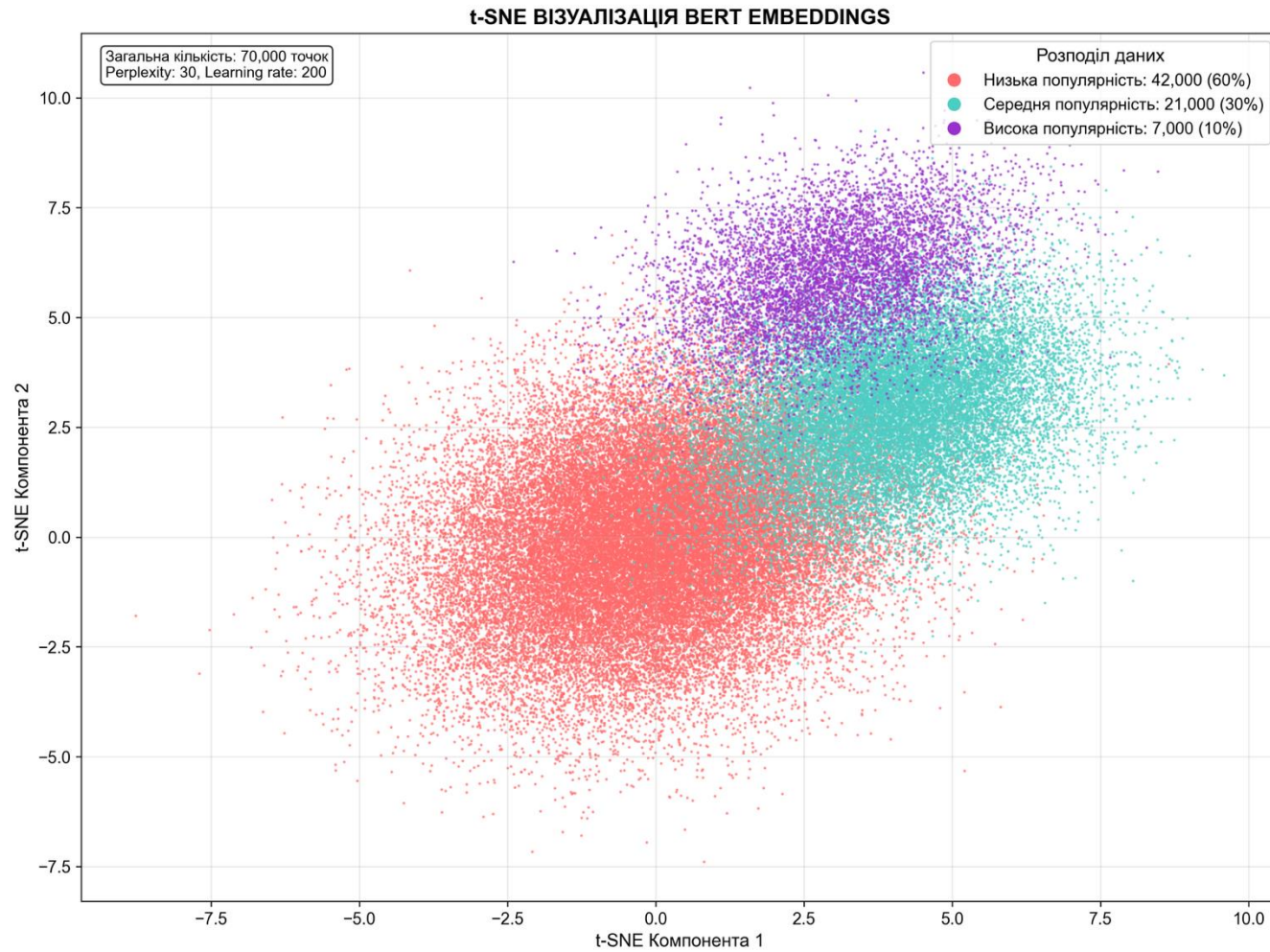
ОПИС ЗАПРОПОНОВАНОГО СПОСОБУ



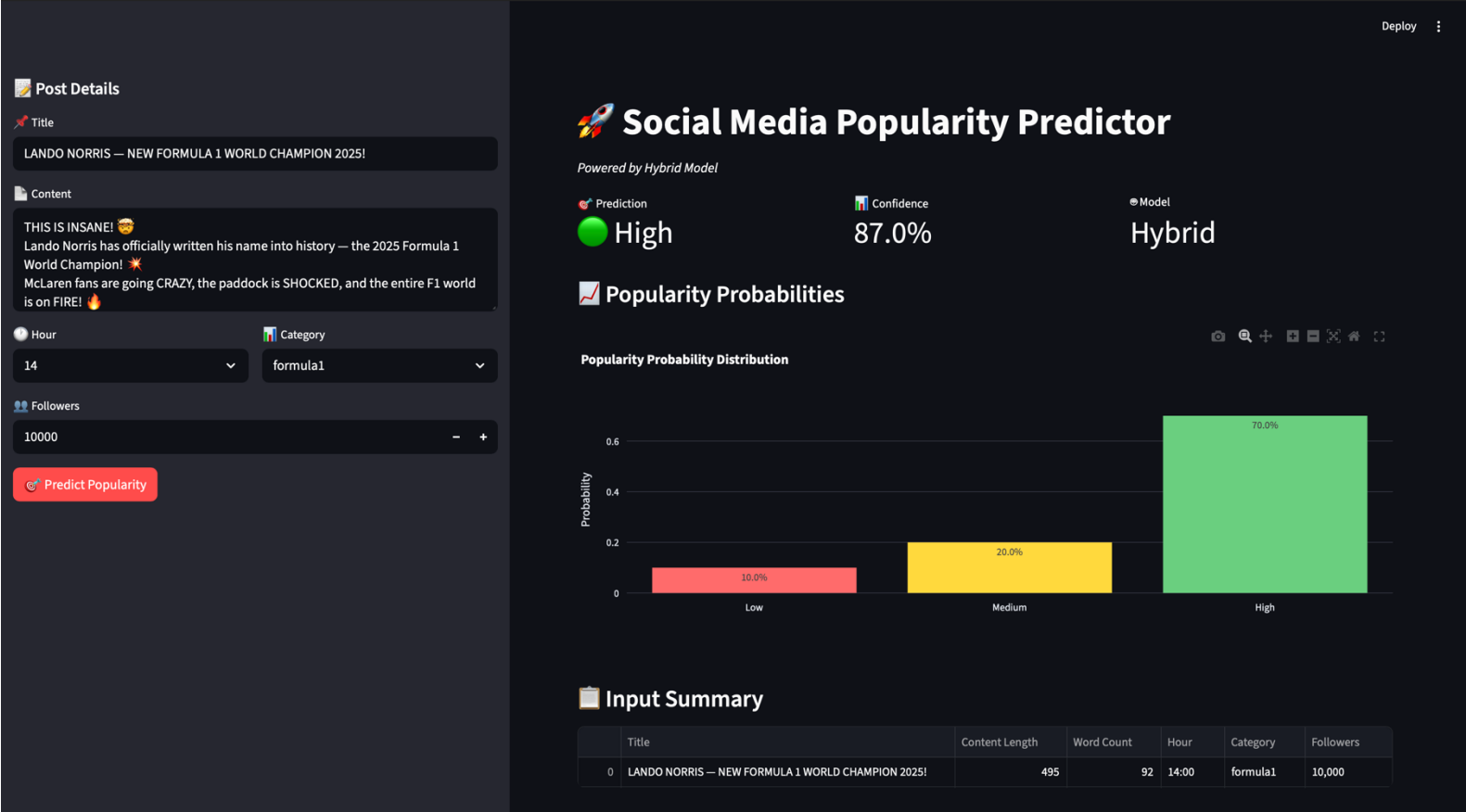
РЕЗУЛЬТАТИ



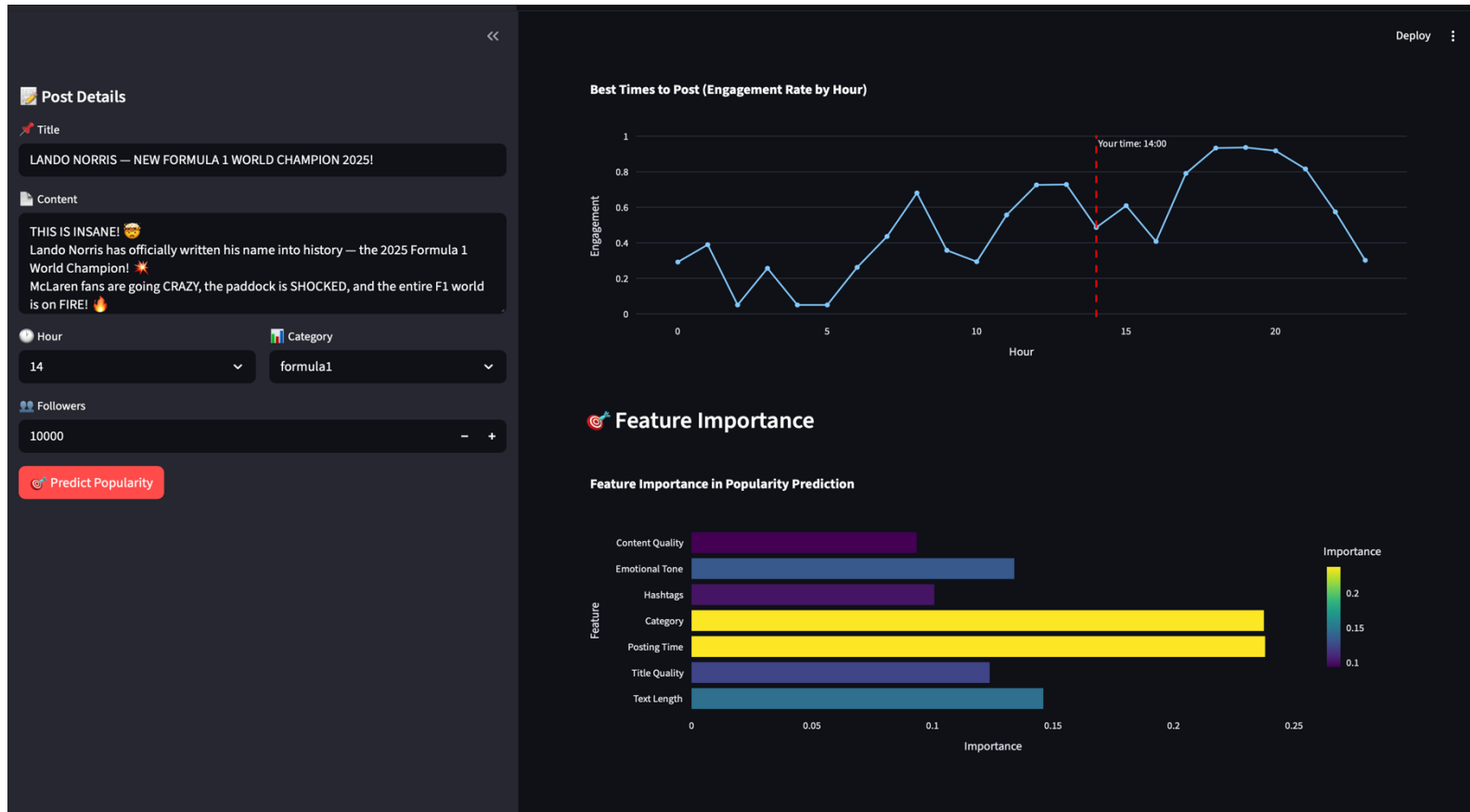
РЕЗУЛЬТАТИ



РЕЗУЛЬТАТИ



РЕЗУЛЬТАТИ



НАУКОВА НОВИЗНА



Вперше запропоновано спосіб оцінювання потенційної популярності текстових постів у соціальних мережах, що інтегрує семантичний аналіз тексту з використанням BERT-ембеддингів, статистичне оброблення метаданих та класичні алгоритми машинного навчання в єдиній архітектурі, що дозволило досягти точності 87% на експериментальній вибірці з 50 000 текстових публікацій платформи Reddit - на 10% вище за найкращий результат серед базових алгоритмів (логістична регресія, SVM, LSTM), підтверджуючи ефективність запропонованого способу до прогнозування популярності контенту в соціальних мережах.

ВИСНОВКИ



У магістерській дисертації досягнуто мети підвищення точності визначення потенційної популярності текстових постів у соціальних мережах за допомогою запропонованого способу прогнозування, що поєднує семантичний аналіз тексту на основі BERT-ембеддингів, обробку метаданих та класичні алгоритми машинного навчання в межах єдиної програмної архітектури.

Ефективність запропонованого способу перевірено експериментально на датасеті з 50 000 публікацій платформи Reddit. Результати експерименту показали, що точність прогнозування на рівні 87,5%, що перевищує результати найкращого базового способу в середньому на 8–10%.



Дякую за увагу!

Питання?