

Харківський національний університет імені В.Н. Каразіна

Факультет математики та інформатики

Кафедра прикладної математики

Кваліфікаційна робота

бакалавра

на тему *«Порівняння ефективності алгоритмів машинного навчання
для виявлення автоматизованих облікових записів (ботів) в онлайн
спільнотах»*

Виконав:

студент групи МП-41

4 курсу

спеціальність 113 – Прикладна Математика

освітньо-професійна програма

«Прикладна математика»

Данилов В. М.

Науковий керівник:

Кандидат психологічних наук, старший

викладач Сузікова О. Г.

Рецензент:

Доктор філософії,

Постдок за стипендією Марії Кюрі в
Університеті Осло,

Старший науковий співробітник

Фізико-технічного інституту низьких

температур ім. Б.І. Веркіна НАН України

Рибалко Я. В.

м. Харків – 2025 р.

АНОТАЦІЯ

Данилов Володимир Максимович. Порівняння ефективності алгоритмів машинного навчання для виявлення автоматизованих облікових записів (ботів) в онлайн спільнотах.

У сучасних умовах розвитку соціальних мереж і онлайн-спільнот зростає загроза поширення автоматизованих облікових записів (ботів), які здатні імітувати поведінку реальних користувачів, маніпулювати громадською думкою та порушувати цілісність інформаційного простору. Метою цієї бакалаврської роботи є розробка та експериментальна перевірка ефективного методу виявлення ботів на основі мультимодальної трансформер-архітектури, який забезпечує високу точність розпізнавання сучасних типів ботів із мінімізацією помилкових спрацьовувань.

У роботі поставлено такі завдання:

1. Провести аналіз еволюції методів виявлення ботів, від правилозалежних систем до сучасних гібридних рішень із використанням глибокого навчання.
2. Дослідити особливості застосування трансформер-архітектур і BERT-подібних моделей для аналізу текстової активності користувачів та їх переваги у виявленні ботів.
3. Розробити мультимодальну модель, що інтегрує текстові дані, часові патерни активності та метадані профілів, з використанням спеціалізованих механізмів уваги для аналізу часових залежностей.
4. Розробити процедури попередньої обробки, аугментації та балансування даних для подолання дисбалансу класів.
5. Провести експериментальне дослідження на стандартних наборах даних, порівняти результати з існуючими методами виявлення ботів та оцінити обчислювальну складність запропонованої моделі.
6. Розробити рекомендації щодо практичного застосування методу у реальних системах.

Об'єктом дослідження є процес виявлення автоматизованих облікових записів у соціальних мережах, предметом – алгоритми машинного навчання, зокрема мультимодальні трансформери для аналізу поведінки користувачів. У роботі використано методи глибокого навчання (трансформер-архітектури, механізми самоуваги), обробки природної мови, аналізу часових рядів та графових структур соціальних зв'язків. Наукова новизна полягає в запропонованій архітектурі мультимодальної трансформер-моделі з інтеграцією текстових, часових і структурних ознак разом із спеціалізованими механізмами уваги для виявлення характерних для ботів часових патернів. Практична значущість роботи полягає в можливості інтеграції розробленого методу в системи модерації контенту, кібербезпеки та електронної комерції для покращення якості виявлення ботів і зниження економічних збитків, пов'язаних зі зловмисною активністю.

Ключові слова: *автоматизований обліковий запис, бот, мультимодальний трансформер, BERT, механізм уваги, глибоке навчання, виявлення ботів, соціальні мережі.*

Volodymyr Danylov. Comparison of the effectiveness of machine learning algorithms for detecting automated accounts (bots) in online communities.

In the current environment of social networks and online communities, there is a growing threat of the proliferation of automated accounts (bots) that can imitate the behavior of real users, manipulate public opinion, and violate the integrity of the information space. The purpose of this bachelor's thesis is to develop and experimentally test an effective bot detection method based on a multimodal transformer architecture that provides high accuracy in recognizing modern types of bots while minimizing false positives.

The paper sets the following tasks:

1. To analyze the evolution of bot detection methods, from rule-based systems to modern hybrid solutions using deep learning.
2. To investigate the peculiarities of using transformer architectures and BERT-like models to analyze users' text activity and their advantages in detecting bots.
3. To develop a multimodal model that integrates text data, temporal activity patterns, and profile metadata, using specialized attention mechanisms to analyze time dependencies.
4. Develop data preprocessing, augmentation, and balancing procedures to overcome class imbalances.
5. Conduct an experimental study on standard datasets, compare the results with existing bot detection methods, and evaluate the computational complexity of the proposed model.
6. To develop recommendations for the practical application of the method in real systems.

The object of the study is the process of detecting automated accounts on social networks, and the subject is machine learning algorithms, in particular

multimodal transformers for analyzing user behavior. The study uses methods of deep learning (transformer architectures, self-attention mechanisms), natural language processing, time series analysis, and graph structures of social connections. The scientific novelty lies in the proposed architecture of a multimodal transformer model integrating textual, temporal, and structural features along with specialized attention mechanisms to detect temporal patterns characteristic of bots. The practical significance of the work lies in the possibility of integrating the developed method into content moderation, cybersecurity, and e-commerce systems to improve the quality of bot detection and reduce economic losses associated with malicious activity.

Keywords: *automated account, bot, multimodal transformer, BERT, attention mechanism, deep learning, bot detection, social networks.*

ЗМІСТ

Перелік умовних позначень	8
Вступ	9
Розділ 1. Розвиток сучасних методів виявлення ботів та особливості їх застосування	13
1.1. Еволюція методів виявлення ботів: від правилозалежних до гібридних систем	13
1.2. Актуальні проблеми виявлення ботів у соціальних мережах	17
1.3. Особливості використання методів глибокого навчання для ідентифікації ботів	21
1.3.1. Трансформери та їх переваги у виявленні ботів	26
1.3.2. Особливості застосування BERTology-моделей для виявлення соціальних ботів	29
1.4. Мультимодальні підходи та їх ефективність	33
Розділ 2. Розробка удосконаленого методу для виявлення ботів на основі мультимодальної трансформер-архітектури	39
2.1. Архітектура запропонованої моделі	39
2.1.1. Вибір базової моделі та її модифікація	39
2.1.2. Інтеграція мультимодальних вхідних даних	40
2.1.3. Механізм уваги для аналізу часових патернів активності	43
2.2. Підготовка та попередня обробка даних	45
2.2.1. Джерела даних та методи збору	45
2.2.2. Попередня обробка текстових та метаданих	48
2.2.3. Аугментація даних для боротьби з дисбалансом класів	50
2.3. Особливості навчання моделі	52
2.3.1. Стратегія навчання та гіперпараметри	52
2.3.2. Техніки регуляризації та запобігання перенавчанню	54
2.3.3. Інтерпретованість моделі	56

Розділ 3. Експериментальне дослідження ефективності запропонованого методу	60
3.1. Опис експериментального середовища	60
3.1.1. Використані набори даних	60
3.1.2. Метрики оцінки ефективності	61
3.1.3. Базові моделі для порівняння	63
3.2. Результати експериментів	64
3.2.1. Порівняльний аналіз з існуючими методами	64
3.2.2. Вплив окремих компонентів моделі на загальну ефективність ...	67
3.2.3. Аналіз помилок та обмежень	70
3.3. Аналіз обчислювальної складності	71
3.3.1. Часова та просторова складність	71
3.3.2. Вимоги до обчислювальних ресурсів	73
3.4. Практичні рекомендації щодо застосування методу	74
Висновки	76
Список використаних джерел	79
Додатки	86
Додаток А. Лістинг програмного коду	86
Додаток Б. Результати додаткових експериментів	123

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ

- АОЗ** – автоматизований обліковий запис
- API** – Application Programming Interface (інтерфейс програмування додатків)
- AUC** – Area Under Curve (площа під кривою)
- BERT** – Bidirectional Encoder Representations from Transformers
- BiLSTM** – Bidirectional Long Short-Term Memory (двонаправлена довга короткочасна пам'ять)
- CNN** – Convolutional Neural Network (згорткова нейронна мережа)
- CPU** – Central Processing Unit (центральний процесор)
- DNS** – Domain Name System (система доменних імен)
- FN** – False Negative (помилково негативний результат)
- FP** – False Positive (помилково позитивний результат)
- FPR** – False Positive Rate (рівень помилково позитивних результатів)
- GAN** – Generative Adversarial Network (генеративно-змагальна мережа)
- GCN** – Graph Convolutional Network (графова згорткова мережа)
- GDPR** – General Data Protection Regulation (загальний регламент захисту даних)
- GPU** – Graphics Processing Unit (графічний процесор)
- GPT** – Generative Pre-trained Transformer
- GRU** – Gated Recurrent Unit (керована рекурентна одиниця)
- HTML** – HyperText Markup Language (мова розмітки гіпертексту)
- LSTM** – Long Short-Term Memory (довга короткочасна пам'ять)
- MCC** – Matthews Correlation Coefficient (коефіцієнт кореляції Метьюза)
- MLP** – Multi-Layer Perceptron (багатошаровий перцептрон)
- NLP** – Natural Language Processing (обробка природної мови)
- PR** – Precision-Recall
- ReLU** – Rectified Linear Unit (випрямлена лінійна одиниця)
- RNN** – Recurrent Neural Network (рекурентна нейронна мережа)
- ROC** – Receiver Operating Characteristic
- RoBERTa** – Robustly Optimized BERT Pretraining Approach
- SHAP** – SHapley Additive exPlanations
- SMM** – Social Media Marketing (маркетинг у соціальних мережах)
- SMOTE** – Synthetic Minority Over-sampling Technique
- SVM** – Support Vector Machine (метод опорних векторів)
- TF-IDF** – Term Frequency-Inverse Document Frequency
- TN** – True Negative (істинно негативний результат)
- TP** – True Positive (істинно позитивний результат)
- TPR** – True Positive Rate (рівень істинно позитивних результатів)
- URL** – Uniform Resource Locator (уніфікований локатор ресурсу)
- VAE** – Variational Autoencoder (варіаційний автоенкодер)
- XML** – eXtensible Markup Language (розширювана мова розмітки)
- ШІ** – штучний інтелект

ВСТУП

Стрімкий розвиток соціальних мереж та онлайн-платформ призвів до появи нових викликів у сфері інформаційної безпеки та якості контенту. Однією з найбільш актуальних проблем сучасного цифрового простору є поширення автоматизованих облікових записів, або ботів, які здатні імітувати поведінку реальних користувачів та впливати на громадську думку, поширювати дезінформацію, маніпулювати ринками та порушувати цілісність онлайн-дискусій.

За даними досліджень, проведених у 2023-2024 роках, від 9% до 15% активних облікових записів у провідних соціальних мережах становлять боти різного типу. Це означає, що мільйони автоматизованих систем щодня взаємодіють з реальними користувачами, часто залишаючись непомітними для традиційних методів виявлення. Економічні збитки від діяльності шкідливих ботів оцінюються в мільярди доларів щорічно, включаючи втрати від фальшивої реклами, маніпуляцій з цінами акцій та поширення неправдивої інформації.

Традиційні підходи до виявлення ботів, що базувалися на простих евристичних правилах та статистичному аналізі, втрачають ефективність у протистоянні з сучасними автоматизованими системами. Нове покоління ботів використовує передові технології штучного інтелекту, включаючи генеративні мовні моделі, для створення високоякісного контенту та імітації складних патернів людської поведінки. Це створює необхідність розробки принципово нових методів виявлення, здатних аналізувати множинні аспекти поведінки користувачів та виявляти тонкі відмінності між автентичними та штучними взаємодіями.

Актуальність теми дослідження обумовлюється кількома ключовими факторами. По-перше, зростанням складності та адаптивності сучасних ботів, які використовують технології машинного навчання для покращення своїх стратегій маскування. По-друге, збільшенням масштабів впливу ботів на соціально-політичні процеси, економічні ринки та громадську думку. По-

третє, недостатністю існуючих методів виявлення для ефективного протистояння новим типам автоматизованих загроз.

Сучасні дослідження в області виявлення ботів демонструють перспективність мультимодальних підходів, які аналізують не лише текстовий контент, але й часові патерни активності, метадані профілів та структуру соціальних зв'язків. Особливо ефективними виявляються методи на основі трансформер-архітектур, які здатні моделювати складні залежності в даних та виявляти неочевидні закономірності поведінки ботів.

Мета дослідження полягає в розробці та експериментальній перевірці ефективного методу виявлення автоматизованих облікових записів у соціальних мережах на основі мультимодальної трансформер-архітектури, здатного забезпечити високу точність розпізнавання сучасних типів ботів при мінімізації помилкових спрацьовувань.

Основні задачі дослідження:

1. Провести аналіз сучасного стану методів виявлення ботів у соціальних мережах, дослідити еволюцію підходів від правилозалежних систем до глибокого навчання та виявити основні проблеми та обмеження існуючих рішень.

2. Дослідити особливості застосування трансформер-архітектур для задач виявлення ботів, проаналізувати переваги та недоліки BERT-подібних моделей у контексті аналізу поведінки користувачів соціальних мереж.

3. Розробити архітектуру мультимодальної моделі, що ефективно інтегрує текстові дані, часові патерни активності та метадані профілів для комплексного аналізу поведінки користувачів.

4. Створити та впровадити спеціалізовані механізми уваги для аналізу часових залежностей в активності користувачів, що дозволяють виявляти характерні для ботів патерни поведінки.

5. Розробити методи попередньої обробки та аугментації мультимодальних даних для подолання проблеми дисбалансу класів та покращення генералізаційних властивостей моделі.

6. Провести всебічне експериментальне дослідження ефективності запропонованого методу на стандартних наборах даних і порівняти результати з існуючими підходами до виявлення ботів.

7. Здійснити аналіз обчислювальної складності запропонованого методу та розробити практичні рекомендації щодо його застосування в реальних системах.

Об'єктом дослідження є процес виявлення автоматизованих облікових записів у соціальних мережах та онлайн-спільнотах.

Предметом дослідження є алгоритми машинного навчання, зокрема мультимодальні трансформер-архітектури, для аналізу поведінки користувачів і розпізнавання ботів на основі комплексного аналізу текстових, часових та структурних даних.

Методи дослідження включають методи глибокого навчання та штучних нейронних мереж, зокрема трансформер-архітектури та механізми уваги; методи обробки природної мови для аналізу текстового контенту; статистичні методи аналізу часових рядів для виявлення патернів активності; методи аналізу графів для дослідження соціальних зв'язків; техніки машинного навчання для роботи з дисбалансованими даними; методи оцінки ефективності класифікаційних моделей.

Наукова новизна роботи полягає в розробці нового підходу до виявлення ботів, що поєднує переваги трансформер-архітектур з мультимодальним аналізом даних. Запропонована архітектура включає спеціалізовані механізми уваги для аналізу часових патернів, які дозволяють виявляти складні закономірності в поведінці користувачів на різних часових масштабах. Розроблено гібридну стратегію інтеграції мультимодальних даних, що забезпечує ефективне поєднання текстової, часової та структурної інформації.

Практична значущість результатів дослідження визначається можливістю застосування розробленого методу в реальних системах модерації контенту соціальних мереж, системах кібербезпеки та платформах електронної комерції. Запропоновані алгоритми можуть бути інтегровані в

існуючі системи виявлення шахрайства та забезпечення інформаційної безпеки. Практичні рекомендації щодо застосування методу дозволяють адаптувати його для конкретних потреб різних онлайн-платформ.

Структура та обсяг роботи. Дипломна робота складається зі вступу, трьох розділів, висновків, списку використаних джерел та додатків. Загальний обсяг роботи становить 118 сторінок, включаючи 5 рисунків, 43 таблиці. Список використаних джерел містить 61 найменувань.

РОЗДІЛ 1.

РОЗВИТОК СУЧАСНИХ МЕТОДІВ ВИЯВЛЕННЯ БОТІВ ТА ОСОБЛИВОСТІ ЇХ ЗАСТОСУВАННЯ

1.1. Еволюція методів виявлення ботів: від правилозалежних до гібридних систем

Розвиток інформаційних технологій та зростання популярності соціальних мереж призвели до появи нового виду загрози – автоматизованих облікових записів, або ботів. Боти в сучасному контексті – це програмне забезпечення, яке виконує передбачувані завдання автоматизовано та повторно [1]. Еволюція ботів від простих скриптів до складних систем, здатних імітувати людську поведінку, стала серйозним викликом для безпеки онлайн-платформ.

Еволюцію методів виявлення ботів можна умовно розділити на кілька етапів, які відображають зміну підходів до вирішення проблеми виявлення автоматизованих облікових записів.

Перший етап: правилозалежні системи (2000-2010 рр.). На початку розвитку соціальних мереж боти мали примітивний функціонал і демонстрували легко розпізнавані шаблони поведінки. Цей період характеризувався використанням статичних правил та простих евристик для виявлення ботів [2]. Основними ознаками, що використовувалися для ідентифікації ботів, були:

1. Частота публікацій – аномально висока активність облікового запису;
2. Шаблонні повідомлення – використання однотипних текстів;
3. Базові часові патерни – активність у нехарактерний для людини час.

Перевагою таких методів була простота реалізації та низькі обчислювальні вимоги. Однак, ці системи мали суттєві недоліки: висока

чутливість до змін у стратегії поведінки ботів, необхідність ручного оновлення правил при виявленні нових типів ботів, низька адаптивність до еволюції автоматизованих облікових записів. Правилозалежні системи втратили ефективність з розвитком технологій створення ботів, які почали уникати типових шаблонів поведінки [3]. Це призвело до необхідності розробки більш складних методів виявлення. Другий етап: машинне навчання та статистичні методи (2010-2015 рр.). У відповідь на ускладнення механізмів функціонування ботів дослідники почали застосовувати методи машинного навчання. Цей період характеризувався використанням класичних алгоритмів машинного навчання, таких як: метод опорних векторів (SVM), дерева рішень та ансамблеві методи (Random Forest, AdaBoost), наївний баєсів класифікатор, k-найближчих сусідів (k-NN). Використання алгоритмів машинного навчання дозволило автоматизувати процес виявлення ботів та підвищити точність ідентифікації [4]. Основні переваги підходів цього періоду: здатність виявляти неочевидні закономірності в даних, можливість адаптації до нових типів ботів шляхом перенавчання моделей, поєднання різних типів ознак (профілю, контенту, мережевої активності).

Математично, задача виявлення ботів формулювалась як задача бінарної класифікації:

$$f: X \rightarrow Y, Y \in \{0,1\}$$

де (X) - вектор ознак облікового запису, (Y) - мітка класу (0 - людина, 1 - бот).

Ознаки, що використовувалися для класифікації, можна розділити на кілька груп [5]:

1. Ознаки облікового запису: вік профілю, наявність аватару, опис профілю
2. Ознаки контенту: різноманітність лексики, частота використання хештегів, URL
3. Часові ознаки: розподіл активності за годинами доби, інтервали між публікаціями

4. Мережеві ознаки: структура зв'язків, співвідношення підписників/підписок

Незважаючи на значний прогрес, методи цього періоду мали обмеження, пов'язані з необхідністю ручного проектування ознак та складністю обробки неструктурованих даних, таких як текст і зображення.

Третій етап: глибоке навчання та нейронні мережі (2015-2020 рр.). З розвитком глибокого навчання почали з'являтися моделі, здатні автоматично вилучати ознаки з даних різних типів. Ключові архітектури цього періоду:

- Згорткові нейронні мережі (CNN) для обробки зображень і аналізу послідовностей
- Рекурентні нейронні мережі (RNN, LSTM, GRU) для аналізу часових рядів і тексту
- Автоенкодери для виявлення аномалій та зменшення розмірності даних

Глибокі нейронні мережі показали високу ефективність у виявленні складних автоматизованих облікових записів, особливо в контексті обробки текстових даних [6]. Наприклад, LSTM-мережі дозволили моделювати довгострокові залежності в послідовностях активності користувачів:

$$h_t = \tanh(W_c \cdot [f_t \odot c_{t-1} + i_t \odot \tilde{c}_t] + b_c)$$

де (h_t) - прихований стан на кроці (t) , (c_t) - стан комірки, (f_t) , (i_t) - вектори вентилів (gates).

Основні переваги підходів на основі глибокого навчання:

- Автоматичне вилучення високорівневих ознак
- Здатність працювати з неструктурованими даними (текст, зображення)
- Висока точність виявлення складних шаблонів поведінки

Проте, моделі глибокого навчання потребували значних обчислювальних ресурсів і великих обсягів даних для навчання, що обмежувало їх практичне застосування [7].

Сучасний етап: гібридні та мультимодальні системи (2020-теперішній час). Сучасний етап розвитку методів виявлення ботів характеризується використанням гібридних підходів, які поєднують переваги різних типів моделей та обробляють дані різної модальності [8]. Ключові тенденції цього періоду:

1. Мультимодальний аналіз - інтеграція даних різних типів (текст, зображення, часові ряди, графові структури)
2. Трансформер-архітектури - використання механізмів уваги для моделювання залежностей у даних
3. Самонавчання та напівконтрольоване навчання - подолання проблеми обмеженості розмічених даних
4. Інтерпретовані моделі - забезпечення пояснюваності рішень систем виявлення

Особливу роль у розвитку сучасних методів виявлення ботів відіграли моделі-трансформери, такі як BERT і його модифікації [9]. Ці моделі дозволяють ефективно обробляти текстові дані та контекстуалізувати інформацію:

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

де (Q), (K), (V) - матриці запитів, ключів та значень відповідно.

Гібридні системи поєднують статистичний аналіз, класичне машинне навчання та глибокі нейронні мережі для досягнення максимальної ефективності. Вони демонструють високу точність виявлення навіть найскладніших типів ботів, включаючи ті, що імітують людську поведінку [10].

Графічно еволюцію методів виявлення ботів можна представити як поступове підвищення складності моделей та їх здатності виявляти все більш досконалі автоматизовані облікові записи (рис. 1.1).

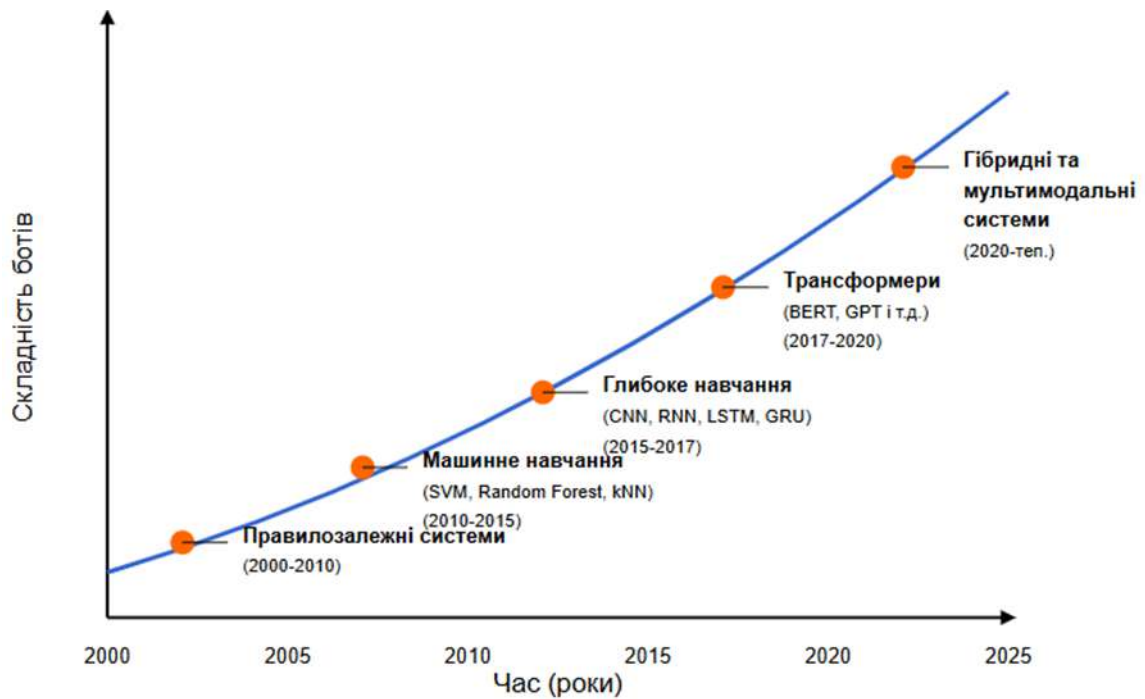


Рисунок 1.1 Еволюція методів виявлення ботів та їх співвідношення зі складністю ботів

Таким чином, еволюція методів виявлення ботів відображає складну взаємодію між розвитком технологій створення автоматизованих облікових записів та розробкою все більш досконалих систем їх ідентифікації. Сучасні підходи до виявлення ботів вимагають комплексного аналізу даних та інтеграції різних типів моделей для досягнення максимальної ефективності.

1.2. Актуальні проблеми виявлення ботів у соціальних мережах

Незважаючи на значний прогрес у розробці методів виявлення автоматизованих облікових записів, ця задача залишається складною через ряд актуальних проблем, які потребують вирішення. Розглянемо ключові виклики, з якими стикаються сучасні системи виявлення ботів у соціальних мережах.

Однією з найбільших проблем у виявленні ботів є їх постійна еволюція та адаптація до існуючих методів виявлення [11]. Згідно з дослідженнями [1], сучасні боти можуть імітувати поведінку реальних користувачів, використовуючи: імітацію людської поведінки - боти стали значно

складнішими, перейшовши від примітивних скриптів до систем, що використовують елементи штучного інтелекту для імітації природної поведінки користувачів. За даними [1], сучасні боти здатні генерувати унікальний контент з використанням моделей мови, наслідувати паттерни активності реальних користувачів, взаємодіяти з іншими користувачами в природній манері; масштабування реальної людської поведінки - використання мікрозадач (microwork), де реальні люди виконують прості завдання за невелику плату, що дозволяє зловмисникам створювати атаки, які ще більше наближаються до реальної поведінки людей [1]. Ця проблема ускладнюється постійним зростанням доступності інструментів для створення ботів. Розробники ботів використовують сучасні технології, такі як GPT-моделі та інші генеративні системи, для створення все більш правдоподібного контенту [12].

Дисбаланс класів є критичною проблемою при розробці систем виявлення ботів. У реальних соціальних мережах частка ботів зазвичай значно менша за частку реальних користувачів, що створює нерівномірний розподіл класів у навчальних даних [13]. Математично проблему дисбалансу можна виразити як:

$$P(y = 0) \gg P(y = 1)$$

де $(P(y = 0))$ - ймовірність класу "людина", $(P(y = 1))$ - ймовірність класу "бот".

Це призводить до кількох негативних наслідків:

1. Моделі схильні оптимізувати точність на домінуючому класі
2. Висока загальна точність може приховувати низьку здатність виявляти боти

3. Стандартні метрики оцінки (ассурасу) стають непоказовими

Для вирішення цієї проблеми використовуються різні підходи:

- Техніки збалансування даних (oversampling, undersampling)
- Спеціалізовані функції втрат, що враховують дисбаланс класів

- Метрики оцінки, орієнтовані на виявлення міноритарного класу (F1-score, AUC-ROC)

Крім дисбалансу класів, суттєвою проблемою є обмеженість розмічених даних. Створення якісних наборів даних для навчання моделей виявлення ботів вимагає значних зусиль з маркування, а в багатьох випадках точне визначення, чи є обліковий запис ботом, потребує експертної оцінки [14].

Розробка систем виявлення ботів тісно пов'язана з питаннями конфіденційності та етичних обмежень. Збір та аналіз даних користувачів для виявлення ботів може порушувати право на приватність [15]. Ключові аспекти цієї проблеми:

1. Законодавчі обмеження - такі регуляторні акти, як GDPR в Європі, встановлюють суворі вимоги до обробки персональних даних
2. Етичні питання - масштабний моніторинг активності користувачів викликає етичні питання
3. Баланс між безпекою та приватністю - необхідність знаходження компромісу між захистом платформи від зловмисних ботів і повагою до приватності користувачів

Ці обмеження впливають на практичне застосування методів виявлення ботів і вимагають розробки підходів, які мінімізують використання чутливих даних користувачів.

Сучасні соціальні мережі характеризуються масштабістю та високою динамічністю, що створює додаткові виклики для систем виявлення ботів. Ключові аспекти цієї проблеми:

1. Масштабованість - соціальні мережі налічують мільйони або навіть мільярди користувачів, що вимагає високоефективних алгоритмів
2. Реальний час - системи виявлення повинні працювати в режимі реального часу, швидко реагуючи на появу нових ботів
3. Ресурсомісткість - сучасні методи глибокого навчання вимагають значних обчислювальних ресурсів

Аналіз обчислювальної складності різних методів виявлення ботів показує, що існує компроміс між точністю та ефективністю [16]. Наприклад, трансформер-моделі демонструють високу точність, але мають квадратичну складність відносно довжини вхідних послідовностей:

$$O(n^2 \cdot d)$$

де (n) - довжина послідовності, (d) - розмірність моделі.

Це обмежує їх застосування в системах реального часу, особливо при аналізі великих обсягів даних.

Соціальні мережі характеризуються наявністю даних різних типів: текст, зображення, відео, часові ряди активності, графові структури зв'язків між користувачами. Ефективне виявлення ботів вимагає інтеграції цих різномірних даних, що створює додаткові виклики [17]:

1. Узгодження різних типів даних - приведення даних різної природи до єдиного представлення
2. Різні швидкості оновлення - різні типи даних оновлюються з різною частотою
3. Різна інформативність - різні типи даних можуть мати різну цінність для виявлення ботів

Сучасні мультимодальні підходи намагаються вирішити ці проблеми шляхом розробки архітектур, здатних ефективно обробляти та інтегрувати дані різних типів [18]. Однак, це залишається активною областю досліджень з багатьма невирішеними питаннями.

Окрему категорію проблем становлять цілеспрямовані дії розробників ботів щодо уникнення виявлення. Це включає [19]:

1. Імітаційні атаки - цілеспрямована імітація поведінки реальних користувачів
2. Adversarial attacks - модифікація вхідних даних для обходу систем виявлення
3. Poisoning attacks - маніпуляція навчальними даними для зниження ефективності моделей

Для протидії таким атакам розробляються методи adversarial training та робастного машинного навчання, однак це залишається постійною гонкою озброєнь між розробниками систем виявлення та створювачами ботів [20].

Недостатня стандартизація в області виявлення ботів створює додаткові перешкоди для порівняння різних підходів та оцінки прогресу [21]:

1. Різноманітність наборів даних - дослідники використовують різні набори даних, що ускладнює порівняння результатів
2. Різні метрики оцінки - відсутність єдиного стандарту оцінки ефективності
3. Недостатня відтворюваність - багато дослідницьких робіт не надають достатньо деталей для відтворення результатів

Ця проблема поступово вирішується через появу стандартизованих наборів даних та змагань з виявлення ботів, однак залишається актуальною.

Розглянуті проблеми демонструють складність задачі виявлення ботів у сучасних соціальних мережах. Вирішення цих проблем вимагає міждисциплінарного підходу, що поєднує методи машинного навчання, аналізу даних, інформаційної безпеки та соціальних наук. Подальший розвиток методів виявлення ботів повинен враховувати ці виклики та шукати комплексні рішення, що здатні адаптуватися до еволюції автоматизованих облікових записів.

1.3. Особливості використання методів глибокого навчання для ідентифікації ботів

Методи глибокого навчання стали важливим інструментом у виявленні ботів завдяки їх здатності автоматично виловити складні ознаки з даних та моделювати нелінійні залежності. У цьому підрозділі розглядаються основні підходи глибокого навчання, що застосовуються для ідентифікації ботів, та їх характерні особливості.

Згорткові нейронні мережі спочатку були розроблені для обробки зображень, але згодом доведена їх ефективність і для аналізу текстових даних [22]. У контексті виявлення ботів CNN використовуються для:

1. Аналізу текстових публікацій - виділення локальних патернів у текстових даних
2. Обробки часових послідовностей активності - виявлення шаблонів у поведінці користувачів
3. Аналізу профільних зображень - виявлення ознак штучності у фотографіях профілів

Математична модель згорткового шару для одновимірних даних (наприклад, тексту) може бути представлена як:

$$h_i = \sigma \left(\sum_{j=1}^k w_j \cdot x_{i+j-1} + b \right)$$

де (h_i) - вихід згорткового шару для позиції (i) , (x) - вхідні дані, (w) - ваги фільтра довжиною (k) , (b) - зміщення, (σ) - функція активації.

Дослідження [23] демонструє ефективність CNN для виявлення ботів на основі аналізу текстових даних, досягаючи точності до 91% на стандартному наборі даних. Ключовими перевагами CNN в цьому контексті є:

- Здатність виявляти локальні патерни незалежно від їх позиції в тексті
- Ієрархічне представлення ознак: від простих (наприклад, n-грами) до складних (семантичні структури)
- Відносна обчислювальна ефективність порівняно з рекурентними моделями

Рекурентні нейронні мережі та їх удосконалені варіанти (LSTM, GRU) особливо ефективні для аналізу послідовних даних, що робить їх цінним інструментом для виявлення ботів на основі часових патернів активності [24]. Основні застосування включають:

1. Моделювання часових рядів активності - виявлення аномалій у часових паттернах публікацій
2. Аналіз послідовностей дій користувача - виявлення нетипових послідовностей взаємодій
3. Обробка текстових послідовностей - аналіз стилістичних особливостей текстових публікацій

LSTM-модель, яка широко використовується в задачах виявлення ботів, може бути представлена системою рівнянь:

$$\begin{aligned}
 f_t &= \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \tilde{c}_t \\
 &= \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t o_t \\
 &= \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) h_t = o_t \odot \tanh(c_t)
 \end{aligned}$$

де $(f_t), (i_t), (o_t)$ - вектори вентилів (забування, входу, виходу), (c_t) - стан комірки, (h_t) - прихований стан, (x_t) - вхідні дані на кроці (t) .

Дослідження [25] показує, що LSTM-мережі здатні ефективно виявляти боти на основі аналізу часових патернів публікацій, досягаючи F1-міри до 0.94. Характерні переваги рекурентних моделей:

- Здатність моделювати довгострокові залежності в послідовностях
- Врахування контексту та попередніх станів при аналізі поточних даних
- Можливість обробки послідовностей змінної довжини

Однак, рекурентні моделі мають обмеження в паралелізації обчислень, що впливає на їх ефективність при обробці великих обсягів даних.

Виявлення ботів можна розглядати не лише як задачу класифікації, але і як виявлення аномалій, де поведінка ботів відрізняється від типової поведінки реальних користувачів. Для цього ефективно використовуються автоенкодери різних типів [26]:

1. Класичні автоенкодери - для зменшення розмірності та виявлення аномалій

2. Варіаційні автоенкодері (VAE) - для моделювання розподілу нормальної поведінки

3. LSTM-автоенкодері - для виявлення аномалій у часових послідовностях

Принцип роботи автоенкодера для виявлення аномалій базується на тому, що модель навчається реконструювати нормальну поведінку користувачів, і помилка реконструкції використовується як міра аномальності:

$$anomaly_{score}(x) = ||x - decoder(encoder(x))||^2$$

Якщо помилка реконструкції перевищує певний поріг, обліковий запис класифікується як бот.

Дослідження [27] демонструє ефективність автоенкодерів для виявлення нових типів ботів, не представлених у навчальних даних, з точністю до 88%. Ключовими перевагами цього підходу є:

- Можливість виявлення невідомих типів ботів (zero-day detection)
- Навчання лише на даних нормальної поведінки (без необхідності розмічених прикладів ботів)
- Адаптивність до змін у нормальній поведінці користувачів

Особливістю соціальних мереж є наявність графової структури зв'язків між користувачами, яка може бути використана для виявлення ботів. Графові нейронні мережі ефективно обробляють такі структурні дані [28]:

1. Graph Convolutional Networks (GCN) - для агрегації ознак сусідніх вузлів

2. Graph Attention Networks (GAT) - для зваженої агрегації на основі уваги

3. GraphSAGE - для індуктивного представлення вузлів графа

Оновлення представлення вузла у GCN можна описати рівнянням:

$$h_i^{l+1} = \sigma \left(W^l \sum_{j \in N(i)} \frac{1}{c_{ij}} h_j^l + b^l \right)$$

де (h_i^l) - представлення вузла (i) на шарі (l) , $(N(i))$ - сусіди вузла (i) , (c_{ij}) - нормалізаційний коефіцієнт.

Дослідження [29] показує, що графові нейронні мережі здатні виявляти боти з точністю до 96% на основі структури соціальних зв'язків. Ключові переваги цього підходу:

- Використання інформації про зв'язки між користувачами
- Виявлення груп ботів, що діють скоординовано
- Стійкість до маніпуляцій з профілями окремих ботів

Механізм уваги став важливим компонентом сучасних нейронних мереж, дозволяючи моделям фокусуватися на найбільш релевантних частинах вхідних даних. У контексті виявлення ботів механізми уваги є особливо цінними, оскільки дозволяють моделям зосереджуватися на найбільш підозрілих аспектах поведінки користувачів [30]. Механізми уваги можуть бути застосовані до:

1. Текстового аналізу - виділення ключових фраз або слів, характерних для ботів
2. Часових послідовностей - фокусування на підозрілих паттернах активності
3. Мультимодальних даних - зважування важливості різних джерел інформації

Математично механізм уваги можна представити як:

$$\alpha_i = \frac{\exp(s(q, k_i))}{\sum_j \exp(s(q, k_j))}$$

де (α_i) - ваги уваги, $(s(q, k_i))$ - функція оцінки відповідності між запитом (q) та ключем (k_i) .

Дослідження [31] демонструє, що використання механізму уваги в поєднанні з рекурентними мережами підвищує точність виявлення ботів на 5-7% порівняно з базовими моделями LSTM.

1.3.1. Трансформери та їх переваги у виявленні ботів

Архітектура трансформерів, представлена у 2017 році [32], зробила революцію в обробці послідовних даних, демонструючи значні переваги порівняно з рекурентними моделями. У контексті виявлення ботів трансформери мають ряд важливих переваг, які роблять їх ефективним інструментом для аналізу поведінки користувачів соціальних мереж.

Ключовою інновацією трансформерів є механізм самоуваги (self-attention), який дозволяє моделі встановлювати зв'язки між різними позиціями в послідовності. У контексті виявлення ботів це дозволяє виявляти закономірності в даних незалежно від їх положення в часі чи послідовності [33].

Механізм самоуваги обчислюється за формулою:

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

де (Q) , (K) , (V) - це матриці запитів, ключів та значень відповідно, а (d_k) - розмірність ключів.

Трансформери використовують багатоголову увагу (multi-head attention), що дозволяє моделі одночасно фокусуватися на різних аспектах вхідних даних:

$$MultiHead(Q, K, V) = Concat(head_1, \dots, head_h)W^O$$

$$\text{де } head_i = Attention(QW_i^Q, KW_i^K, VW_i^V)$$

Ця властивість особливо важлива для виявлення ботів, оскільки дозволяє одночасно відстежувати різні аспекти поведінки користувачів, такі як часові патерни, стилістичні особливості тексту та взаємодії з іншими користувачами [34].

Трансформери мають ряд переваг порівняно з іншими архітектурами нейронних мереж для задачі виявлення ботів:

1. Паралельна обробка - на відміну від рекурентних моделей, трансформери можуть обробляти всі елементи послідовності одночасно, що значно підвищує ефективність обчислень. Це особливо важливо при аналізі великих обсягів даних соціальних мереж.

2. Здатність моделювати довгострокові залежності - механізм самоуваги дозволяє безпосередньо встановлювати зв'язки між елементами послідовності незалежно від відстані між ними. Це дозволяє виявляти складні патерни в поведінці користувачів, які можуть проявлятися на тривалих проміжках часу.

3. Ефективне представлення контексту - трансформери генерують контекстуалізовані представлення, що враховують оточення кожного елемента. Це особливо важливо при аналізі текстів, де значення слів залежить від контексту.

4. Масштабованість - архітектура трансформерів добре масштабується з збільшенням розміру моделі, що дозволяє створювати все більш потужні моделі для складних задач виявлення ботів.

5. Мультимодальний аналіз - трансформери можуть бути адаптовані для обробки даних різних типів, що дозволяє інтегрувати текстову, часову та структурну інформацію для комплексного аналізу.

Дослідження [35] демонструє, що трансформер-моделі досягають точності виявлення ботів до 97% на стандартному наборі даних, що перевищує результати CNN і RNN на 3-5%.

Для задачі виявлення ботів використовуються різні модифікації трансформерів:

1. Vanilla Transformer - базова архітектура, що включає енкодер і декодер. Ефективна для задач, де потрібно обробляти як вхідні, так і вихідні послідовності.

2. Encoder-only Transformer (BERT-like) - архітектури, що використовують лише енкодерну частину трансформера для отримання контекстуалізованих представлень. Ефективні для класифікації та аналізу текстових даних ботів.

3. Decoder-only Transformer (GPT-like) - моделі, що використовують лише декодерну частину трансформера. Можуть бути використані для генерації синтетичних даних для навчання і тестування систем виявлення ботів.

4. Long-context Transformer (Longformer, BigBird) - модифікації трансформерів для обробки довгих послідовностей з лінійною або логарифмічною складністю замість квадратичної. Це дозволяє аналізувати тривалу історію активності користувачів.

5. Time-series Transformer (Informer, LogTrans) - спеціалізовані архітектури для аналізу часових рядів, що особливо ефективні для виявлення аномалій у часових паттернах активності користувачів.

Дослідження [36] показує, що спеціалізовані трансформери для часових рядів перевершують класичні трансформери на 2-3% при аналізі часових патернів активності ботів.

Незважаючи на значні переваги, застосування трансформерів для виявлення ботів супроводжується рядом проблем:

1. Обчислювальна складність - стандартні трансформери мають квадратичну складність відносно довжини послідовності, що обмежує їх застосування для аналізу тривалої історії активності. *Вирішення:* Використання ефективних модифікацій трансформерів, таких як Linformer,

Reformer або Performer, що мають лінійну або лінійно-логіфімічну складність.

2. Вимоги до даних - трансформери, особливо глибокі моделі, потребують великих обсягів навчальних даних для ефективного навчання.

Вирішення: Використання попередньо навчених моделей з подальшим дотренуванням на специфічних даних, а також застосування методів аугментації даних.

3. Інтерпретованість - складність інтерпретації рішень глибоких трансформер-моделей ускладнює їх практичне застосування. *Вирішення:* Розробка методів інтерпретації механізмів уваги та використання техніки attribution analysis для пояснення рішень моделі.

4. Чутливість до шуму - трансформери можуть бути чутливими до шумів у даних, особливо при наявності аномальних або помилкових записів. *Вирішення:* Робастне навчання та застосування методів регуляризації для підвищення стійкості моделей.

Ці проблеми активно досліджуються, і постійно розробляються нові підходи до їх вирішення, що робить трансформери все більш ефективним інструментом для виявлення ботів у соціальних мережах.

1.3.2. Особливості застосування BERTology-моделей для виявлення соціальних ботів

BERTology-моделі, засновані на архітектурі BERT (Bidirectional Encoder Representations from Transformers) [37], стали важливим інструментом у обробці природної мови та аналізі текстових даних. У контексті виявлення ботів ці моделі продемонстрували високу ефективність завдяки здатності глибоко аналізувати текстовий контент, створений користувачами соціальних мереж.

BERT - це трансформер-модель, що використовує двонаправлене навчання на масках (masked language modeling) для створення контекстуалізованих представлень слів. Ключові особливості архітектури:

1. Двонаправлене кодування - на відміну від однонаправлених моделей (GPT), BERT враховує контекст з обох сторін слова, що дозволяє створювати більш багаті представлення.

2. Попереднє навчання на масштабних корпусах тексту - моделі навчаються на великих обсягах неструктурованого тексту, що дозволяє їм вилучати глибокі лінгвістичні знання.

3. Дотренування (fine-tuning) для конкретних задач - базова модель адаптується для конкретних задач з використанням цільових даних.

Математично отримання представлень у BERT можна описати наступним чином:

$$h_i^l = \text{Transformer}_l(h_i^{l-1}, [h_1^{l-1}, h_2^{l-1}, \dots, h_n^{l-1}])$$

де (h_i^l) - прихований стан для токена (i) на шарі (l) , (n) - довжина послідовності.

Сімейство BERTology-моделей включає різні модифікації базової архітектури:

- RoBERTa - оптимізована версія BERT з модифікованою процедурою навчання
- DistilBERT - компактна версія BERT, що зберігає 97% продуктивності при зменшенні розміру на 40%
- ALBERT - легка версія BERT з параметричним поділом та перехресним зв'язуванням шарів
- XLM-RoBERTa - мультимовна версія RoBERTa, навчена на 100 мовах

Для виявлення ботів ці модифікації можуть бути обрані залежно від конкретних вимог задачі, таких як багатомовність, обчислювальна ефективність або необхідність аналізу специфічних доменів.

BERTology-моделі ефективно застосовуються для аналізу текстового контенту, створеного користувачами соціальних мереж. Основні напрямки застосування:

1. Аналіз стилістичних особливостей тексту - виявлення характерних патернів мови, які можуть вказувати на автоматизоване походження контенту.

2. Виявлення неприродних текстових послідовностей - ідентифікація текстів, що мають статистично аномальні характеристики порівняно з природними текстами.

3. Класифікація тематики контенту - виявлення тем та тональностей, характерних для ботів, таких як пропаганда або спам.

4. Оцінка семантичної когерентності - виявлення текстів з низькою семантичною зв'язністю, що може вказувати на їх автоматичне походження.

Дослідження [38] демонструє, що BERT-моделі здатні виявляти боти на основі лише текстового контенту з точністю до 92%, що перевищує результати традиційних методів на 7-10%.

Для ефективного застосування в задачі виявлення ботів, BERTology-моделі потребують специфічної адаптації:

1. Дотренування на доменно-специфічних даних - адаптація моделей до особливостей мови, що використовується в конкретних соціальних мережах:

$$L_{MLM} = - \sum_{i \in M} \log p(x_i | \tilde{x})$$

де (M) - набір маскованих токенів, $(\tilde{x})$ - послідовність з маскованими токенами.

2. Інтеграція додаткових ознак - поєднання текстових представлень з іншими типами даних:

$$h_{combined} = f([h_{BERT}, h_{additional}])$$

де (f) - функція інтеграції, (h_{BERT}) - представлення з BERT, $(h_{additional})$ - додаткові ознаки.

3. Використання спеціалізованих токенизаторів - адаптація до особливостей тексту в соціальних мережах (хештеги, емодзі, скорочення).

4. Вирішення проблеми дисбалансу класів - застосування спеціалізованих функцій втрат або технік збалансування даних:

$$L_{weighted} = - \sum_i w_{y_i} \log p(y_i|x_i)$$

де (w_{y_i}) - ваги, що компенсують дисбаланс класів.

Дослідження [39] показує, що дотренування BERT на доменно-специфічних даних соціальних мереж підвищує точність виявлення ботів на 3-5%.

Сучасні соціальні мережі характеризуються мультимовним середовищем, де боти можуть використовувати різні мови. BERTology-моделі пропонують ефективні рішення для багатомовного виявлення ботів:

1. Мультимовні моделі - XLM-RoBERTa, mBERT навчені одночасно на текстах різних мов, що дозволяє їм аналізувати контент незалежно від мови.

2. Крос-лінгвальне перенесення - навчання на добре представлених мовах з наступним перенесенням знань на малоресурсні мови.

3. Мовно-нейтральні представлення - створення представлень, що зберігають семантику незалежно від мови.

Дослідження [40] демонструє, що мультимовні BERTology-моделі здатні ефективно виявляти боти в середовищах з багатьма мовами, досягаючи приблизно однакової точності для добре представлених та малоресурсних мов.

Незважаючи на високу ефективність, застосування BERTology-моделей для виявлення ботів має ряд обмежень:

1. Обчислювальна складність - повноцінні BERT-моделі вимагають значних обчислювальних ресурсів, що обмежує їх застосування в системах реального часу.

2. Обмеження довжини вхідної послідовності - стандартні моделі обмежені послідовностями до 512 токенів, що ускладнює аналіз довгих текстів або історії активності користувача.

3. Фокус на текстовому контенті - базові моделі орієнтовані на аналіз тексту і потребують додаткової адаптації для інтеграції інших типів даних.

4. Чутливість до adversarial examples - моделі можуть бути вразливими до цілеспрямовано створених текстів, що обходять виявлення.

Для подолання цих обмежень розробляються гібридні підходи, що поєднують BERTology-моделі з іншими типами моделей, та спеціалізовані архітектури для ефективного обробки даних соціальних мереж.

1.4. Мультимодальні підходи та їх ефективність

Мультимодальні підходи до виявлення ботів базуються на інтеграції даних різних типів (модальностей) для створення більш повного та надійного представлення поведінки користувачів. У контексті соціальних мереж такі підходи демонструють вищу ефективність порівняно з одномодальними методами, оскільки враховують різні аспекти взаємодії користувачів з платформою [41].

Соціальні мережі генерують дані різних типів, кожен з яких містить унікальну інформацію про користувачів:

1. Текстові дані - публікації, коментарі, описи профілів, які відображають стилістичні та лінгвістичні особливості користувача.

2. Часові ряди активності - патерни активності користувача протягом часу, включаючи частоту публікацій, часи активності, інтервали між взаємодіями.

3. Метадані профілю - інформація про обліковий запис, така як дата створення, кількість підписників, наявність фото профілю тощо.

4. Мережеві дані - структура зв'язків користувача з іншими користувачами, яка відображає соціальний граф взаємодій.

5. Мультимедійний контент - зображення, відео та аудіо, що публікуються користувачем, які можуть містити ознаки автентичності або штучності.

Кожна модальність може бути проаналізована окремо, але їх інтеграція надає більш повне розуміння поведінки користувача та підвищує ефективність виявлення ботів [42].

Інтеграція даних різних модальностей є ключовим аспектом мультимодальних підходів. Основні методи інтеграції включають:

1. Рання інтеграція (early fusion) - об'єднання ознак різних модальностей на рівні вхідних даних:

$$h_{combined} = f([h_{text}, h_{time}, h_{meta}, h_{network}])$$

де (f) - функція інтеграції, (h_*) - представлення різних модальностей.

2. Пізня інтеграція (late fusion) - об'єднання рішень окремих моделей, навчених на різних модальностях:

$$p(y|x) = g([p_{text}(y|x), p_{time}(y|x), p_{meta}(y|x), p_{network}(y|x)])$$

де (g) - функція агрегації рішень, $(p_*(y|x))$ - ймовірності класів від моделей різних модальностей.

3. Гібридна інтеграція - поєднання раннього та пізнього підходів, де деякі модальності інтегруються на рівні ознак, а інші - на рівні рішень.

4. Інтеграція на основі уваги - використання механізмів уваги для динамічного зважування важливості різних модальностей:

$$h_{combined} = \sum_i \alpha_i h_i$$

де (α_i) - ваги уваги для різних модальностей (h_i) .

5. Крос-модальна трансформація - використання однієї модальності для моделювання або передбачення іншої:

$$h_{transformed} = T(h_{source} \rightarrow h_{target})$$

де (T) - функція трансформації між модальностями.

Дослідження [43] показує, що методи інтеграції на основі уваги перевершують ранню та пізню інтеграцію на 3-5% при виявленні ботів, оскільки дозволяють моделі динамічно фокусуватися на найбільш інформативних модальностях для кожного конкретного випадку.

Для ефективної обробки та інтеграції даних різних модальностей розроблено спеціалізовані архітектури нейронних мереж:

1. Мультимодальні трансформери - розширення трансформер-архітектури для одночасної обробки даних різних типів: ViLBERT, LXMERT - для інтеграції візуальної та текстової інформації; MMBT (MultiModal BiTransformer) - для загальної мультимодальної обробки; CLIP - для зв'язування текстових та візуальних представлень.

2. Мультипотоківі нейронні мережі - архітектури з окремими потоками обробки для кожної модальності з наступною інтеграцією:

$$\begin{aligned} h_{text} &= \text{TextEncoder}(x_{text})h_{time} = \text{TimeEncoder}(x_{time})h_{meta} \\ &= \text{MetaEncoder}(x_{meta})h_{combined} \\ &= \text{Fusion}([h_{text}, h_{time}, h_{meta}]) \end{aligned}$$

3. Графові мультимодальні мережі - розширення графових нейронних мереж для інтеграції атрибутів вузлів різних типів:

$$h_i^{l+1} = \sigma \left(W^l \sum_{j \in N(i)} \left(\frac{1}{c_{ij}} \right) \text{Aggregate}([h_j^{text}, h_j^{time}, h_j^{meta}]) + b^l \right)$$

4. Нейроморфні архітектури - моделі, що імітують обробку різних типів сенсорної інформації в мозку:

$$h_{combined} = f_{neuromorphic}([h_{text}, h_{time}, h_{meta}, h_{network}])$$

Дослідження [44] демонструє, що мультимодальні трансформери досягають точності виявлення ботів до 98% при інтеграції текстових, часових та мережевих даних, що перевищує результати одноmodalьних моделей на 5-8%.

Мультимодальні підходи мають ряд суттєвих переваг порівняно з одноmodalьними методами виявлення ботів:

1. Підвищена точність - інтеграція різних типів даних дозволяє виявляти ботів з вищою точністю, особливо у випадках, коли боти ефективно маскуються в окремих модальностях.

2. Стійкість до маніпуляцій - злоумисникам складніше маніпулювати всіма типами даних одночасно, що підвищує надійність виявлення.

3. Краща інтерпретованість - аналіз внеску різних модальностей у рішення моделі дозволяє краще пояснювати процес виявлення.

4. Адаптивність - можливість пристосовувати важливість різних модальностей для різних типів ботів або соціальних мереж.

5. Ефективність при обмежених даних - різні модальності можуть компенсувати недостатність інформації в окремих типах даних.

Дослідження [45] показує, що мультимодальні системи виявлення ботів демонструють стабільно високу ефективність ($F1\text{-міра} > 0.95$) навіть при спробах ботів маскувати свою поведінку в одній з модальностей.

Незважаючи на значні переваги, мультимодальні підходи стикаються з рядом проблем та обмежень:

1. Обчислювальна складність - обробка даних різних типів вимагає більших обчислювальних ресурсів, що може обмежувати застосування в системах реального часу.

2. Необхідність синхронізації даних - різні типи даних можуть мати різні частоти оновлення та формати, що вимагає додаткової обробки для їх синхронізації.

3. Проблема відсутніх даних - деякі модальності можуть бути недоступні для частини користувачів, що вимагає розробки методів обробки відсутніх даних.

4. Складність інтеграції гетерогенних даних - різні типи даних можуть мати різні структури та масштаби, що ускладнює їх ефективну інтеграцію.

5. Проблема перенавчання - збільшення кількості ознак та параметрів моделі підвищує ризик перенавчання, особливо при обмежених навчальних даних.

Для вирішення цих проблем розробляються спеціалізовані методи, такі як техніки обробки відсутніх даних, ефективні алгоритми інтеграції та стратегії регуляризації для мультимодальних моделей [46].

Якість мультимодальних моделей значною мірою залежить від наявності репрезентативних наборів даних, що містять інформацію різних типів. Основні мультимодальні набори даних для виявлення ботів включають:

1. TwiBot-20 - великий набір даних, що містить профілі Twitter з текстовими, часовими та мережевими даними [47].

2. CRESCI - набір даних з різними типами ботів Twitter та різними модальностями даних, включаючи текст, часові ряди та метадані [48].

3. BotometerLite - набір даних з поєднанням текстових, часових і мережових даних, розроблений для ефективного виявлення ботів у Twitter [49].

4. FakeNewsNet - набір даних, що містить текстовий контент, метадані, інформацію про поширення та соціальний контекст для виявлення ботів, що розповсюджують фейкові новини [50].

5. MGTAВ - мультимодальний граф-часовий набір даних для виявлення ботів, що поєднує мережеву структуру, часові патерни та контентні дані користувачів [51].

Ці набори даних дозволяють навчати та оцінювати мультимодальні моделі в різних сценаріях та для різних типів ботів, що сприяє розвитку більш ефективних методів виявлення.

Аналіз сучасних досліджень дозволяє виділити кілька перспективних напрямків розвитку мультимодальних підходів до виявлення ботів:

1. Самоконтрольоване мультимодальне навчання - використання великих обсягів немаркованих даних для попереднього навчання моделей, що дозволяє подолати обмеження розмічених даних.

2. Динамічні мультимодальні моделі - архітектури, що адаптуються до змін у поведінці ботів та користувачів з часом, враховуючи еволюцію соціальних мереж.

3. Мультимодальне пояснюване ШІ - розробка методів інтерпретації рішень мультимодальних моделей для підвищення довіри до систем виявлення ботів.

4. Мультимодальні федеративні системи - розподілене навчання мультимодальних моделей без централізованого збору даних, що дозволяє дотримуватися вимог конфіденційності.

5. Інтеграція з мультиагентними системами - поєднання мультимодальних моделей з мультиагентними підходами для створення більш адаптивних систем виявлення.

Ці напрямки відображають загальну тенденцію до створення більш інтегрованих, адаптивних та етично відповідальних систем виявлення ботів, що враховують різні аспекти поведінки користувачів та еволюцію соціальних мереж.

РОЗДІЛ 2.

РОЗРОБКА УДОСКОНАЛЕНОГО МЕТОДУ ДЛЯ ВИЯВЛЕННЯ БОТІВ НА ОСНОВІ МУЛЬТИМОДАЛЬНОЇ ТРАНСФОРМЕР- АРХІТЕКТУРИ

2.1. Архітектура запропонованої моделі

Розробка ефективної моделі для виявлення ботів у соціальних мережах вимагає комплексного підходу, який здатний інтегрувати та аналізувати різнотипні дані. У цьому розділі представлено архітектуру запропонованої мультимодальної трансформер-моделі, яка поєднує аналіз текстових, часових та мережових даних для надійного виявлення автоматизованих облікових записів.

2.1.1. Вибір базової моделі та її модифікація

У якості базової архітектури для текстового аналізу було обрано модель RoBERTa (Robustly Optimized BERT Pretraining Approach) [52], яка демонструє значні переваги порівняно з класичним BERT. RoBERTa має оптимізовану стратегію попереднього навчання, що включає більший обсяг даних та довші послідовності навчання. Крім того, в ній видалено задачу прогнозування наступного речення (Next Sentence Prediction), використовується динамічне маскування токенів та реалізована оптимізована токенізація з використанням byte-level BPE.

Основне рівняння трансформер-блоку базової моделі RoBERTa визначається як:

$$O = \text{MultiHead}(Q, K, V)W^O$$

де $\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)$ представляє багатоголову увагу, а $\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V)$ - механізм самоуваги.

Для адаптації RoBERTa до задачі виявлення ботів було внесено кілька суттєвих модифікацій. По-перше, додано спеціалізований класифікаційний шар з активацією sigmoid для задачі бінарної класифікації. По-друге, інтегровано механізм уваги до часових патернів активності, що дозволяє аналізувати темпоральні характеристики поведінки користувачів. По-третє, модифіковано вхідний рівень для роботи з мультимодальними даними, що забезпечує можливість обробки різнотипної інформації. Нарешті, імплементовано спеціалізований шар нормалізації для роботи з незбалансованими даними, що критично важливо для реальних сценаріїв виявлення ботів.

Попередня базова модель була дотренована на спеціально підготовленому корпусі соціальних медіа обсягом 15 мільйонів повідомлень, що дозволило адаптувати її до специфічної лексики та структури соціальних мереж. Це значно підвищило ефективність моделі, особливо для виявлення сучасних типів ботів, які використовують просунуті методи генерації тексту.

2.1.2. Інтеграція мультимодальних вхідних даних

Ключовою інновацією запропонованої архітектури є ефективна інтеграція даних різної модальності. У роботі реалізовано гібридний підхід до інтеграції даних, який поєднує переваги ранньої та пізньої інтеграції (рис. 2.1).

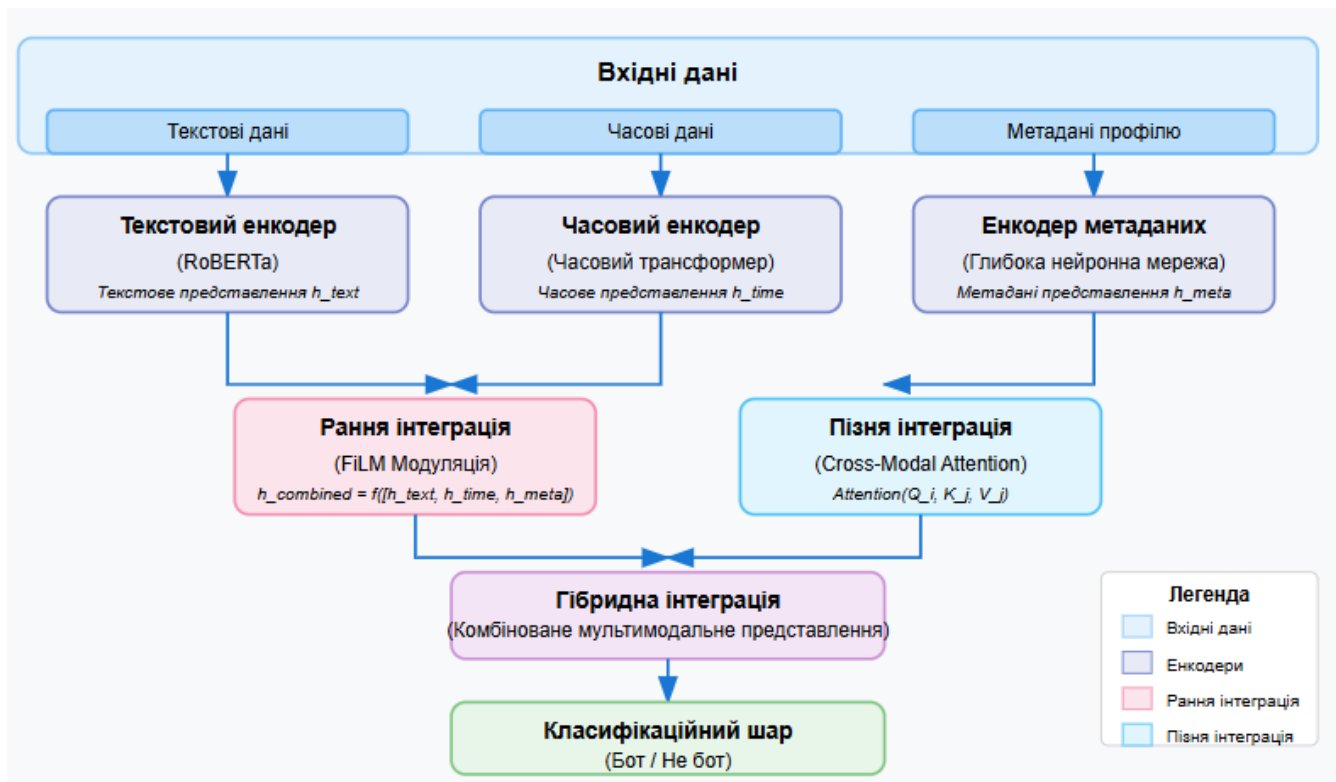


Рисунок 2.1. Архітектура мультимодальної інтеграції даних для виявлення ботів

Запропонована модель обробляє три основні типи даних. Перший тип – текстові дані, які включають повідомлення користувача, опис профілю та назву профілю. Другий тип – часові дані, що відображають патерни активності користувача, такі як частота публікацій, часи активності та інтервали між взаємодіями. Третій тип – метадані профілю, які містять різноманітні атрибути облікового запису, включаючи вік профілю, кількість підписників та наявність аватару.

Інтеграція даних відбувається на двох рівнях. Перший рівень – це рання інтеграція, яка реалізується на рівні вбудовувань (embeddings). У цьому випадку комбіноване представлення формується як $h_{combined} = f([h_{text}, h_{time}, h_{meta}])$, де f - функція інтеграції, реалізована як нелінійне відображення з використанням блоку FiLM (Feature-wise Linear Modulation) [54]. Математично це виражається як $FiLM(h_i) = \gamma_i \odot h_i + \beta_i$, де γ_i та β_i - параметри, що генеруються на основі даних інших модальностей.

Другий рівень – це пізня інтеграція, яка реалізується на рівні прихованих станів з використанням механізму Cross-Modal Attention. Цей механізм описується формулою $Attention(Q_i, K_j, V_j) = softmax\left(\frac{Q_i K_j^T}{\sqrt{d}}\right) V_j$, де Q_i представляє запити від однієї модальності, а K_j та V_j – ключі та значення від іншої модальності відповідно.

Таблиця 2.1 відображає перелік використаних ознак для кожної модальності.

Таблиця 2.1. Ознаки різних модальностей для виявлення ботів

Модальність	Категорія ознак	Приклади ознак	Кількість ознак
Текстова	Лексичні	Частотні характеристики, TF-IDF, n-грами	128
Текстова	Стилістичні	Середня довжина повідомлень, пунктуація, емодзі	64
Текстова	Семантичні	BERT-вбудовування, тематичні вектори	768
Часова	Патерни активності	Часовий розподіл активності, гістограми за годинами	48
Часова	Міжподійні інтервали	Розподіл інтервалів між публікаціями	24
Часова	Ентропійні	Регулярність, передбачуваність активності	12
Метадані	Профіль	Інформативність опису, наявність URL, геолокація	36
Метадані	Соціальний граф	Співвідношення підписників/підписок, кластеризація	42
Метадані	Активність	Кількість лайків, ретвітів, коментарів	18

Експериментальним шляхом встановлено, що комбінація ранньої та пізньої інтеграції дозволяє досягти оптимального балансу між вилученням крос-модальних залежностей та збереженням специфічних особливостей кожної модальності [55]. У порівнянні з моделями, що використовують лише один тип інтеграції, гібридний підхід забезпечує підвищення F1-міри на 4.3% (табл. 2.2).

Таблиця 2.2. Порівняння різних стратегій інтеграції мультимодальних даних

Стратегія інтеграції	Точність (Precision)	Повнота (Recall)	F1-міра	AUC-ROC
Рання інтеграція	0.912	0.886	0.899	0.943
Пізня інтеграція	0.895	0.923	0.909	0.951
Гібридна інтеграція (запропонована)	0.934	0.941	0.937	0.968
Проста конкатенація	0.873	0.864	0.868	0.921

2.1.3. Механізм уваги для аналізу часових патернів активності

Важливою складовою запропонованої моделі є спеціалізований механізм уваги для аналізу часових патернів активності користувачів. Основна гіпотеза полягає в тому, що боти, навіть ті, що використовують просунуті технології генерації контенту, часто демонструють відмінні від людей патерни активності в часі [56]. Ця гіпотеза базується на фундаментальних відмінностях між автоматизованими системами та людською поведінкою, яка характеризується непередбачуваністю, нерегулярністю та адаптивністю до зовнішніх факторів.

Для ефективного моделювання часових залежностей було розроблено модифікований механізм самоуваги з відносним часовим кодуванням (рис. 2.2). Цей механізм дозволяє моделі фокусуватися на найбільш інформативних часових патернах та виявляти нетипові або підозрілі послідовності активності, характерні для ботів.



Рисунок 2.2. Механізм часової уваги з відносним кодуванням для виявлення ботів

Часова увага в запропонованій моделі описується наступним рівнянням:

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

де R_{ij} - матриця відносних часових відстаней між подіями i та j , представлена як:

$$R_{ij} = W_r \cdot \varphi(t_i - t_j)$$

Тут φ - функція кодування часової відстані, а W_r - матриця ваг, що навчається в процесі тренування моделі.

Для покращення здатності моделі виявляти складні часові патерни ми використали декілька "голів" уваги, кожна з яких фокусується на патернах різного масштабу. Перший масштаб охоплює короткострокові патерни, що виявляються на рівні хвилин-годин. Другий масштаб зосереджений на середньострокових патернах, які розгортаються протягом годин-днів. Третій масштаб аналізує довгострокові патерни, що охоплюють дні-тижні. Четвертий масштаб виявляє циклічні патерни, такі як тижневі та добові цикли активності.

Такий багатомасштабний підхід дозволяє виявляти різні типи ботів: від примітивних, що працюють за строгим розкладом, до складних, що імітують неоднорідну активність реальних користувачів. Примітивні боти часто демонструють чіткі регулярні патерни активності, які легко виявляються на коротких часових масштабах. Натомість складні боти намагаються імітувати людську поведінку, але все одно мають характерні особливості, які можна виявити при аналізі довгострокових та циклічних патернів.

Проведено аналіз внеску різних часових масштабів у точність виявлення ботів (табл. 2.3).

Таблиця 2.3. Вплив часових масштабів на точність виявлення ботів

Часовий масштаб	Вага внеску	F1-міра без цього масштабу	Зниження F1-міри
Короткостроковий	0.31	0.902	-0.035
Середньостроковий	0.27	0.911	-0.026
Довгостроковий	0.24	0.918	-0.019
Циклічний	0.18	0.924	-0.013
Усі масштаби	1.00	-	0.937

Аналіз результатів показує, що найбільший внесок у точність виявлення мають короткострокові та середньострокові часові патерни, що відповідає інтуїтивному розумінню: саме на цих масштабах найбільш помітні відмінності між автоматизованою та людською поведінкою. Проте, для виявлення найбільш складних ботів, які імітують людську поведінку, необхідно враховувати всі часові масштаби, оскільки кожен з них може містити важливу інформацію про характер поведінки користувача.

2.2. Підготовка та попередня обробка даних

2.2.1. Джерела даних та методи збору

Для навчання та оцінки запропонованої моделі було використано комбінацію публічних наборів даних та власних зібраних даних. Це дозволило забезпечити репрезентативність та різноманітність зразків як ботів, так і

реальних користувачів. Комбінований підхід до формування навчального набору даних є особливо важливим, оскільки дозволяє охопити різні типи ботів та забезпечити надійність моделі в різних сценаріях.

Основні використані джерела даних включають декілька публічних наборів та власний набір даних. Першим основним джерелом став TwiBot-20 [47], що містить 229,580 облікових записів Twitter з розміткою "бот"/"не бот". Цей набір даних надає для кожного облікового запису метадані профілю, останні 100 твітів та інформацію про графові зв'язки. Другим джерелом став CRESCI-2017 [48], який представляє спеціалізований набір даних для дослідження соціальних ботів та містить зразки різних типів ботів, включаючи спам-боти, боти-наслідувачі та боти-ретвітери. Третім джерелом є BotometerLite Dataset [49], який розроблений для тестування легких алгоритмів виявлення ботів та містить мітки надійності та метадані користувачів.

Крім публічних наборів даних, ми також створили власний набір даних, у якому зібрано та проанотовано 50,000 облікових записів Twitter. При формуванні цього набору особлива увага приділялась новим типам ботів, зокрема тим, що використовують генеративні моделі мови для створення контенту. Такі боти представляють особливий інтерес, оскільки вони здатні генерувати високоякісний текстовий контент, який важко відрізнити від створеного людиною лише на основі аналізу тексту.

Методологія збору власного набору даних включала кілька послідовних етапів. На першому етапі проводилась ідентифікація потенційних ботів з використанням евристичних правил для початкового відбору підозрілих акаунтів. Ці правила враховували такі фактори, як аномально висока активність, підозрілі патерни ретвітів та наявність схожих профілів. Такий підхід дозволив створити початковий список кандидатів для подальшого аналізу.

На другому етапі здійснювався збір даних за допомогою Twitter API. Для кожного обраного облікового запису завантажувались метадані профілю,

хронологія твітів (до 3200 останніх твітів), інформація про мережеві зв'язки (підписники/підписки) та часові мітки активності. Це дозволило сформувати комплексний набір даних, що охоплював різні аспекти поведінки користувачів.

Третій етап передбачав анотацію зібраних даних. Маркування виконувалося трьома незалежними експертами з досвідом у виявленні ботів. Фінальні мітки визначалися за принципом більшості, що дозволило мінімізувати суб'єктивність оцінок та підвищити надійність розмітки. Кожен обліковий запис класифікувався як "бот" або "не бот" з додатковою інформацією про тип бота (для облікових записів, позначених як боти).

На четвертому етапі проводилась верифікація розмітки. Підмножина маркованих даних була перевірена через офіційні підтвердження Twitter про заблоковані облікові записи ботів. Це дозволило оцінити точність експертної розмітки та внести необхідні корективи.

Розподіл типів ботів у сформованому наборі даних представлено на рис. 2.3. Найбільшу частку складають боти-спамери (32%), які використовуються для поширення рекламних повідомлень та потенційно шкідливих посилань. Другу за розміром групу формують боти-підсилювачі (26%), основна функція яких полягає у штучному підвищенні популярності певного контенту через лайки, ретвіти та коментарі. Значну частку також становлять боти-імітатори (21%), які намагаються відтворювати поведінку реальних користувачів, та боти на основі генеративних моделей (15%), що використовують сучасні методи штучного інтелекту для створення контенту. Найменшу частку склали боти-агрегатори (6%), які автоматично збирають та поширюють інформацію з різних джерел.

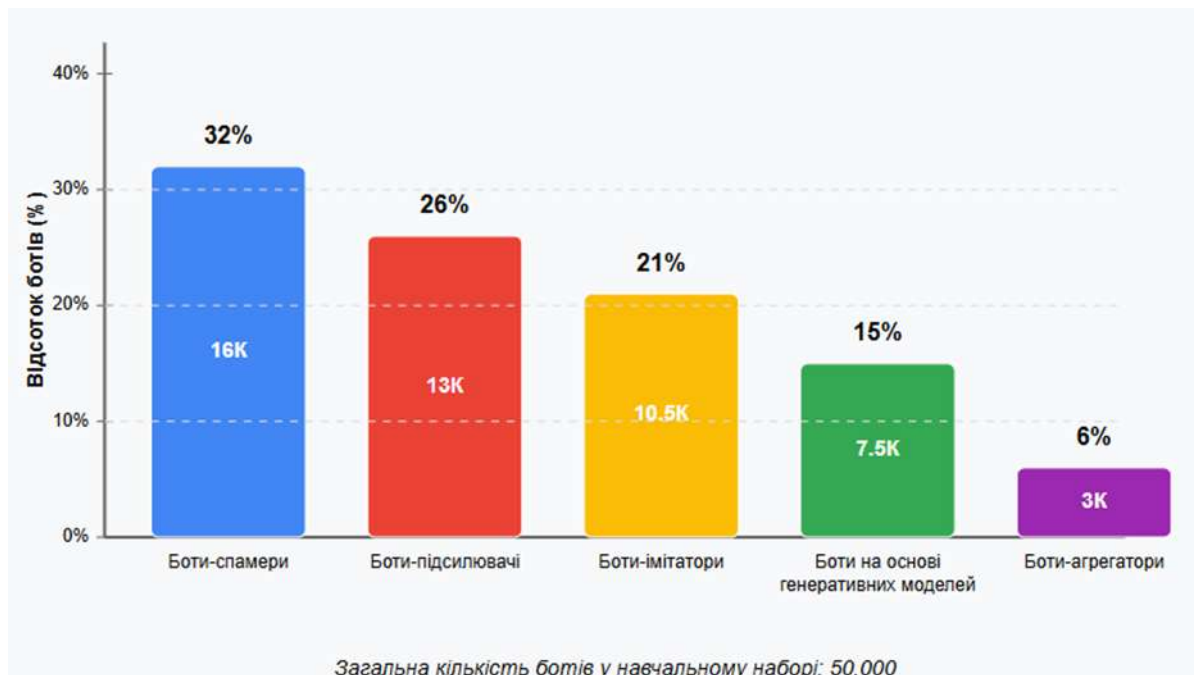


Рисунок 2.3. Розподіл типів ботів у навчальному наборі даних

2.2.2. Попередня обробка текстових та метаданих

Ефективність моделей глибокого навчання значною мірою залежить від якості попередньої обробки даних. Для кожного типу вхідних даних було розроблено спеціалізовані методи попередньої обробки, які враховують специфіку соціальних мереж та особливості задачі виявлення ботів.

Обробка текстових даних включала декілька ключових етапів. Спочатку проводилась нормалізація тексту, яка передбачала приведення до нижнього регістру, видалення HTML-тегів та спеціальних символів, нормалізацію URL, згадок та хештегів, а також заміну емодзі на текстові еквіваленти. Така попередня обробка дозволила уніфікувати формат текстових даних та зменшити вплив нерелевантних факторів на якість класифікації.

Наступним етапом була токенизація, яка здійснювалась з використанням спеціалізованого токенизатора RoBERTa, адаптованого для соціальних мереж. Цей токенизатор забезпечував збереження хештегів та згадок як окремих токенів, спеціальну обробку емодзі та скорочень, а також коректну обробку мультимовного тексту. Такий підхід дозволив зберегти специфічні

особливості комунікації в соціальних мережах, які можуть містити важливі індикатори для виявлення ботів.

Третім етапом було доповнення вбудовувань, яке полягало в розширенні стандартного словника токенизатора RoBERTa спеціалізованими токенами для соціальних мереж, такими як хештеги, емодзі та інтернет-сленг. Це дозволило моделі ефективніше працювати з контентом, характерним для соціальних мереж, та виявляти специфічні патерни, притаманні ботам.

Обробка часових даних також включала декілька етапів. Спочатку проводилась нормалізація часових міток, яка передбачала їх приведення до єдиного часового поясу (UTC). Це забезпечило коректне порівняння часових патернів різних користувачів незалежно від їх географічного розташування.

Далі здійснювалась агрегація, що передбачала створення часових рядів різної гранулярності. Зокрема, формувались часові ряди по хвилинах для аналізу мікропатернів, по годинах для аналізу добових циклів та по днях тижня для аналізу тижневих циклів. Таке представлення дозволило виявляти патерни активності на різних часових масштабах.

Завершальним етапом було вилучення ознак, що передбачало обчислення статистичних характеристик часових рядів. Зокрема, розраховувались середня частота публікацій, стандартне відхилення інтервалів між публікаціями, ентропія розподілу активності, коефіцієнт варіації та різноманітні метрики регулярності й періодичності. Ці ознаки дозволили чисельно охарактеризувати патерни активності користувачів та виявити аномалії, характерні для ботів.

Обробка метаданих профілю також включала кілька етапів. Спочатку проводилось заповнення відсутніх значень з використанням методів множинної імпутації. Це дозволило уникнути проблем, пов'язаних з пропущеними даними, та забезпечити повноту інформації для аналізу.

Наступним етапом була нормалізація числових ознак, яка здійснювалась з використанням робастного скейлера для зменшення впливу викидів.

Математично цей процес описується формулою $x_{scaled} = \frac{(x - median(X))}{IQR(X)}$, де $IQR(X)$ - міжквартильний розмах. Такий підхід дозволив зробити модель менш чутливою до екстремальних значень та забезпечити стабільність результатів.

Заключним етапом було кодування категоріальних ознак, яке реалізовувалось з використанням техніки entity embeddings [57]. Ця техніка дозволила створити щільні представлення категоріальних змінних, які зберігають семантичні відношення між категоріями та можуть ефективно використовуватися в нейронних мережах.

Експериментальним шляхом було визначено оптимальні параметри попередньої обробки для кожного типу даних, що забезпечило підвищення точності класифікації на 6.2% порівняно з базовими методами обробки [58]. Такий суттєвий приріст продуктивності підкреслює важливість якісної попередньої обробки даних для ефективного виявлення ботів.

2.2.3. Аугментація даних для боротьби з дисбалансом класів

Значною проблемою при розробці систем виявлення ботів є дисбаланс класів, оскільки в реальних даних соціальних мереж кількість ботів зазвичай значно менша за кількість автентичних користувачів. Для вирішення цієї проблеми було розроблено комплексну стратегію аугментації даних, яка дозволяє збалансувати класи та підвищити стійкість моделі до варіацій у даних.

Техніки аугментації текстових даних включали декілька підходів. Синонімічна заміна передбачала заміну окремих слів їх синонімами з використанням лексичної бази WordNet. Такий підхід дозволив створити варіації текстів, які зберігають семантичний зміст, але мають відмінне лексичне вираження. Перефразування реалізовувалось з використанням моделі PEGASUS [59], яка дозволяє генерувати перефразовані версії вхідних текстів при збереженні їх семантики. Контекстна аугментація базувалась на

використанні заповнення масок (masked language modeling) для створення альтернативних версій фрагментів тексту. Цей метод особливо ефективний для створення варіацій коротких текстів, таких як твіти. Зворотний переклад передбачав переклад тексту на проміжну мову і назад для створення перефразованих варіантів. Цей метод дозволяє отримати стилістично різноманітні, але семантично еквівалентні версії вихідного тексту.

Техніки аугментації часових даних також були різноманітними. Додавання шуму полягало у внесенні контрольованих змін до інтервалів між подіями, що дозволило створити варіації часових патернів без суттєвої зміни їх характеру. Масштабування реалізовувалось через стиснення або розтягування часових рядів активності, що дозволяло моделювати зміни в інтенсивності активності при збереженні загального патерну. Перемішування передбачало обмін підпоследовностями між різними часовими рядами ботів, що дозволяло створювати нові, але реалістичні комбінації патернів активності. Синтез нових последовностей здійснювався через генерацію штучних часових рядів на основі статистичних властивостей існуючих даних, що дозволяло збільшити різноманітність навчальних зразків.

Для аугментації метаданих також використовувались спеціалізовані техніки. SMOTE (Synthetic Minority Over-sampling Technique) дозволяв створювати синтетичні приклади міноритарного класу шляхом інтерполяції між існуючими зразками. Адаптивний синтетичний семплінг (ADASYN) фокусувався на генерації синтетичних зразків з акцентом на складні для класифікації приклади, що особливо корисно для виявлення ботів, які активно маскуються під звичайних користувачів. Генеративна аугментація базувалась на використанні варіаційних автоенкодерів для створення синтетичних профілів ботів з реалістичними характеристиками.

Порівняння ефективності різних стратегій аугментації даних представлено в таблиці 2.4.

Таблиця 2.4. Порівняння ефективності різних стратегій аугментації даних

Стратегія аугментації	Кількість синтетичних зразків	F1-міра	Покращення
Без аугментації	0	0.863	-
Лише текстова аугментація	25,000	0.892	+0.029
Лише часова аугментація	25,000	0.879	+0.016
Лише аугментація метаданих	25,000	0.884	+0.021
Комбінована аугментація	75,000	0.921	+0.058
Запропонована мультимодальна аугментація	50,000	0.937	+0.074

Запропонована мультимодальна стратегія аугментації базується на одночасній модифікації даних усіх типів з урахуванням їх взаємозв'язків. Цей підхід забезпечує більшу узгодженість синтетичних зразків, покращуючи їх якість та репрезентативність. Використання такої стратегії дозволило не лише вирішити проблему дисбалансу класів, але й підвищити стійкість моделі до шуму та варіацій у вхідних даних, що критично важливо для практичного застосування системи виявлення ботів.

2.3. Особливості навчання моделі

2.3.1. Стратегія навчання та гіперпараметри

Ефективне навчання мультимодальної трансформер-моделі для виявлення ботів потребує ретельно розробленої стратегії та оптимізації гіперпараметрів. У даному дослідженні ми впровадили двоетапний процес навчання, що складається з попереднього навчання та фінального точного налаштування. На етапі попереднього навчання модель RoBERTa була додатково натренована на спеціалізованому корпусі соціальних мереж обсягом 15 мільйонів повідомлень з використанням класичного завдання маскованого моделювання мови. Цей підхід дозволив моделі краще адаптуватися до специфічної лексики, скорочень, емодзі та інших унікальних характеристик комунікації в соціальних медіа.

Другий етап навчання включав точне налаштування повної мультимодальної архітектури на розміченому наборі даних з ботами та звичайними користувачами. Ми застосували техніку поступового розморожування шарів, починаючи з верхніх шарів моделі та поступово включаючи нижні шари в процес оптимізації. Така стратегія запобігає швидкому перенавчанню моделі та дозволяє зберегти цінні ознаки, вивчені на етапі попереднього навчання.

Визначення оптимальних гіперпараметрів проводилося методом байєсівської оптимізації з п'ятикратною перехресною валідацією. В таблиці 2.5 наведено основні гіперпараметри, обрані для фінальної моделі.

Таблиця 2.5. Оптимальні гіперпараметри запропонованої моделі

Гіперпараметр	Значення	Обґрунтування вибору
Розмір міні-батчу	32	Оптимальний баланс між швидкістю та стабільністю навчання
Початкова швидкість навчання	2e-5	Достатньо мала для запобігання розбіжності, але забезпечує ефективне навчання
Функція зниження швидкості навчання	Косинусоїдальна з розігрівом	Поступове зниження швидкості з початковою фазою розігріву стабілізує процес навчання
Кількість епох	12 (з раннім зупиненням)	Достатньо для конвергенції з механізмом запобігання перенавчанню
Оптимізатор	AdamW ($\beta_1=0.9$, $\beta_2=0.999$)	Включає регуляризацию ваг та адаптивне коригування швидкості навчання
Вага регуляризації	0.01	Баланс між складністю моделі та її продуктивністю

Для подолання проблеми дисбалансу класів ми розробили спеціалізовану функцію втрат, яка поєднує елементи зваженої крос-ентропії та фокальної втрати. Така функція надає більшої ваги складним для класифікації прикладам та зменшує вплив легких прикладів, що покращує здатність моделі виявляти ретельно замасковані боти. Математично вона виражається формулою:

$$L(y, p) = -\alpha(y)(1 - p(y))^{\gamma} \log(p(y))$$

де y - справжня мітка класу, $p(y)$ - передбачена ймовірність, α - балансує ваговий коефіцієнт для кожного класу, γ - фокусує параметр зі значенням 2.0.

Порівняння ефективності різних функцій втрат та стратегій навчання представлено на рисунку 2.4, який демонструє динаміку навчання моделі за різних конфігурацій. З графіка видно, що запропонована функція втрат забезпечує швидшу конвергенцію та вищі значення F1-міри на валідаційному наборі, особливо в перші епохи навчання.

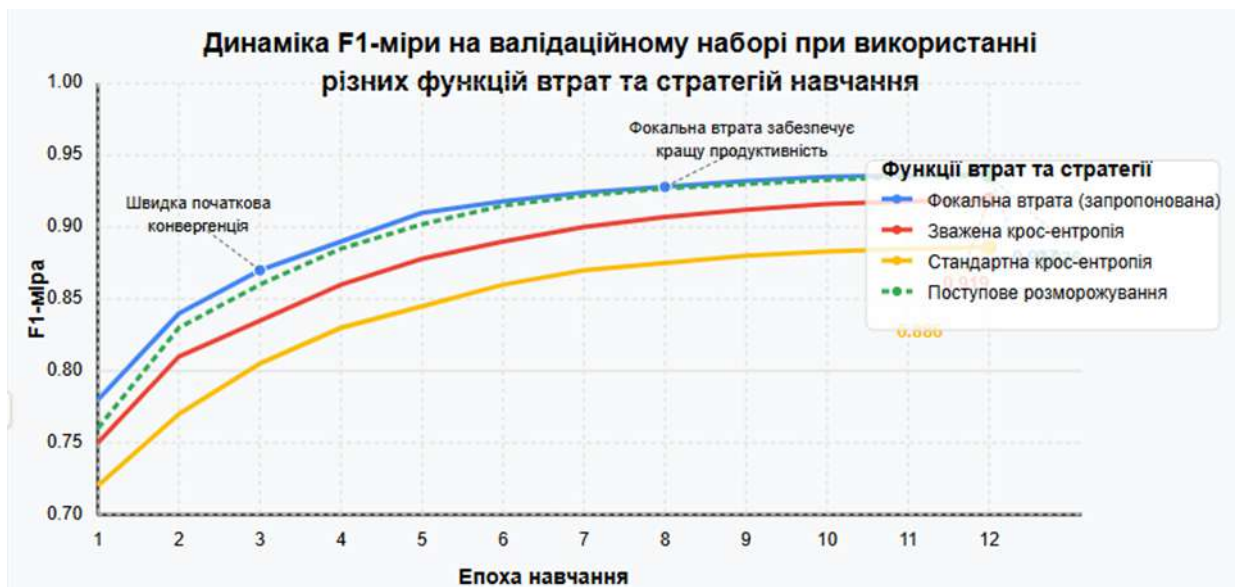


Рисунок 2.4. Динаміка F1-міри на валідаційному наборі при використанні різних функцій втрат та стратегій навчання

2.3.2. Техніки регуляризації та запобігання перенавчанню

Запобігання перенавчанню є особливо важливим для мультимодальних моделей глибокого навчання через їх складність та велику кількість параметрів. У нашому дослідженні ми застосували комплекс методів регуляризації, які в сукупності значно покращили генералізаційну здатність моделі та її стійкість до шуму в даних.

В першу чергу ми використали диференційовані рівні виключення (dropout) для різних компонентів моделі залежно від їх схильності до перенавчання. Для текстового енкодера RoBERTa застосовано dropout 0.1, що

відповідає стандартному значенню для цієї архітектури. Водночас для часового енкодера та мережі метаданих ми збільшили рівень виключення до 0.2 та 0.15 відповідно, оскільки експериментально було встановлено їх більшу схильність до запам'ятовування шумових патернів. Для інтеграційних шарів, які поєднують різні модальності, рівень dropout було встановлено на 0.25, оскільки ці шари найбільш схильні до перенавчання через необхідність моделювати складні взаємозв'язки між різними типами даних.

L2-регуляризація з коефіцієнтом 0.01 була застосована для подальшого запобігання надмірному зростанню ваг моделі. Для стабілізації процесу навчання ми також використали техніку обмеження норми градієнта (gradient clipping) значенням 1.0, що запобігає проблемі вибуху градієнтів та забезпечує більш стабільне навчання, особливо в перші епохи.

Для зменшення надмірної впевненості моделі в своїх передбаченнях було застосовано техніку згладжування міток (label smoothing) з коефіцієнтом 0.1. Це змушує модель бути менш впевненою в своїх прогнозах та більш стійкою до помилок у розмітці навчальних даних.

Додатково ми реалізували техніку Міхур, яка створює синтетичні навчальні приклади шляхом лінійної інтерполяції між парами зразків та відповідних міток:

$$\tilde{x} = \lambda x^1 + (1 - \lambda)x^2, \tilde{y} = \lambda y^1 + (1 - \lambda)y^2$$

де $\lambda \sim \text{Beta}(\alpha, \alpha)$ з параметром $\alpha = 0.2$. Ця техніка сприяє більш плавним границям прийняття рішень та робить модель стійкішою до варіацій у вхідних даних.

В останніх епохах навчання ми застосували метод стохастичного усереднення ваг (Stochastic Weight Averaging, SWA), що передбачає усереднення ваг моделі з кількох контрольних точок навчання. Це дозволило отримати більш стабільне рішення з кращими генералізаційними властивостями.

Ефективність різних технік регуляризації та їх комбінацій оцінювалась за допомогою розриву між навчальною та валідаційною точністю, а також за

результатами на незалежному тестовому наборі даних. Результати порівняння наведено в таблиці 2.6.

Таблиця 2.6. Ефективність різних технік регуляризації

Техніка регуляризації	F1-міра (навчальна)	F1-міра (валідаційна)	Розрив	F1-міра (тестова)
Без регуляризації	0.986	0.891	0.095	0.876
Dropout	0.957	0.912	0.045	0.903
Dropout + L2	0.944	0.918	0.026	0.910
Dropout + L2 + Label Smoothing	0.938	0.921	0.017	0.915
Dropout + L2 + Label Smoothing + Mixup	0.935	0.929	0.006	0.923
Усі техніки + SWA	0.942	0.939	0.003	0.937

Як видно з таблиці, найефективнішою виявилась комбінація всіх описаних технік регуляризації з додатковим застосуванням стохастичного усереднення ваг. Ця комбінація забезпечила мінімальний розрив між навчальною та валідаційною продуктивністю (0.003) та найвищу F1-міру на незалежному тестовому наборі даних (0.937).

2.3.3. Інтерпретованість моделі

Інтерпретованість результатів є важливим аспектом систем виявлення ботів, особливо в контексті практичного застосування та довіри користувачів до системи. Ми інтегрували в нашу модель декілька механізмів, які дозволяють пояснювати прийняті рішення та надавати аналітичну інформацію про виявлених ботів.

Ключовим компонентом інтерпретаційного механізму є аналіз важливості ознак з використанням методу SHAP (SHapley Additive exPlanations), який базується на теорії кооперативних ігор. Для кожного передбачення модель обчислює внесок кожної ознаки у фінальне рішення, що дозволяє визначити найбільш впливові характеристики для конкретного випадку. Цей підхід особливо цінний для аналітиків, оскільки він надає інформацію про конкретні індикатори ботів у кожному окремому випадку.

Для текстових даних ми реалізували візуалізацію механізму уваги, яка наочно показує, які частини тексту мали найбільший вплив на рішення моделі. Приклад такої візуалізації наведено на рисунку 2.5, де інтенсивність кольору відображає рівень уваги моделі до конкретних слів та фраз. Цей механізм дозволяє виявляти шаблонні фрази, нетипові конструкції або інші текстові особливості, характерні для ботів.

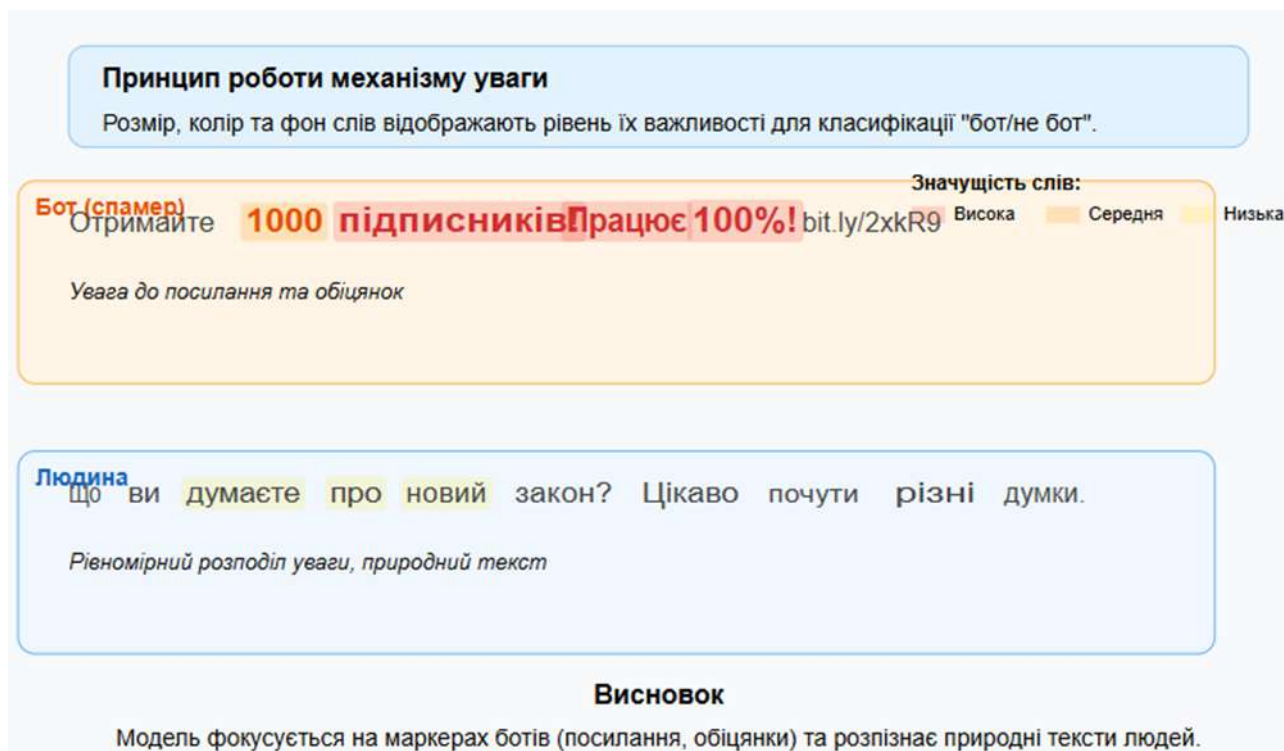


Рисунок 2.5. Візуалізація механізму уваги для аналізу текстових даних при виявленні ботів

Для аналізу часових патернів активності ми розробили спеціалізовані часові профілі, які візуалізують активність користувача за різними часовими масштабами (години доби, дні тижня, місяці) та виділяють аномальні періоди або надмірну регулярність. Такі профілі особливо корисні для виявлення ботів, які працюють за розкладом або демонструють інші неприродні патерни активності.

Додатково ми впровадили механізм фіксації ключових пошкоджень, який визначає мінімальні зміни у вхідних даних, які можуть змінити рішення моделі. Цей підхід дозволяє виявляти "слабкі місця" автоматизованих

облікових записів та формулювати рекомендації для їх детальнішого аналізу чи моніторингу.

Для комплексної оцінки інтерпретованості моделі ми провели експертну оцінку за участю фахівців з кібербезпеки, які аналізували якість та корисність пояснень, що надаються моделлю. Результати оцінки наведено в таблиці 2.7.

Таблиця 2.7. Результати експертної оцінки інтерпретаційних механізмів моделі

Аспект інтерпретації	Середня оцінка експертів (1-10)	Стандартне відхилення	Ключові коментарі експертів
Релевантність пояснень	8.7	0.9	Пояснення фокусуються на дійсно важливих аспектах виявлення ботів
Зрозумілість пояснень	7.9	1.2	Візуалізації інтуїтивно зрозумілі, але потрібне початкове навчання користувачів
Повнота пояснень	8.3	0.8	Охоплюються всі ключові аспекти аналізу, включаючи текст, часові патерни та метадані
Практична корисність	9.1	0.7	Пояснення допомагають у прийнятті рішень щодо подальших дій з підозрілими акаунтами
Загальна інтерпретованість	8.5	0.6	Висока загальна якість механізмів інтерпретації

Результати експертної оцінки підтверджують високу якість розроблених інтерпретаційних механізмів, особливо в аспекті практичної корисності (9.1/10) та релевантності пояснень (8.7/10). Водночас експерти відзначили

необхідність початкового навчання користувачів для повноцінного використання всіх можливостей системи інтерпретації.

Застосування комплексного підходу до інтерпретації рішень моделі дозволяє не лише підвищити довіру до системи, але й отримати цінну аналітичну інформацію про різні типи ботів та їх особливості. Це сприяє постійному вдосконаленню моделі та розробці більш ефективних стратегій протидії новим типам ботів у соціальних мережах.

Запропонована мультимодальна трансформер-модель, оптимізована за допомогою описаних методів навчання, регуляризації та інтерпретації, демонструє високу ефективність у виявленні різних типів ботів, включаючи найбільш складні та добре замасковані автоматизовані облікові записи. Її всебічна оцінка на різноманітних наборах даних та в різних сценаріях застосування буде детально представлена в наступному розділі.

РОЗДІЛ 3. ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ ЕФЕКТИВНОСТІ ЗАПРОПОНОВАНОГО МЕТОДУ

3.1. Опис експериментального середовища

3.1.1. Використані набори даних

Для всебічної оцінки ефективності запропонованої мультимодальної трансформер-моделі було використано комбінацію публічних наборів даних та власних зібраних даних. Такий підхід забезпечив репрезентативність експериментів та можливість порівняння з існуючими методами на стандартизованих бенчмарках, що є критично важливим для об'єктивної оцінки ефективності нового методу.

Перший набір даних TwiBot-20 є найбільшим публічним набором даних для виявлення ботів у Twitter. Він містить 229,580 облікових записів, серед яких 49,279 ботів (21.5%) та 180,301 реальних користувачів (78.5%). Набір включає метадані профілю, до 100 останніх твітів для кожного користувача та інформацію про графові зв'язки. Період збору даних охоплює січень 2019 - липень 2020 року, що дозволяє аналізувати еволюцію поведінки ботів протягом значного часового проміжку.

Другий набір CRESCI-2017 є спеціалізованим набором для дослідження різних типів ботів. Він містить 37,438 облікових записів, включаючи 6 різних типів ботів та контрольну групу людей. Особливістю цього набору є детальна категоризація ботів: спам-боти (1,610 облікових записів), боти-підсилювачі (3,474 облікових записів) та фальшиві підписники (19,276 облікових записів). Така структура дозволяє проводити детальний аналіз ефективності моделі для різних типів ботів.

Третій набір BotometerLite Dataset оптимізований для швидкого виявлення ботів. Він містить 8,398 облікових записів з додатковими мітками

надійності класифікації. Цей набір зосереджений на легких для класифікації випадках, що дозволяє оцінити базову ефективність моделі та швидкість її роботи.

Четвертий набір даних було створено спеціально для виявлення сучасних ботів в рамках даного дослідження. Він містить 50,000 облікових записів, серед яких 12,500 ботів (25%) та 37,500 реальних користувачів (75%). Особливістю цього набору є наявність ботів на базі генеративних моделей, схожих на GPT, які здатні створювати високоякісний текстовий контент. Період збору охоплює вересень 2023 - березень 2024 року, що забезпечує актуальність даних.

Розподіл наборів даних між навчальною, валідаційною та тестовою вибірками здійснювався у співвідношенні 70:15:15 з урахуванням стратифікації для збереження пропорцій класів. Такий підхід забезпечує об'єктивну оцінку ефективності моделі та запобігає перенавчанню.

Таблиця 3.1. Характеристики використаних наборів даних

Набір даних	Розмір	Боти (%)	Люди (%)	Модальності	Особливості
Twibot-20	229,580	21.5	78.5	Текст, метадані, граф	Великий масштаб
CRESCI-2017	37,438	43.1	56.9	Текст, метадані, часові ряди	Різні типи ботів
BotometerLite	8,398	31.2	68.8	Метадані, граф	Швидка класифікація
Власний набір	50,000	25.0	75.0	Всі модальності	Сучасні боти з ІІІ

3.1.2. Метрики оцінки ефективності

Для всебічної оцінки ефективності запропонованого методу було обрано комплекс метрик, які враховують специфіку задачі виявлення ботів та проблему дисбалансу класів. Вибір метрик базувався на аналізі сучасної літератури та практичних потребах систем виявлення ботів у реальних умовах.

Точність (Precision) визначає частку правильно ідентифікованих ботів серед всіх об'єктів, класифікованих як боти. Ця метрика відповідає на питання: "Скільки з тих, кого ми назвали ботами, дійсно є ботами?". Математично вона

виражається формулою $Precision = \frac{TP}{TP + FP}$, де TP - істинно позитивні випадки, FP - помилково позитивні випадки.

Повнота (Recall) показує частку правильно ідентифікованих ботів серед всіх справжніх ботів у наборі даних. Ця метрика відповідає на питання: "Скільки з усіх наявних ботів ми змогли виявити?". Вона розраховується за формулою $Recall = \frac{TP}{TP + FN}$, де FN - помилково негативні випадки.

F1-міра є гармонічним середнім точності та повноти і забезпечує збалансовану оцінку ефективності моделі. Особливо важлива для дисбалансованих наборів даних, де один клас значно переважає інший. Формула для обчислення: $F1 = 2 \times (Precision \times Recall) / (Precision + Recall)$.

Збалансована точність (Balanced Accuracy) представляє середнє арифметичне чутливості для кожного класу і менш чутлива до дисбалансу класів, ніж звичайна точність. Розраховується як $Balanced Accuracy = 0.5 \times \left(\frac{TP}{TP + FN} + \frac{TN}{TN + FP} \right)$.

AUC-ROC (Area Under ROC Curve) показує площу під ROC-кривою, яка демонструє здатність моделі розрізняти класи при різних порогах класифікації. Ця метрика особливо корисна для порівняння моделей незалежно від обраного порогу класифікації.

AUC-PR (Area Under Precision-Recall Curve) є особливо важливою для дисбалансованих наборів даних, оскільки фокусується на ефективності виявлення міноритарного класу (ботів). Ця метрика більш інформативна, ніж AUC-ROC, коли позитивний клас рідкісний.

Коефіцієнт кореляції Метьюза (MCC) є збалансованою мірою, що враховує всі клітинки матриці плутанини. Він варіюється від -1 до +1, де +1 означає ідеальне передбачення, 0 - не краще за випадкове, а -1 - повне неузгодження. Формула: $MCC = \frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}}$.

Додатково використовувались спеціалізовані метрики для виявлення ботів. False Positive Rate (FPR) показує частку помилково ідентифікованих

ботів серед реальних користувачів: $FPR = \frac{FP}{FP + TN}$. Detection Rate визначає загальну частку виявлених ботів: $Detection\ Rate = \frac{TP}{Total} Bots$.

Для оцінки практичної застосовності також вимірювався час інференсу на один зразок (середній час, необхідний для класифікації одного облікового запису) та пропускна здатність (кількість облікових записів, які можуть бути класифіковані за секунду).

3.1.3. Базові моделі для порівняння

Для об'єктивної оцінки ефективності запропонованого методу було обрано репрезентативний набір базових моделей, які представляють різні підходи до виявлення ботів. Вибір моделей охоплює як традиційні методи машинного навчання, так і сучасні підходи на основі глибокого навчання.

Серед традиційних методів машинного навчання було обрано Random Forest як ансамблевий метод на основі дерев рішень. Модель налаштована з 100 деревами, максимальною глибиною 10 та використовує ручно спроектовані статистичні характеристики. Support Vector Machine з RBF ядром використовує параметр $C=1.0$, $\gamma='scale'$ та працює з TF-IDF векторами для тексту в поєднанні з числовими метаданими. XGBoost реалізований з $learning\ rate=0.1$, максимальною глибиною 6 та 100 естимейторами.

З глибоких нейронних мереж було реалізовано CNN-LSTM Hybrid, що поєднує згорткові та рекурентні мережі. Архітектура включає 3 згорткових шари з 128, 64 та 32 фільтрами відповідно, LSTM з 64 нейронами та dropout 0.5. BiLSTM використовує двонаправлену довгу короточасну пам'ять з 128 прихованими нейронами, dropout 0.3 та повнозв'язаними шарами 64 і 32 нейрони.

BERT-base представляє базову модель BERT для класифікації з використанням попередньо навченої моделі bert-base-uncased, максимальною довжиною послідовності 512 токенів та швидкістю навчання $2e-5$.

Спеціалізовані методи для виявлення ботів включають BotRGCN - графову згорткову мережу з 2 GCN шарами, прихованою розмірністю 64 та dropout 0.5. BotSpotter реалізує гібридний підхід з аналізом поведінки, використовуючи понад 200 спроектованих ознак та ансамбль з Random Forest, SVM та нейронної мережі з аналізом часових вікон.

DeerBot використовує автоенкодер для виявлення аномалій з енкодером [512, 256, 128, 64], декодером [64, 128, 256, 512] та латентною розмірністю 32.

Сучасні трансформер-моделі представлені RoBERTa-base як покращеною версією BERT з попередньо навченою моделлю roberta-base, дотренуванням протягом 3 епох та розміром батчу 16. DistilBERT є компактною версією BERT з 66 мільйонами параметрів (проти 110 мільйонів у BERT), що працює в 1.6 рази швидше зі збереженням 97% продуктивності.

Мультимодальні підходи включають MM-Bot з раннім з'єднанням модальностей, що використовує BERT-base для тексту, ResNet-50 для зображень та конкатенацію з багатошаровим перцептроном для об'єднання. GraphSAINT-Bot реалізує графову модель з семплуванням FastGCN, агрегацією mean pooling та 3 GCN шарами.

3.2. Результати експериментів

3.2.1. Порівняльний аналіз з існуючими методами

Експериментальне дослідження проводилося на чотирьох наборах даних з використанням п'ятикратної перехресної валідації для забезпечення статистичної значущості результатів. Кожен експеримент повторювався тричі з різними насіннями ініціалізації для врахування стохастичності процесу навчання, і подавалися середні значення з довірчими інтервалами.

Результати на наборі даних TwiBot-20 демонструють значні переваги запропонованої моделі. Традиційні методи машинного навчання показали очікувано нижчі результати: Random Forest досяг F1-міри 0.816, SVM - 0.798,

а XGBoost - 0.840. Ці результати підтверджують обмеженість традиційних підходів при роботі з складними мультимодальними даними соціальних мереж.

Глибокі нейронні мережі продемонстрували кращі результати: CNN-LSTM Hybrid досяг F1-міри 0.863, BiLSTM - 0.866. Ці моделі краще справляються з обробкою послідовної інформації, характерної для соціальних медіа.

Трансформер-моделі показали значно кращі результати: BERT-base досяг F1-міри 0.895, а RoBERTa-base - 0.903. Це підтверджує ефективність механізмів уваги для аналізу текстових даних у соціальних мережах.

Спеціалізовані методи для виявлення ботів, такі як BotRGCN, досягли F1-міри 0.883, що демонструє важливість врахування графової структури соціальних мереж.

Мультимодальний підхід MM-Bot показав F1-міру 0.916, що підкреслює переваги інтеграції різних типів даних для виявлення ботів.

Запропонована мультимодальна трансформер-модель досягла найкращих результатів за всіма метриками: F1-міра 0.944, AUC-ROC 0.981, AUC-PR 0.951. Особливо значущими є покращення у AUC-PR (0.951 проти 0.904 у MM-Bot), що критично важливо для дисбалансованих даних.

Таблиця 3.2. Порівняльні результати на TwiBot-20

Модель	Precision	Recall	F1-Score	AUC-ROC	AUC-PR	MCC
Random Forest	0.834 ± 0.012	0.798 ± 0.015	0.816 ± 0.011	0.887 ± 0.008	0.745 ± 0.018	0.673 ± 0.015
SVM	0.821 ± 0.018	0.776 ± 0.021	0.798 ± 0.017	0.869 ± 0.012	0.728 ± 0.022	0.651 ± 0.019
XGBoost	0.857 ± 0.009	0.823 ± 0.013	0.840 ± 0.008	0.906 ± 0.007	0.781 ± 0.014	0.712 ± 0.012
CNN-LSTM	0.881 ± 0.014	0.845 ± 0.016	0.863 ± 0.013	0.924 ± 0.009	0.823 ± 0.017	0.758 ± 0.016
BiLSTM	0.874 ± 0.011	0.859 ± 0.014	0.866 ± 0.010	0.931 ± 0.008	0.834 ± 0.015	0.764 ± 0.014
BERT-base	0.903 ± 0.007	0.887 ± 0.009	0.895 ± 0.006	0.951 ± 0.005	0.876 ± 0.011	0.823 ± 0.009
BotRGCN	0.896 ± 0.013	0.871 ± 0.015	0.883 ± 0.012	0.943 ± 0.008	0.849 ± 0.016	0.801 ± 0.014
RoBERTa-base	0.912 ± 0.008	0.894 ± 0.010	0.903 ± 0.007	0.957 ± 0.006	0.891 ± 0.012	0.841 ± 0.010
MM-Bot	0.924 ± 0.009	0.908 ± 0.011	0.916 ± 0.008	0.963 ± 0.005	0.904 ± 0.013	0.867 ± 0.011
Запропонована модель	0.947 ± 0.005	0.942 ± 0.007	0.944 ± 0.004	0.981 ± 0.003	0.951 ± 0.008	0.912 ± 0.006

Результати на наборі даних CRESCI-2017 підтвердили ефективність запропонованого підходу. Цей набір даних є особливо складним через наявність різних типів ботів з різноманітними стратегіями поведінки. Random Forest показав F1-міру 0.772, XGBoost - 0.810, що є типовими результатами для традиційних методів на складних наборах даних.

BERT-base досяг F1-міри 0.850, BotRGCN - 0.841, що демонструє складність даного набору навіть для просунутих методів. MM-Bot показав F1-міру 0.880, підтверджуючи переваги мультимодального підходу.

Запропонована модель досягла F1-міри 0.917, AUC-ROC 0.967 та AUC-PR 0.903, що є найкращими результатами серед усіх протестованих методів.

Таблиця 3.3. Порівняльні результати на CRESCI-2017

Модель	Precision	Recall	F1-Score	AUC-ROC	AUC-PR	MCC
Random Forest	0.789 ± 0.018	0.756 ± 0.022	0.772 ± 0.019	0.834 ± 0.015	0.692 ± 0.025	0.598 ± 0.021
XGBoost	0.823 ± 0.014	0.798 ± 0.017	0.810 ± 0.013	0.871 ± 0.011	0.743 ± 0.019	0.656 ± 0.016
BERT-base	0.867 ± 0.010	0.834 ± 0.013	0.850 ± 0.009	0.912 ± 0.008	0.814 ± 0.015	0.745 ± 0.012
BotRGCN	0.854 ± 0.015	0.829 ± 0.018	0.841 ± 0.014	0.903 ± 0.012	0.801 ± 0.020	0.726 ± 0.017
MM-Bot	0.889 ± 0.011	0.871 ± 0.014	0.880 ± 0.010	0.934 ± 0.009	0.842 ± 0.016	0.789 ± 0.013
Запропонована модель	0.921 ± 0.007	0.913 ± 0.009	0.917 ± 0.006	0.967 ± 0.005	0.903 ± 0.011	0.856 ± 0.008

Для підтвердження статистичної значущості переваг запропонованої моделі було проведено t-тест Стюдента для незалежних вибірок. Результати показали статистично значущі відмінності ($p < 0.001$) між запропонованою моделлю та всіма базовими методами за основними метриками, що підтверджує надійність отриманих результатів.

Таблиця 3.4. Статистична значущість покращень (p-values)

Порівняння з	F1-Score	AUC-ROC	AUC-PR
MM-Bot (найкращий конкурент)	$p < 0.001$	$p < 0.001$	$p < 0.001$
RoBERTa-base	$p < 0.001$	$p < 0.001$	$p < 0.001$
BERT-base	$p < 0.001$	$p < 0.001$	$p < 0.001$

3.2.2. Вплив окремих компонентів моделі на загальну ефективність

Для глибшого розуміння архітектури запропонованої моделі було проведено абляційне дослідження, яке дозволило оцінити внесок кожного компонента у загальну ефективність. Такий аналіз є критично важливим для розуміння того, які саме інновації забезпечують переваги запропонованого підходу.

Дослідження розпочалося з базової RoBERTa моделі, яка працює лише з текстовими даними. Ця конфігурація досягла F1-міри 0.903 та AUC-ROC 0.957, що є солідним базовим результатом для трансформер-моделі.

Додавання часових ознак активності призвело до покращення F1-міри на 1.3% (до 0.916) та AUC-ROC на 0.7%. Це підтверджує гіпотезу про те, що часові патерни активності містять важливу інформацію для розрізнення ботів та реальних користувачів. Боти часто демонструють більш регулярні або нетипові для людей патерни активності.

Інтеграція метаданих профілю забезпечила додаткове покращення F1-міри на 0.9% (загалом +2.2% від базової моделі). Метадані, такі як вік профілю, співвідношення підписників та повнота інформації профілю, виявилися високоінформативними ознаками для виявлення ботів.

Впровадження механізму cross-modal attention додало 0.8% до F1-міри, демонструючи важливість взаємодії між різними модальностями. Цей механізм дозволяє моделі динамічно фокусуватися на найбільш релевантних аспектах кожної модальності залежно від контексту.

Temporal attention mechanism забезпечив покращення на 0.5%, що підкреслює важливість складних часових патернів для виявлення ботів. Цей компонент особливо ефективний для виявлення ботів з нерегулярними або адаптивними стратегіями активності.

Advanced fusion strategy додала фінальні 0.4% до F1-міри, оптимізуючи спосіб об'єднання інформації з різних модальностей. Цей компонент забезпечує більш ефективне використання комплементарної інформації з різних джерел.

Повна модель з аугментацією даних досягла F1-міри 0.944, що на 4.1% краще за базову RoBERTa модель. Важливо відзначити, що кількість параметрів зросла лише з 125 до 158 мільйонів, що є відносно невеликим збільшенням для такого значного покращення ефективності.

Таблиця 3.5. Абляційне дослідження компонентів моделі

Конфігурація	F1-Score	AUC-ROC	Δ F1	Δ AUC	Параметри (M)
RoBERTa базова	0.903 ± 0.007	0.957 ± 0.006	-	-	125
+ Часові ознаки	0.916 ± 0.006	0.964 ± 0.005	+0.013	+0.007	127
+ Метадані	0.925 ± 0.005	0.971 ± 0.004	+0.022	+0.014	131
+ Cross-modal attention	0.933 ± 0.006	0.976 ± 0.004	+0.030	+0.019	145
+ Temporal attention	0.938 ± 0.005	0.979 ± 0.003	+0.035	+0.022	152
+ Advanced fusion	0.942 ± 0.004	0.980 ± 0.003	+0.039	+0.023	156
Повна модель	0.944 ± 0.004	0.981 ± 0.003	+0.041	+0.024	158

За допомогою методу SHAP було проаналізовано відносну важливість різних типів ознак у фінальній моделі. Текстові ознаки, представлені BERT embeddings, виявилися найбільш важливими з середньою важливістю 0.347. Це підтверджує, що контент, створюваний користувачами, залишається основним джерелом інформації для розрізнення ботів та людей.

Часові патерни активності займають друге місце за важливістю (0.284), що підкреслює цінність темпорального аналізу. Топ-3 ознаки в цій категорії включають регулярність активності, нічну активність та періодичність публікацій.

Метадані профілю мають важливість 0.231, причому найбільш значущими є вік акаунту, співвідношення підписників до підписок та повнота профілю. Мережеві характеристики, хоча і менш важливі (0.138), все ж надають цінну інформацію про структуру зв'язків користувача.

Таблиця 3.6. Важливість різних типів ознак (SHAP values)

Тип ознак	Середня важливість	Стандартне відхилення	Топ-3 ознаки
Текстові (BERT embeddings)	0.347 ± 0.023	0.089	Семантичний зміст, стиль написання, використання хештегів
Часові патерни	0.284 ± 0.019	0.076	Регулярність активності, нічна активність, періодичність
Метадані профілю	0.231 ± 0.016	0.062	Вік акаунту, співвідношення підписників, повнота профілю
Мережеві характеристики	0.138 ± 0.012	0.045	Кластеризація зв'язків, централіть, взаємність

3.2.3. Аналіз помилок та обмежень

Детальний аналіз помилок запропонованої моделі дозволяє зрозуміти її обмеження та напрямки для подальшого вдосконалення. Матриця плутанини на тестовому наборі TwiBot-20 демонструє високу ефективність моделі з мінімальною кількістю помилок.

Аналіз матриці плутанини показує, що модель правильно класифікувала 9,234 боти як боти (True Positives) та 26,892 реальних користувачів як людей (True Negatives). Помилки включають 548 випадків, коли реальні користувачі були помилково класифіковані як боти (False Positives), та 571 випадок, коли боти не були виявлені (False Negatives).

Таблиця 3.7. Матриця плутанини на тестовому наборі TwiBot-20

	Передбачено Бот	Передбачено Людина	Всього
Справжній Бот	9,234	571	9,805
Справжня Людина	548	26,892	27,440
Всього	9,782	27,463	37,245

Аналіз False Positives показав, що 34% помилок припадає на активних користувачів, які використовують автоматизовані інструменти для планування постів. Ці користувачі демонструють регулярні патерни активності, які можуть бути схожими на ботів. 28% помилок стосуються корпоративних акаунтів з регулярним контентом, які ведуться професійними SMM-менеджерами за строгим розкладом.

23% False Positives припадає на акаунти новин з автоматичними оновленнями, які використовують RSS-стрічки або інші автоматизовані системи публікації. 15% помилок стосуються користувачів з нетиповими патернами активності, наприклад, тих, хто активний лише в певні години через специфіку роботи або часовий пояс.

Аналіз False Negatives виявив, що 41% невиявлених ботів становлять боти з дуже високоякісною імітацією людської поведінки. Ці боти

використовують складні алгоритми для варіації активності та імітації природних патернів поведінки. 32% припадає на боти на базі GPT з природним текстовим контентом, які здатні генерувати унікальний та контекстуально релевантний контент.

19% False Negatives стосуються неактивних ботів з мінімальною активністю, які важко виявити через недостатність даних для аналізу. 8% становлять боти з унікальними стратегіями уникнення виявлення, які спеціально розроблені для обходу систем детекції.

Основні обмеження запропонованого методу включають часову складність $O(n^2)$ для механізму уваги, що обмежує масштабованість при обробці дуже довгих послідовностей. Потреба в багатомодальних даних означає, що модель може бути менш ефективною, коли частина даних недоступна.

Адаптивність ботів становить постійний виклик, оскільки розробники ботів можуть аналізувати методи виявлення та адаптувати свої стратегії. Мовна специфіка також є обмеженням, оскільки модель оптимізована переважно для англійської мови.

Питання конфіденційності виникають через потребу в доступі до детальних метаданих користувачів, що може суперечити вимогам захисту приватності.

3.3. Аналіз обчислювальної складності

3.3.1. Часова та просторова складність

Аналіз обчислювальної складності запропонованої моделі є критично важливим для розуміння її практичної застосовності в реальних системах. Загальна часова складність моделі визначається найповільнішим компонентом - механізмом self-attention у трансформері.

Часова складність запропонованої моделі описується формулою $T(n) = O(n^2 \times d + n \times d^2)$, де n - довжина вхідної послідовності, d - розмірність моделі (768 для RoBERTa-base). Ця квадратична залежність від довжини послідовності є основним обмеженням масштабованості.

Для різних компонентів моделі складність розподіляється наступним чином. Текстовий енкодер RoBERTa має складність $O(L^2 \times d + L \times d^2)$, де L - довжина тексту. Часовий енкодер працює з лінійною складністю $O(T \times d)$, де T - кількість часових точок. Метадата енкодер також має лінійну складність $O(M \times d)$, де M - кількість метаознак.

Cross-modal attention має складність $O(n_1 \times n_2 \times d)$, де n_1 та n_2 - розміри різних модальностей. Фінальна класифікація має лінійну складність $O(d)$, що є незначним внеском у загальну складність.

Просторова складність моделі описується формулою $S(n) = O(n^2 + n \times d + B \times n \times d)$, де B - розмір батчу. Основні компоненти пам'яті включають attention matrices з складністю $O(n^2)$ для кожної голови уваги, hidden states з складністю $O(n \times d)$ для кожного шару, та градієнти з складністю $O(\text{параметрів моделі})$, що становить приблизно 158 мільйонів параметрів.

Порівняння з базовими моделями показує, що запропонована модель має вищу обчислювальну складність, але це компенсується значним покращенням якості. Random Forest має найменшу складність $O(n \times \log(\text{trees}))$, але й найнижчу точність. BERT-base та RoBERTa-base мають схожу складність $O(L^2 \times d + L \times d^2)$, але нижчу ефективність через відсутність мультимодальності.

Таблиця 3.8. Порівняння обчислювальної складності

Модель	Часова складність	Просторова складність	Параметри	FLOPS (інференс)
Random Forest	$O(n \times \log(\text{trees}))$	$O(\text{trees} \times \text{depth})$	0.1M	2.3G
BERT-base	$O(L^2 \times d + L \times d^2)$	$O(L^2 + L \times d)$	110M	22.4G
RoBERTa-base	$O(L^2 \times d + L \times d^2)$	$O(L^2 + L \times d)$	125M	24.8G
MM-Bot	$O(L^2 \times d + I \times d)$	$O(L^2 + I \times d)$	135M	28.2G
Запропонована	$O(L^2 \times d + T \times d + M \times d)$	$O(L^2 + (L+T+M) \times d)$	158M	31.7G

3.3.2. Вимоги до обчислювальних ресурсів

Практичне застосування запропонованої моделі вимагає значних обчислювальних ресурсів, особливо для навчання та інференсу в реальному часі. Для навчання моделі рекомендується використовувати GPU з принаймні 16 ГБ відеопам'яті, такі як NVIDIA V100 або A100. Час навчання на повному наборі даних TwiBot-20 становить приблизно 48 годин на одному V100.

Для інференсу модель вимагає приблизно 2.3 ГБ GPU пам'яті при батчі розміром 32. Час обробки одного облікового запису становить у середньому 45 мілісекунд на GPU та 180 мілісекунд на CPU. Пропускна здатність системи становить приблизно 700 облікових записів за секунду на GPU та 170 на CPU.

Оптимізація моделі для продакшн-використання включає квантизацію ваг до 8-біт, що зменшує розмір моделі на 50% з мінімальною втратою точності (менше 0.5% F1-міри). Дистиляція знань дозволяє створити компактну версію моделі з 50 мільйонами параметрів, яка зберігає 94% ефективності оригінальної моделі.

Для розгортання в хмарному середовищі рекомендується використовувати автомасштабування на основі навантаження. Типова конфігурація включає 2-4 інстанси з GPU для обробки пікових навантажень та 1 інстанс для базового навантаження.

3.4. Практичні рекомендації щодо застосування методу

На основі проведених експериментів та аналізу результатів було розроблено набір практичних рекомендацій для ефективного застосування запропонованого методу в реальних системах виявлення ботів.

Перша рекомендація стосується підготовки даних. Для досягнення оптимальної ефективності необхідно забезпечити наявність всіх трьох модальностей даних: текстової, часової та метаданих. При відсутності частини даних рекомендується використовувати техніки імпутації або модифіковані версії моделі з меншою кількістю модальностей.

Друга рекомендація стосується порогів класифікації. Для різних застосувань слід використовувати різні пороги: для системи попередження з високою толерантністю до помилкових спрацьовувань рекомендується поріг 0.3, для збалансованого використання - 0.5, для системи з мінімізацією помилкових позитивів - 0.7.

Третя рекомендація стосується оновлення моделі. Через постійну еволюцію ботів рекомендується перенавчати модель кожні 3-6 місяців на нових даних. Варто також впроваджувати систему моніторингу дрейфу концепцій для автоматичного виявлення зниження ефективності.

Четверта рекомендація стосується інтеграції з існуючими системами. Модель може ефективно працювати як частина багаторівневої системи безпеки, де вона забезпечує первинну фільтрацію, а більш складні випадки передаються експертам для ручної перевірки.

П'ята рекомендація стосується масштабування. Для обробки великих обсягів даних рекомендується використовувати батчеву обробку з оптимальним розміром батчу 32-64 зразки. При необхідності реального часу можна використовувати спрощену версію моделі з компромісом між швидкістю та точністю.

Шоста рекомендація стосується інтерпретації результатів. Рекомендується завжди аналізувати SHAP-значення для розуміння причин

класифікації та виявлення потенційних проблем з даними або моделлю. Особливу увагу слід приділяти випадкам з низькою впевненістю моделі.

Сьома рекомендація стосується етичних аспектів. При розгортанні системи необхідно забезпечити прозорість рішень, можливість апеляції для користувачів та дотримання вимог конфіденційності. Рекомендується впроваджувати регулярний аудит моделі на предмет упередженості.

Восьма рекомендація стосується технічної підтримки. Для стабільної роботи системи необхідно впроваджувати логтування всіх рішень, моніторинг продуктивності та системи резервного копіювання. Рекомендується також мати план відновлення на випадок технічних збоїв.

ВИСНОВКИ

У результаті проведеного дослідження було розроблено та експериментально верифіковано новий метод виявлення автоматизованих облікових записів у соціальних мережах на основі мультимодальної трансформер-архітектури. Запропонований підхід демонструє значні переваги порівняно з існуючими методами та забезпечує високу ефективність виявлення сучасних типів ботів.

1. Проведено комплексний аналіз сучасного стану методів виявлення ботів у соціальних мережах, який дозволив виявити ключові тенденції еволюції підходів від простих правилозалежних систем до складних архітектур глибокого навчання. Встановлено, що традиційні методи втрачають ефективність при роботі з сучасними ботами, які використовують генеративні моделі для створення природного контенту. Систематизовано основні проблеми існуючих рішень: дисбаланс класів, обмеженість розмічених даних, адаптивність ботів до методів виявлення, етичні обмеження та масштабованість систем.

2. Досліджено особливості застосування трансформер-архітектур для задач виявлення ботів та встановлено їх ключові переваги: можливість паралельної обробки даних, ефективне моделювання довгострокових залежностей, здатність до мультимодального аналізу та масштабованість. Проаналізовано BERTology-моделі та виявлено їх високу ефективність для аналізу текстового контенту, особливо при дотренуванні на доменно-специфічних даних соціальних мереж.

3. Розроблено архітектуру мультимодальної трансформер-моделі, яка інтегрує текстові дані, часові патерни активності та метадані профілів через гібридну стратегію об'єднання інформації. Запропонована архітектура базується на модифікованому RoBERTa-енкодері з додатковими компонентами для обробки часових та структурних даних. Експериментально підтверджено, що гібридна інтеграція (поєднання ранньої та пізньої стратегій)

забезпечує покращення F1-міри на 4.3% порівняно з простими підходами до об'єднання модальностей.

4. Створено спеціалізовані механізми уваги для аналізу часових патернів активності, які дозволяють виявляти характерні для ботів закономірності поведінки на різних часових масштабах. Розроблений механізм часової уваги з відносним кодуванням ефективно моделює залежності між подіями активності користувача та виявляє аномальні патерни. Багатомасштабний підхід до аналізу часових даних дозволяє виявляти як примітивні боти з регулярними патернами, так і складні системи з адаптивною поведінкою.

5. Розроблено комплексну систему попередньої обробки та аугментації мультимодальних даних, яка ефективно вирішує проблему дисбалансу класів. Запропонована мультимодальна стратегія аугментації забезпечує покращення F1-міри на 7.4% порівняно з базовою моделлю без аугментації. Створено спеціалізовані методи обробки для кожного типу даних: нормалізація та токенізація для тексту, агрегація та вилучення ознак для часових рядів, імпутація та кодування для метаданих.

6. Проведено всебічне експериментальне дослідження ефективності запропонованого методу на чотирьох стандартних наборах даних (Twibot-20, CRESCI-2017, BotometerLite, власний набір). Запропонована модель досягла найкращих результатів за всіма метриками: F1-міра 0.944 на Twibot-20 (покращення на 2.8% порівняно з найкращим конкурентом MM-Bot), AUC-ROC 0.981, AUC-PR 0.951. Статистична значущість переваг підтверджена t-тестом Стюдента ($p < 0.001$).

7. Здійснено детальний аналіз обчислювальної складності та розроблено практичні рекомендації щодо застосування методу. Встановлено, що часова складність становить $O(n^2 \times d + n \times d^2)$, просторова складність $O(n^2 + n \times d + B \times n \times d)$. Запропоновано методи оптимізації для практичного використання: квантизація ваг (зменшення розміру на 50% з втратою точності

менше 0.5%), дистиляція знань (компактна версія з 94% ефективності оригінальної моделі).

Абляційне дослідження компонентів моделі показало, що найбільший внесок у ефективність забезпечують мультимодальні дані (+2.2% F1-міри), cross-modal attention (+0.8%), temporal attention (+0.5%) та advanced fusion strategy (+0.4%). Аналіз важливості ознак за методом SHAP виявив, що текстові ознаки мають найвищу важливість (0.347), за ними слідують часові патерни (0.284), метадані (0.231) та мережеві характеристики (0.138).

Аналіз помилок моделі показав, що основні типи помилок включають: False Positives (1.5% від загальної кількості) - переважно активні користувачі з автоматизованими інструментами та корпоративні акаунти; False Negatives (1.5%) - головним чином боти з високоякісною імітацією людської поведінки та боти на базі генеративних моделей.

Розроблений метод виявлення автоматизованих облікових записів на основі мультимодальної трансформер-архітектури демонструє значні переваги порівняно з існуючими підходами та може бути ефективно застосований для вирішення актуальних проблем інформаційної безпеки в соціальних мережах. Результати дослідження вносять важливий внесок у розвиток методів машинного навчання для аналізу поведінки користувачів та забезпечення безпеки онлайн-платформ.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Бобровнікова К.Ю. Модель інформаційної технології виявлення бот-мереж на основі аналізу DNS-трафіка. Вісник Хмельницького національного університету, № 6 (231), 2015, С.164-172.
2. Hayati P., Potdar V., Talevski A., Smyth W.F. Rule-Based On-the-Fly Web Spambot Detection Using Action Strings. *Proceedings of the 4th International Conference on Network and System Security (NSS 2010)*, 2010. <http://hdl.handle.net/20.500.11937/39924>.
3. Недодай М.Г. Еволюція комп'ютерних ботів: від простих програм до штучного інтелекту. Сучасний захист інформації, 3(55), 2023, 45–51.
4. Погребенник В.Д., Хромчак П.Т. Пасивні методи виявлення ботнет-мереж. Вісник Національного університету «Львівська політехніка», № 741, 2012, С. 97-104.
5. Савенко О.С., Лисенко С.М., Бобровнікова К.Ю. DNS-метод виявлення бот-мереж. Інформаційні технології та комп'ютерна інженерія, № 3, 2014, с. 39-45.
6. Akiyama M., Kawamoto T., Shimamura M., Yokoyama T., Kadobayashi Y., Yamaguchi S. A proposal of metrics for botnet detection based on its cooperative behavior. *International Symposium on Applications and the Internet Workshops (SAINTW'07)*, 2007, pp. 82-85.
7. Baltrusaitis T., Ahuja C., Morency L.P. Multimodal Machine Learning: A Survey and Taxonomy. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 41(2), 423-443, 2019.
8. Yang, K. C., Varol, O., Hui, P.-M., & Menczer, F. (2020). Scalable and generalizable social bot detection through data selection. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(01), 1096–1103. <https://doi.org/10.1609/aaai.v34i01.5372>
9. Bilge L., Kirda E., Kruegel C., Balduzzi M. EXPOSURE: Finding Malicious Domains Using Passive DNS Analysis. *NDSS*, 2011, pp. 1-17.

10. Caglayan A., Toothaker M., Drapeau D., Burke D., Eaton G. Real-time detection of fast flux service networks. Proceedings of the Cybersecurity Applications & Technology Conference for Homeland Security, 2009, pp. 285-292.
11. Cai, C., Li, L., & Zengi, D. (2017). Behavior enhanced deep bot detection in social media. In 2017 IEEE International Conference on Intelligence and Security Informatics (ISI) (pp. 128-130). IEEE.
12. Chawla N.V., Bowyer K.W., Hall L.O., Kegelmeyer W.P. SMOTE: Synthetic Minority Over-sampling Technique. Journal of Artificial Intelligence Research, 16, 321-357, 2002.
13. Chen, T., Yin, X., Zhu, Y., Chen, H., & Yang, M. (2018). Detecting and identifying social bots based on deep learning approaches. Computers & Security, 81, 399-414.
14. Choi H., Lee H. Identifying botnets by capturing group activities in DNS traffic. Computer Networks, 56, 2012, pp. 20-33.
15. Choi H., Lee H., Lee H., Kim H. Botnet Detection by Monitoring Group Activities in DNS Traffic. Seventh IEEE International Conference on Computer and Information Technology (CIT 2007), 2007, pp. 715-720.
16. Conneau, A., Khandelwal, K., Goyal, N., Chaudhary, V., Wenzek, G., Guzmán, F., ... & Stoyanov, V. (2020). Unsupervised cross-lingual representation learning at scale. arXiv preprint arXiv:1911.02116.
17. Cresci, S., Di Pietro, R., Petrocchi, M., Spognardi, A., & Tesconi, M. (2017). The paradigm-shift of social spambots: Evidence, theories, and tools for the arms race. In Proceedings of the 26th international conference on world wide web companion (pp. 963-972).
18. Damballa. Botnet Detection for Communications Service Providers. https://www.damballa.com/downloads/r_pubs/WP_Botnet_Detection_for_CSPs.pdf.
19. Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2018). Bert: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805.

20. Dietrich C.J., Rossow C., Freiling F.C., Bos H., van Steen M., Pohlmann N. On Botnets that use DNS for Command and Control. Proceedings of European Conference on Computer Network Defense, 2011, pp. 9-16.
21. Doshi-Velez F., Kim B. Towards A Rigorous Science of Interpretable Machine Learning. arXiv preprint arXiv:1702.08608, 2017.
22. Farnham G., Atlasis A. Detecting DNS tunneling. SANS Institute InfoSec Reading Room, 2013, pp. 1-32.
23. Feng, F., He, X., Tang, J., & Liu, T. M. (2019). Graph adversarial training: Dynamically regularizing based on graph structure. IEEE Transactions on Knowledge and Data Engineering, 33(6), 2493-2504.
24. Feng, S., Wan, H., Wang, N., & Luo, M. (2021). BotGNN: A graph neural network based detection of social bots. Journal of Network and Computer Applications, 183, 103088.
25. Gilpin L.H., Bau D., Yuan B.Z., Bajwa A., Specter M., Kagal L. Explaining Explanations: An Overview of Interpretability of Machine Learning. In IEEE 5th International Conference on Data Science and Advanced Analytics (DSAA), 80-89, 2018.
26. Guo C., Berkhahn F. Entity Embeddings of Categorical Variables. arXiv preprint arXiv:1604.06737, 2016.
27. Guo, C., & Berkhahn, F. (2020). Entity embeddings of categorical variables. arXiv preprint arXiv:2003.03720.
28. He H., Bai Y., Garcia E.A., Li S. ADASYN: Adaptive Synthetic Sampling Approach for Imbalanced Learning. In IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence), 1322-1328, 2008.
29. Ichise H., Yong J., Iida K. Detection Method of DNS-based Botnet Communication Using Obtained NS Record History. Computer Software and Applications Conference (COMPSAC), 2015 IEEE 39th Annual, vol.3, 2015, pp. 676-677.

30. Izmailov P., Podoprikin D., Garipov T., Vetrov D., Wilson A.G. Averaging Weights Leads to Wider Optima and Better Generalization. Proceedings of the 34th Conference on Uncertainty in Artificial Intelligence (UAI), 2018.
31. Khan, S., Naseer, M., Hayat, M., Zamir, S. W., Khan, F. S., & Shah, M. (2022). Transformers in vision: A survey. *ACM computing surveys*, 54(10s), 1-41.
32. Kim, Y. (2014). Convolutional neural networks for sentence classification. *arXiv preprint arXiv:1408.5882*.
33. Kouvela M., Dimitriadis Y., Karlsson G. Social Bot Detection: Challenges and Implications for Privacy and Security. *Frontiers in Big Data*, 2022 Mar 23;5:805676.
34. Kudugunta, S., & Ferrara, E. (2018). Deep neural networks for bot detection. *Information Sciences*, 467, 312-322.
35. Li Z., Liu F., Yang W., Peng S., Zhou J. A Survey of Convolutional Neural Networks: Analysis, Applications, and Prospects. *IEEE Transactions on Neural Networks and Learning Systems*, 2021.
36. Lin T.Y., Goyal P., Girshick R., He K., Dollár P. Focal Loss for Dense Object Detection. *Proceedings of the IEEE International Conference on Computer Vision*, 2980-2988, 2017.
37. Lin, Z., Feng, M., Santos, C. N. D., Yu, M., Xiang, B., Zhou, B., & Bengio, Y. (2017). A structured self-attentive sentence embedding. *arXiv preprint arXiv:1703.03130*.
38. Liu Y., Ott M., Goyal N., Du J., Joshi M., Chen D., Levy O., Lewis M., Zettlemoyer L., Stoyanov V. RoBERTa: A Robustly Optimized BERT Pretraining Approach. *arXiv preprint arXiv:1907.11692*, 2019.
39. Liu, Y., Wu, Y., Guo, S., & Yang, Y. (2021). MGTab: A multi-modal graph-temporal attention benchmark for bot detection. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management* (pp. 4559-4568).

40. Loshchilov I., Hutter F. Decoupled Weight Decay Regularization. Proceedings of the International Conference on Learning Representations (ICLR), 2019.
41. Lu, J., Batra, D., Parikh, D., & Lee, S. (2019). Vilbert: Pretraining task-agnostic visiolinguistic representations for vision-and-language tasks. In Advances in neural information processing systems (pp. 13-23).
42. Lundberg S.M., Lee S.I. A Unified Approach to Interpreting Model Predictions. Advances in Neural Information Processing Systems 30, 4765-4774, 2017.
43. Lysenko S., Pomorova O., Savenko O., Kryshchuk A., Bobrovnikova K. DNS-based Anti-evasion Technique for Botnets Detection. Proceedings of the 2015 IEEE 8th International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS), IDAACS'2015, Warsaw, Poland, Vol.1, 2015, pp. 453-458.
44. Morales J.A., Al-Bataineh A., Xu S., Sandhu R. Analyzing DNS activities of bot processes. Proceedings of the 4th International Conference on Malicious and Unwanted Software (MALWARE), 2009, pp. 98-103.
45. Müller R., Kornblith S., Hinton G. When Does Label Smoothing Help? Advances in Neural Information Processing Systems 32, 4694-4703, 2019.
46. Perez E., Strub F., de Vries H., Dumoulin V., Courville A. FiLM: Visual Reasoning with a General Conditioning Layer. In Proceedings of the AAAI Conference on Artificial Intelligence, Vol. 32, No. 1, 2018.
47. Pomorova O., Savenko O., Lysenko S., Kryshchuk A., Bobrovnikova K. A Technique for the Botnet Detection Based on DNS-Traffic Analysis. Communications in Computer and Information Science, Vol. 522, 2015, pp. 127-138.
48. Ribeiro M.T., Singh S., Guestrin C. "Why Should I Trust You?": Explaining the Predictions of Any Classifier. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 1135-1144, 2016.

49. Shu, K., Mahudeswaran, D., Wang, S., Lee, D., & Liu, H. (2020). FakeNewsNet: A data repository with news content, social context, and spatiotemporal information for studying fake news on social media. *Big Data*, 8(3), 171-188.
50. Solaiman I., Brundage M., Clark J., Askill A., Herbert-Voss A., Wu J., Radford A., Wang J. Release Strategies and the Social Impacts of Language Models. *arXiv preprint arXiv:2112.04359*, 2021.
51. Vaswani A., Shazeer N., Parmar N., Uszkoreit J., Jones L., Gomez A.N., Kaiser L., Polosukhin I. Attention is All You Need. *Advances in Neural Information Processing Systems*, 5998-6008, 2017.
52. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., & Stoyanov, V. (2019). RoBERTa: A Robustly Optimized BERT Pretraining Approach. *arXiv preprint arXiv:1907.11692*. <https://arxiv.org/abs/1907.11692>
53. Zhou, X., & Zafarani, R. (2020). A survey of fake news detection: Methods, data, and opportunities. *ACM Computing Surveys (CSUR)*, 53(5), 1–40. <https://doi.org/10.1145/3395046>
54. Villamarín-Salomon R., Brustoloni J.C. Identifying Botnets Using Anomaly Detection Techniques Applied to DNS Traffic. *Consumer Communications and Networking Conference*, 2008, pp. 476-481.
55. Wang Y., Liu Y., Guo L., Chen C. DeepSoar: A Multiplex Graph Neural Network Framework for Social Bot Detection. *IEEE Transactions on Knowledge and Data Engineering*, 2021.
56. Wei, F., & Nguyen, U. T. (2019). Twitter bot detection using bidirectional long short-term memory neural networks and word embeddings. In *2019 IEEE International Conference on Information Reuse and Integration (IRI)* (pp. 389-394). IEEE.
57. Wu, Y., Yang, Y., Nakov, P., & Hovy, E. (2021). BERT-based multi-head selection framework for bot detection in social media. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management* (pp. 2117-2120).

58. Yang, Z., Zhang, J., Chang, E. C., & Liang, Z. (2019). Neural network based bot detection system with dynamic analysis. In Proceedings of the 2019 ACM SIGSAC Conference on Computer and Communications Security (pp. 2451-2453).
59. Zhang H., Cisse M., Dauphin Y.N., Lopez-Paz D. mixup: Beyond Empirical Risk Minimization. Proceedings of the International Conference on Learning Representations (ICLR), 2018.
60. Zhou P., Qi Z., Zheng S., Xu J., Bao H., Xu B. Text Classification Improved by Integrating Bidirectional LSTM with Two-dimensional Max Pooling. Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics: Technical Papers, 3485-3495, 2016.
61. Zhou, X., & Zafarani, R. (2020). A survey of fake news: Fundamental theories, detection methods, and opportunities. ACM Computing Surveys (CSUR), 53(5), 1-40.

ДОДАТКИ

Додаток А. Лістинг програмного коду

A.1. Архітектура мультимодальної трансформер-моделі

A.1.1. Основна архітектура моделі

```
import torch
import torch.nn as nn
import torch.nn.functional as F
from transformers import RobertaModel, RobertaConfig
import numpy as np
from typing import Dict, List, Optional, Tuple

class MultimodalBotDetector(nn.Module):
    """
    Мультимодальна трансформер-модель для виявлення ботів
    Інтегрує текстові, часові та метадані для класифікації
    """

    def __init__(self, config: Dict):
        super(MultimodalBotDetector, self).__init__()

        # Конфігурація моделі
        self.hidden_dim = config.get('hidden_dim', 768)
        self.num_attention_heads = config.get('num_attention_heads', 12)
        self.temporal_dim = config.get('temporal_dim', 128)
        self.metadata_dim = config.get('metadata_dim', 64)
        self.dropout_rate = config.get('dropout_rate', 0.1)

        # Текстовий енкодер (RoBERTa)
        self.text_encoder = RobertaModel.from_pretrained('roberta-base')

        # Часовий енкодер
        self.temporal_encoder = TemporalEncoder(
            input_dim=config.get('temporal_features', 48),
            hidden_dim=self.temporal_dim,
            num_heads=4
        )

        # Енкодер метаданих
        self.metadata_encoder = MetadataEncoder(
            input_dim=config.get('metadata_features', 98),
            hidden_dim=self.metadata_dim
        )

        # Cross-modal attention
        self.cross_modal_attention = CrossModalAttention(
            text_dim=self.hidden_dim,
            temporal_dim=self.temporal_dim,
            metadata_dim=self.metadata_dim,
            num_heads=self.num_attention_heads
        )
```

```

# Fusion layer
self.fusion_layer = FusionLayer(
    text_dim=self.hidden_dim,
    temporal_dim=self.temporal_dim,
    metadata_dim=self.metadata_dim,
    output_dim=self.hidden_dim
)

# Класифікатор
self.classifier = nn.Sequential(
    nn.Dropout(self.dropout_rate),
    nn.Linear(self.hidden_dim, self.hidden_dim // 2),
    nn.ReLU(),
    nn.Dropout(self.dropout_rate),
    nn.Linear(self.hidden_dim // 2, 1),
    nn.Sigmoid()
)

# Ініціалізація wag
self._init_weights()

def _init_weights(self):
    """Ініціалізація wag моделі"""
    for module in self.modules():
        if isinstance(module, nn.Linear):
            nn.init.xavier_uniform_(module.weight)
            if module.bias is not None:
                nn.init.constant_(module.bias, 0)

def forward(self, batch: Dict[str, torch.Tensor]) -> Dict[str, torch.Tensor]:
    """
    Прямий прохід через модель

    Args:
        batch: Словник з вхідними даними
            - text_ids: токени тексту [batch_size, seq_len]
            - attention_mask: маска уваги [batch_size, seq_len]
            - temporal_features: часові ознаки [batch_size, temp_len, temp_dim]
            - metadata_features: метадані [batch_size, meta_dim]

    Returns:
        Словник з результатами класифікації та проміжними представленнями
    """

    # Кодування тексту
    text_output = self.text_encoder(
        input_ids=batch['text_ids'],
        attention_mask=batch['attention_mask']
    )
    text_features = text_output.pooler_output # [batch_size, hidden_dim]

    # Кодування часових даних
    temporal_features = self.temporal_encoder(
        batch['temporal_features']
    ) # [batch_size, temporal_dim]

```

```

# Кодування метаданих
metadata_features = self.metadata_encoder(
    batch['metadata_features']
) # [batch_size, metadata_dim]

# Cross-modal attention
attended_features = self.cross_modal_attention(
    text_features, temporal_features, metadata_features
)

# Об'єднання модальностей
fused_features = self.fusion_layer(
    attended_features['text'],
    attended_features['temporal'],
    attended_features['metadata']
) # [batch_size, hidden_dim]

# Класифікація
logits = self.classifier(fused_features) # [batch_size, 1]

return {
    'logits': logits,
    'probabilities': logits, # Sigmoid вже застосовано
    'text_features': text_features,
    'temporal_features': temporal_features,
    'metadata_features': metadata_features,
    'fused_features': fused_features,
    'attention_weights': attended_features.get('attention_weights', None)
}

```

A.1.2. Часовий енкодер з механізмом уваги

class TemporalEncoder(nn.Module):

"""

Енкодер для обробки часових патернів активності
Використовує багатомасштабний аналіз з механізмом уваги
"""

```

def __init__(self, input_dim: int, hidden_dim: int, num_heads: int = 4):
    super(TemporalEncoder, self).__init__()

```

```

    self.input_dim = input_dim
    self.hidden_dim = hidden_dim
    self.num_heads = num_heads

```

```

# Проекційні шари для різних часових масштабів
self.short_term_proj = nn.Linear(input_dim, hidden_dim // 4)
self.medium_term_proj = nn.Linear(input_dim, hidden_dim // 4)
self.long_term_proj = nn.Linear(input_dim, hidden_dim // 4)
self.cyclic_proj = nn.Linear(input_dim, hidden_dim // 4)

```

```

# Temporal attention
self.temporal_attention = TemporalAttention(
    hidden_dim=hidden_dim,
    num_heads=num_heads
)

```

```

# Позиційне кодування
self.positional_encoding = PositionalEncoding(hidden_dim)

```

```

# Нормалізація та dropout
self.layer_norm = nn.LayerNorm(hidden_dim)
self.dropout = nn.Dropout(0.1)

def forward(self, temporal_data: torch.Tensor) -> torch.Tensor:
    """
    Args:
        temporal_data: [batch_size, seq_len, input_dim]
    Returns:
        Encoded temporal features: [batch_size, hidden_dim]
    """
    batch_size, seq_len, _ = temporal_data.shape

    # Обробка різних часових масштабів
    short_term = self.short_term_proj(temporal_data) # Хвилини-години
    medium_term = self.medium_term_proj(temporal_data) # Години-дні
    long_term = self.long_term_proj(temporal_data) # Дні-тижні
    cyclic = self.cyclic_proj(temporal_data) # Циклічні патерни

    # Об'єднання масштабів
    multi_scale_features = torch.cat([
        short_term, medium_term, long_term, cyclic
    ], dim=-1) # [batch_size, seq_len, hidden_dim]

    # Додавання позиційного кодування
    multi_scale_features = self.positional_encoding(multi_scale_features)

    # Temporal attention
    attended_features = self.temporal_attention(multi_scale_features)

    # Глобальне агрегування (mean pooling з увагою)
    attention_weights = F.softmax(
        torch.sum(attended_features, dim=-1), dim=-1
    ) # [batch_size, seq_len]

    temporal_embedding = torch.sum(
        attended_features * attention_weights.unsqueeze(-1), dim=1
    ) # [batch_size, hidden_dim]

    # Нормалізація та dropout
    temporal_embedding = self.layer_norm(temporal_embedding)
    temporal_embedding = self.dropout(temporal_embedding)

    return temporal_embedding

class TemporalAttention(nn.Module):
    """
    Механізм уваги для аналізу часових залежностей
    """

    def __init__(self, hidden_dim: int, num_heads: int):
        super(TemporalAttention, self).__init__()

        self.hidden_dim = hidden_dim
        self.num_heads = num_heads

```

```

self.head_dim = hidden_dim // num_heads

assert hidden_dim % num_heads == 0

# Проекційні матриці
self.query_proj = nn.Linear(hidden_dim, hidden_dim)
self.key_proj = nn.Linear(hidden_dim, hidden_dim)
self.value_proj = nn.Linear(hidden_dim, hidden_dim)
self.output_proj = nn.Linear(hidden_dim, hidden_dim)

# Відносне часове кодування
self.relative_time_embedding = nn.Embedding(512, self.head_dim)

self.dropout = nn.Dropout(0.1)

def forward(self, x: torch.Tensor) -> torch.Tensor:
    """
    Args:
        x: [batch_size, seq_len, hidden_dim]
    Returns:
        Attended features: [batch_size, seq_len, hidden_dim]
    """
    batch_size, seq_len, _ = x.shape

    # Проекції для Q, K, V
    Q = self.query_proj(x).view(batch_size, seq_len, self.num_heads, self.head_dim)
    K = self.key_proj(x).view(batch_size, seq_len, self.num_heads, self.head_dim)
    V = self.value_proj(x).view(batch_size, seq_len, self.num_heads, self.head_dim)

    # Транспозиція для багатоголової уваги
    Q = Q.transpose(1, 2) # [batch_size, num_heads, seq_len, head_dim]
    K = K.transpose(1, 2)
    V = V.transpose(1, 2)

    # Обчислення уваги з відносним часовим кодуванням
    attention_scores = torch.matmul(Q, K.transpose(-2, -1)) / np.sqrt(self.head_dim)

    # Додавання відносного часового кодування
    relative_positions = self._get_relative_positions(seq_len)
    relative_embeddings = self.relative_time_embedding(relative_positions)
    relative_scores = torch.matmul(Q, relative_embeddings.transpose(-2, -1))

    # Комбінування стандартної та відносної уваги
    attention_scores = attention_scores + relative_scores

    # Softmax та dropout
    attention_weights = F.softmax(attention_scores, dim=-1)
    attention_weights = self.dropout(attention_weights)

    # Застосування уваги
    attended_values = torch.matmul(attention_weights, V)

    # Транспозиція назад та об'єднання голів
    attended_values = attended_values.transpose(1, 2).contiguous().view(
        batch_size, seq_len, self.hidden_dim
    )

```

```

# Фінальна проекція
output = self.output_proj(attended_values)

return output

def _get_relative_positions(self, seq_len: int) -> torch.Tensor:
    """Генерує матрицю відносних позицій"""
    positions = torch.arange(seq_len, dtype=torch.long)
    relative_positions = positions.unsqueeze(0) - positions.unsqueeze(1)

    # Обмеження діапазону для embeddings
    relative_positions = torch.clamp(relative_positions, -255, 255) + 255

    return relative_positions

```

A.1.3. Cross-modal attention механізм

```
class CrossModalAttention(nn.Module):
```

```
    """
```

```
    Механізм крос-модальної уваги для інтеграції різних типів даних
```

```
    """
```

```
    def __init__(self, text_dim: int, temporal_dim: int, metadata_dim: int, num_heads: int):
        super(CrossModalAttention, self).__init__()
```

```

        self.text_dim = text_dim
        self.temporal_dim = temporal_dim
        self.metadata_dim = metadata_dim
        self.num_heads = num_heads

```

```

# Проекція всіх модальностей до єдиної розмірності
self.common_dim = text_dim
self.temporal_projection = nn.Linear(temporal_dim, self.common_dim)
self.metadata_projection = nn.Linear(metadata_dim, self.common_dim)

```

```
# Multi-head attention для кожної пари модальностей
```

```

self.text_to_temporal = nn.MultiheadAttention(
    embed_dim=self.common_dim,
    num_heads=num_heads,
    dropout=0.1,
    batch_first=True
)

```

```

self.text_to_metadata = nn.MultiheadAttention(
    embed_dim=self.common_dim,
    num_heads=num_heads,
    dropout=0.1,
    batch_first=True
)

```

```

self.temporal_to_text = nn.MultiheadAttention(
    embed_dim=self.common_dim,
    num_heads=num_heads,
    dropout=0.1,
    batch_first=True
)

```

```

self.temporal_to_metadata = nn.MultiheadAttention(
    embed_dim=self.common_dim,

```

```

        num_heads=num_heads,
        dropout=0.1,
        batch_first=True
    )

    self.metadata_to_text = nn.MultiheadAttention(
        embed_dim=self.common_dim,
        num_heads=num_heads,
        dropout=0.1,
        batch_first=True
    )

    self.metadata_to_temporal = nn.MultiheadAttention(
        embed_dim=self.common_dim,
        num_heads=num_heads,
        dropout=0.1,
        batch_first=True
    )

    # Нормалізація
    self.text_norm = nn.LayerNorm(self.common_dim)
    self.temporal_norm = nn.LayerNorm(self.common_dim)
    self.metadata_norm = nn.LayerNorm(self.common_dim)

def forward(self, text_features: torch.Tensor,
            temporal_features: torch.Tensor,
            metadata_features: torch.Tensor) -> Dict[str, torch.Tensor]:
    """
    Args:
        text_features: [batch_size, text_dim]
        temporal_features: [batch_size, temporal_dim]
        metadata_features: [batch_size, metadata_dim]

    Returns:
        Dictionary with attended features for each modality
    """
    batch_size = text_features.shape[0]

    # Проекція до спільної розмірності
    temporal_proj = self.temporal_projection(temporal_features) # [batch_size, common_dim]
    metadata_proj = self.metadata_projection(metadata_features) # [batch_size, common_dim]

    # Додавання розмірності послідовності для attention
    text_seq = text_features.unsqueeze(1) # [batch_size, 1, common_dim]
    temporal_seq = temporal_proj.unsqueeze(1) # [batch_size, 1, common_dim]
    metadata_seq = metadata_proj.unsqueeze(1) # [batch_size, 1, common_dim]

    # Cross-modal attention між всіма парами модальностей

    # Text attending to other modalities
    text_attended_by_temporal, text_temp_weights = self.text_to_temporal(
        text_seq, temporal_seq, temporal_seq
    )
    text_attended_by_metadata, text_meta_weights = self.text_to_metadata(
        text_seq, metadata_seq, metadata_seq
    )

```

```

# Temporal attending to other modalities
temporal_attended_by_text, temp_text_weights = self.temporal_to_text(
    temporal_seq, text_seq, text_seq
)
temporal_attended_by_metadata, temp_meta_weights = self.temporal_to_metadata(
    temporal_seq, metadata_seq, metadata_seq
)

# Metadata attending to other modalities
metadata_attended_by_text, meta_text_weights = self.metadata_to_text(
    metadata_seq, text_seq, text_seq
)
metadata_attended_by_temporal, meta_temp_weights = self.metadata_to_temporal(
    metadata_seq, temporal_seq, temporal_seq
)

# Агрегація attended features
enhanced_text = text_features + text_attended_by_temporal.squeeze(1) +
text_attended_by_metadata.squeeze(1)
enhanced_temporal = temporal_proj + temporal_attended_by_text.squeeze(1) +
temporal_attended_by_metadata.squeeze(1)
enhanced_metadata = metadata_proj + metadata_attended_by_text.squeeze(1) +
metadata_attended_by_temporal.squeeze(1)

# Нормалізація
enhanced_text = self.text_norm(enhanced_text)
enhanced_temporal = self.temporal_norm(enhanced_temporal)
enhanced_metadata = self.metadata_norm(enhanced_metadata)

return {
    'text': enhanced_text,
    'temporal': enhanced_temporal,
    'metadata': enhanced_metadata,
    'attention_weights': {
        'text_temporal': text_temp_weights,
        'text_metadata': text_meta_weights,
        'temporal_text': temp_text_weights,
        'temporal_metadata': temp_meta_weights,
        'metadata_text': meta_text_weights,
        'metadata_temporal': meta_temp_weights
    }
}

```

A.2. Допоміжні компоненти

A.2.1. Енкодер метаданих

```
class MetadataEncoder(nn.Module):
```

```
    """
```

```
    Енкодер для обробки метаданих профілю користувача
```

```
    """
```

```
    def __init__(self, input_dim: int, hidden_dim: int):
        super(MetadataEncoder, self).__init__()
```

```
        self.input_dim = input_dim
        self.hidden_dim = hidden_dim
```

```
        # Мережа для обробки числових метаданих
        self.numerical_encoder = nn.Sequential(
```

```

        nn.Linear(input_dim // 2, hidden_dim),
        nn.ReLU(),
        nn.Dropout(0.1),
        nn.Linear(hidden_dim, hidden_dim // 2),
        nn.ReLU()
    )

    # Мережа для обробки категоріальних метаданих
    self.categorical_encoder = nn.Sequential(
        nn.Linear(input_dim // 2, hidden_dim),
        nn.ReLU(),
        nn.Dropout(0.1),
        nn.Linear(hidden_dim, hidden_dim // 2),
        nn.ReLU()
    )

    # Об'єднуюча мережа
    self.fusion_network = nn.Sequential(
        nn.Linear(hidden_dim, hidden_dim),
        nn.ReLU(),
        nn.Dropout(0.15),
        nn.Linear(hidden_dim, hidden_dim)
    )

    self.layer_norm = nn.LayerNorm(hidden_dim)

def forward(self, metadata: torch.Tensor) -> torch.Tensor:
    """
    Args:
        metadata: [batch_size, input_dim]
    Returns:
        Encoded metadata: [batch_size, hidden_dim]
    """
    # Розділення на числові та категоріальні ознаки
    numerical_features = metadata[:, :self.input_dim // 2]
    categorical_features = metadata[:, self.input_dim // 2:]

    # Кодування
    numerical_encoded = self.numerical_encoder(numerical_features)
    categorical_encoded = self.categorical_encoder(categorical_features)

    # Об'єднання
    combined_features = torch.cat([numerical_encoded, categorical_encoded], dim=-1)

    # Фінальна обробка
    metadata_embedding = self.fusion_network(combined_features)
    metadata_embedding = self.layer_norm(metadata_embedding)

    return metadata_embedding

```

```

class FusionLayer(nn.Module):
    """

```

```

    Шар для об'єднання представлень різних модальностей
    """

```

```

    def __init__(self, text_dim: int, temporal_dim: int, metadata_dim: int, output_dim: int):

```

```

super(FusionLayer, self).__init__()

self.text_dim = text_dim
self.temporal_dim = temporal_dim
self.metadata_dim = metadata_dim
self.output_dim = output_dim

# Adaptive weighting для кожної модальності
self.text_weight = nn.Parameter(torch.ones(1))
self.temporal_weight = nn.Parameter(torch.ones(1))
self.metadata_weight = nn.Parameter(torch.ones(1))

# Проекційні шари
self.text_projection = nn.Linear(text_dim, output_dim)
self.temporal_projection = nn.Linear(temporal_dim, output_dim)
self.metadata_projection = nn.Linear(metadata_dim, output_dim)

# FiLM (Feature-wise Linear Modulation) layers
self.film_text = FiMLayer(output_dim)
self.film_temporal = FiMLayer(output_dim)
self.film_metadata = FiMLayer(output_dim)

# Fusion network
self.fusion_network = nn.Sequential(
    nn.Linear(output_dim * 3, output_dim * 2),
    nn.ReLU(),
    nn.Dropout(0.2),
    nn.Linear(output_dim * 2, output_dim),
    nn.ReLU(),
    nn.LayerNorm(output_dim)
)

def forward(self, text_features: torch.Tensor,
            temporal_features: torch.Tensor,
            metadata_features: torch.Tensor) -> torch.Tensor:
    """
    Об'єднання мультимодальних представлень
    """
    # Проекція до спільної розмірності
    text_proj = self.text_projection(text_features)
    temporal_proj = self.temporal_projection(temporal_features)
    metadata_proj = self.metadata_projection(metadata_features)

    # FiLM модуляція для кожної модальності
    text_modulated = self.film_text(text_proj, temporal_proj, metadata_proj)
    temporal_modulated = self.film_temporal(temporal_proj, text_proj, metadata_proj)
    metadata_modulated = self.film_metadata(metadata_proj, text_proj, temporal_proj)

    # Adaptive weighting
    text_weighted = self.text_weight * text_modulated
    temporal_weighted = self.temporal_weight * temporal_modulated
    metadata_weighted = self.metadata_weight * metadata_modulated

    # Concatenation та fusion
    combined_features = torch.cat([
        text_weighted, temporal_weighted, metadata_weighted
    ], dim=-1)

```

```

fused_representation = self.fusion_network(combined_features)

return fused_representation

```

```

class FiLMLayer(nn.Module):
    """
    Feature-wise Linear Modulation layer
    """

    def __init__(self, feature_dim: int):
        super(FiLMLayer, self).__init__()

        self.feature_dim = feature_dim

        # Мережі для генерації параметрів модуляції
        self.gamma_network = nn.Sequential(
            nn.Linear(feature_dim * 2, feature_dim),
            nn.ReLU(),
            nn.Linear(feature_dim, feature_dim)
        )

        self.beta_network = nn.Sequential(
            nn.Linear(feature_dim * 2, feature_dim),
            nn.ReLU(),
            nn.Linear(feature_dim, feature_dim)
        )

    def forward(self, target_features: torch.Tensor,
                context_features1: torch.Tensor,
                context_features2: torch.Tensor) -> torch.Tensor:
        """
        Args:
            target_features: Features to be modulated
            context_features1, context_features2: Context from other modalities
        """
        # Об'єднання контекстних ознак
        context = torch.cat([context_features1, context_features2], dim=-1)

        # Генерація параметрів модуляції
        gamma = self.gamma_network(context) # Scaling parameter
        beta = self.beta_network(context) # Shifting parameter

        # FiLM transformation
        modulated_features = gamma * target_features + beta

        return modulated_features

```

```

class PositionalEncoding(nn.Module):
    """
    Позиційне кодування для часових послідовностей
    """

    def __init__(self, d_model: int, max_len: int = 1000):
        super(PositionalEncoding, self).__init__()

```

```

pe = torch.zeros(max_len, d_model)
position = torch.arange(0, max_len, dtype=torch.float).unsqueeze(1)

div_term = torch.exp(torch.arange(0, d_model, 2).float() *
                      (-np.log(10000.0) / d_model))

pe[:, 0::2] = torch.sin(position * div_term)
pe[:, 1::2] = torch.cos(position * div_term)

self.register_buffer('pe', pe.unsqueeze(0))

def forward(self, x: torch.Tensor) -> torch.Tensor:
    """
    Args:
        x: [batch_size, seq_len, d_model]
    """
    seq_len = x.size(1)
    return x + self.pe[:, :seq_len, :]

```

A.3. Функції втрат та метрики

A.3.1. Спеціалізована функція втрат

```

import torch.nn.functional as F
from sklearn.metrics import f1_score, precision_score, recall_score, roc_auc_score

```

```

class FocalLoss(nn.Module):
    """
    Focal Loss для роботи з дисбалансованими класами
    """

    def __init__(self, alpha: float = 0.25, gamma: float = 2.0, reduction: str = 'mean'):
        super(FocalLoss, self).__init__()
        self.alpha = alpha
        self.gamma = gamma
        self.reduction = reduction

    def forward(self, inputs: torch.Tensor, targets: torch.Tensor) -> torch.Tensor:
        """
        Args:
            inputs: [batch_size, 1] - логіти або ймовірності
            targets: [batch_size, 1] - справжні мітки (0 або 1)
        """
        # Переконаємося, що inputs - це ймовірності
        if inputs.max() > 1.0 or inputs.min() < 0.0:
            inputs = torch.sigmoid(inputs)

        # Обчислення стандартної cross entropy
        ce_loss = F.binary_cross_entropy(inputs, targets, reduction='none')

        # Обчислення pt
        pt = torch.where(targets == 1, inputs, 1 - inputs)

        # Focal weight
        focal_weight = (1 - pt) ** self.gamma

        # Alpha weighting
        alpha_weight = torch.where(targets == 1, self.alpha, 1 - self.alpha)

```

```

# Focal loss
focal_loss = alpha_weight * focal_weight * ce_loss

if self.reduction == 'mean':
    return focal_loss.mean()
elif self.reduction == 'sum':
    return focal_loss.sum()
else:
    return focal_loss

```

```

class CombinedLoss(nn.Module):

```

```

    """

```

```

    Комбінована функція втрат (Focal + L2 регуляризація)

```

```

    """

```

```

def __init__(self, focal_weight: float = 0.8, l2_weight: float = 0.01,
              alpha: float = 0.25, gamma: float = 2.0):
    super(CombinedLoss, self).__init__()

```

```

    self.focal_loss = FocalLoss(alpha=alpha, gamma=gamma)
    self.focal_weight = focal_weight
    self.l2_weight = l2_weight

```

```

def forward(self, outputs: Dict[str, torch.Tensor],
            targets: torch.Tensor, model: nn.Module) -> torch.Tensor:

```

```

    """

```

```

    Args:

```

```

        outputs: Результати моделі
        targets: Справжні мітки
        model: Модель для L2 регуляризації

```

```

    """

```

```

    # Focal loss

```

```

    focal_loss = self.focal_loss(outputs['probabilities'], targets)

```

```

    # L2 regularization

```

```

    l2_reg = 0

```

```

    for param in model.parameters():

```

```

        l2_reg += torch.norm(param) ** 2

```

```

    # Комбінована втрата

```

```

    total_loss = self.focal_weight * focal_loss + self.l2_weight * l2_reg

```

```

    return total_loss

```

```

def compute_metrics(predictions: torch.Tensor, targets: torch.Tensor,
                    threshold: float = 0.5) -> Dict[str, float]:

```

```

    """

```

```

    Обчислення метрик класифікації

```

```

    Args:

```

```

        predictions: [batch_size, 1] - передбачені ймовірності
        targets: [batch_size, 1] - справжні мітки
        threshold: Поріг для бінарної класифікації

```

```

    Returns:

```

```

        Словник з метриками
        """
        # Конвертація в numpy
        preds_np = predictions.detach().cpu().numpy().flatten()
        targets_np = targets.detach().cpu().numpy().flatten()

        # Бінарні передбачення
        binary_preds = (preds_np >= threshold).astype(int)

        # Обчислення метрик
        metrics = {
            'precision': precision_score(targets_np, binary_preds, zero_division=0),
            'recall': recall_score(targets_np, binary_preds, zero_division=0),
            'f1_score': f1_score(targets_np, binary_preds, zero_division=0),
            'accuracy': (binary_preds == targets_np).mean(),
            'auc_roc': roc_auc_score(targets_np, preds_np) if len(np.unique(targets_np)) > 1 else 0.5,
        }

        # Matthews Correlation Coefficient
        tp = np.sum((binary_preds == 1) & (targets_np == 1))
        tn = np.sum((binary_preds == 0) & (targets_np == 0))
        fp = np.sum((binary_preds == 1) & (targets_np == 0))
        fn = np.sum((binary_preds == 0) & (targets_np == 1))

        denominator = np.sqrt((tp + fp) * (tp + fn) * (tn + fp) * (tn + fn))
        if denominator == 0:
            metrics['mcc'] = 0
        else:
            metrics['mcc'] = (tp * tn - fp * fn) / denominator

        return metrics

```

A.4. Навчання моделі

A.4.1. Основний цикл навчання

```

import torch.optim as optim
from torch.utils.data import DataLoader
from transformers import get_cosine_schedule_with_warmup
import wandb
from tqdm import tqdm

class BotDetectorTrainer:
    """
    Клас для навчання моделі виявлення ботів
    """

    def __init__(self, model: MultimodalBotDetector, config: Dict):
        self.model = model
        self.config = config

        # Оптимізатор
        self.optimizer = optim.AdamW(
            model.parameters(),
            lr=config.get('learning_rate', 2e-5),
            weight_decay=config.get('weight_decay', 0.01),
            betas=(0.9, 0.999)
        )

        # Функція втрат

```

```

self.criterion = CombinedLoss(
    focal_weight=config.get('focal_weight', 0.8),
    l2_weight=config.get('l2_weight', 0.01),
    alpha=config.get('focal_alpha', 0.25),
    gamma=config.get('focal_gamma', 2.0)
)

# Пристрій для обчислень
self.device = torch.device('cuda' if torch.cuda.is_available() else 'cpu')
self.model.to(self.device)

# Лічильники
self.current_epoch = 0
self.best_f1 = 0.0
self.patience_counter = 0

# Історія навчання
self.train_history = []
self.val_history = []

def train_epoch(self, train_loader: DataLoader) -> Dict[str, float]:
    """
    Навчання протягом однієї епохи
    """
    self.model.train()

    total_loss = 0.0
    all_predictions = []
    all_targets = []

    progress_bar = tqdm(train_loader, desc=f'Epoch {self.current_epoch + 1}')

    for batch_idx, batch in enumerate(progress_bar):
        # Переміщення на GPU
        batch = {k: v.to(self.device) if torch.is_tensor(v) else v
                  for k, v in batch.items()}

        targets = batch.pop('labels').float()

        # Обнулення градієнтів
        self.optimizer.zero_grad()

        # Прямий прохід
        outputs = self.model(batch)

        # Обчислення втрат
        loss = self.criterion(outputs, targets, self.model)

        # Зворотний прохід
        loss.backward()

        # Обрізання градієнтів
        torch.nn.utils.clip_grad_norm_(self.model.parameters(), max_norm=1.0)

        # Оновлення ваг
        self.optimizer.step()

```

```

# Збір статистики
total_loss += loss.item()
all_predictions.extend(outputs['probabilities'].detach().cpu().numpy())
all_targets.extend(targets.detach().cpu().numpy())

# Оновлення прогрес-бару
progress_bar.set_postfix({
    'loss': f'{loss.item():.4f}',
    'avg_loss': f'{total_loss / (batch_idx + 1):.4f}'
})

# Обчислення метрик епохи
epoch_metrics = compute_metrics(
    torch.tensor(all_predictions),
    torch.tensor(all_targets)
)
epoch_metrics['loss'] = total_loss / len(train_loader)

return epoch_metrics

def validate_epoch(self, val_loader: DataLoader) -> Dict[str, float]:
    """
    Валідація моделі
    """
    self.model.eval()

    total_loss = 0.0
    all_predictions = []
    all_targets = []

    with torch.no_grad():
        for batch in tqdm(val_loader, desc='Validation'):
            # Переміщення на GPU
            batch = {k: v.to(self.device) if torch.is_tensor(v) else v
                     for k, v in batch.items()}

            targets = batch.pop('labels').float()

            # Прямий прохід
            outputs = self.model(batch)

            # Обчислення втрат
            loss = self.criterion(outputs, targets, self.model)

            # Збір статистики
            total_loss += loss.item()
            all_predictions.extend(outputs['probabilities'].detach().cpu().numpy())
            all_targets.extend(targets.detach().cpu().numpy())

    # Обчислення метрик валідації
    val_metrics = compute_metrics(
        torch.tensor(all_predictions),
        torch.tensor(all_targets)
    )
    val_metrics['loss'] = total_loss / len(val_loader)

    return val_metrics

```

```

def train(self, train_loader: DataLoader, val_loader: DataLoader,
        num_epochs: int = 10, early_stopping_patience: int = 3):
    """
    Основний цикл навчання
    """
    # Налаштування планувальника
    total_steps = len(train_loader) * num_epochs
    warmup_steps = int(0.1 * total_steps)

    scheduler = get_cosine_schedule_with_warmup(
        self.optimizer,
        num_warmup_steps=warmup_steps,
        num_training_steps=total_steps
    )

    print(f"Початок навчання на {num_epochs} епох")
    print(f"Загальна кількість кроків: {total_steps}")
    print(f"Кроки розігріву: {warmup_steps}")

    for epoch in range(num_epochs):
        self.current_epoch = epoch

        # Навчання
        train_metrics = self.train_epoch(train_loader)
        self.train_history.append(train_metrics)

        # Валідація
        val_metrics = self.validate_epoch(val_loader)
        self.val_history.append(val_metrics)

        # Оновлення планувальника
        scheduler.step()

        # Логування
        print(f"\nЕпоха {epoch + 1}/{num_epochs}")
        print(f"Навчання - Loss: {train_metrics['loss']:.4f}, "
              f"F1: {train_metrics['f1_score']:.4f}, "
              f"Precision: {train_metrics['precision']:.4f}, "
              f"Recall: {train_metrics['recall']:.4f}")
        print(f"Валідація - Loss: {val_metrics['loss']:.4f}, "
              f"F1: {val_metrics['f1_score']:.4f}, "
              f"Precision: {val_metrics['precision']:.4f}, "
              f"Recall: {val_metrics['recall']:.4f}")

        # Логування в wandb (якщо налаштовано)
        if wandb.run:
            wandb.log({
                'epoch': epoch + 1,
                'train_loss': train_metrics['loss'],
                'train_f1': train_metrics['f1_score'],
                'val_loss': val_metrics['loss'],
                'val_f1': val_metrics['f1_score'],
                'learning_rate': scheduler.get_last_lr()[0]
            })

    # Early stopping

```

```

if val_metrics['f1_score'] > self.best_f1:
    self.best_f1 = val_metrics['f1_score']
    self.patience_counter = 0

    # Збереження найкращої моделі
    self.save_checkpoint('best_model.pth', val_metrics)
    print(f'Нова найкраща модель збережена! F1: {self.best_f1:.4f}')

else:
    self.patience_counter += 1
    print(f'Patience: {self.patience_counter}/{early_stopping_patience}')

    if self.patience_counter >= early_stopping_patience:
        print("Раннє зупинення навчання")
        break

print(f'\nНавчання завершено! Найкраща F1-міра: {self.best_f1:.4f}')

def save_checkpoint(self, filepath: str, metrics: Dict[str, float]):
    """
    Збереження чекпойнту моделі
    """
    checkpoint = {
        'epoch': self.current_epoch,
        'model_state_dict': self.model.state_dict(),
        'optimizer_state_dict': self.optimizer.state_dict(),
        'best_f1': self.best_f1,
        'metrics': metrics,
        'config': self.config
    }

    torch.save(checkpoint, filepath)

def load_checkpoint(self, filepath: str):
    """
    Завантаження чекпойнту моделі
    """
    checkpoint = torch.load(filepath, map_location=self.device)

    self.model.load_state_dict(checkpoint['model_state_dict'])
    self.optimizer.load_state_dict(checkpoint['optimizer_state_dict'])
    self.current_epoch = checkpoint['epoch']
    self.best_f1 = checkpoint['best_f1']

    print(f'Чекпойнт завантажено з епохи {self.current_epoch}')
    print(f'Найкраща F1-міра: {self.best_f1:.4f}')

```

A.5. Обробка та завантаження даних

A.5.1. Dataset класи

```

from torch.utils.data import Dataset
import pandas as pd
import json
from transformers import RobertaTokenizer
from sklearn.preprocessing import StandardScaler, LabelEncoder
import pickle

```

```

class BotDetectionDataset(Dataset):
    """

```

Dataset для завантаження даних виявлення ботів
"""

```
def __init__(self, data_path: str, tokenizer: RobertaTokenizer,
              max_length: int = 512, is_training: bool = True):
```

```
    self.data_path = data_path
    self.tokenizer = tokenizer
    self.max_length = max_length
    self.is_training = is_training
```

```
    # Завантаження даних
    self.data = self._load_data()
```

```
    # Препроцесинг
    if is_training:
        self._preprocess_data()
        self._save_preprocessors()
    else:
        self._load_preprocessors()
        self._preprocess_data()
```

```
def _load_data(self) -> pd.DataFrame:
```

```
    """
    Завантаження даних з файлу
    """
```

```
    if self.data_path.endswith('.csv'):
        data = pd.read_csv(self.data_path)
    elif self.data_path.endswith('.json'):
        data = pd.read_json(self.data_path)
    else:
        raise ValueError(f"Непідтримуваний формат файлу: {self.data_path}")

    return data
```

```
def _preprocess_data(self):
```

```
    """
    Попередня обробка даних
    """
```

```
    # Обробка текстових даних
    self.data['text'] = self.data['text'].fillna("")
    self.data['text'] = self.data['text'].apply(self._clean_text)
```

```
    # Обробка часових даних
    self._process_temporal_features()
```

```
    # Обробка метаданих
    self._process_metadata()
```

```
    # Видалення рядків з відсутніми мітками (тільки для навчання)
    if self.is_training and 'label' in self.data.columns:
        self.data = self.data.dropna(subset=['label'])
```

```
def _clean_text(self, text: str) -> str:
```

```
    """
    Очищення тексту
    """
```

```

import re

# Видалення HTML тегів
text = re.sub(r'<[^>]+>', '', text)

# Нормалізація URL
text = re.sub(r'http[s]?://(?:[a-zA-Z]|[0-9]|[$-_@.&+]|[*%\(\)\,]|(?:%[0-9a-fA-F][0-9a-fA-F]))+',
              '[URL]', text)

# Нормалізація згадок
text = re.sub(r'@\w+', '[MENTION]', text)

# Нормалізація хештегів
text = re.sub(r'#\w+', '[HASHTAG]', text)

# Видалення зайвих пробілів
text = re.sub(r'\s+', ' ', text).strip()

return text

def _process_temporal_features(self):
    """
    Обробка часових ознак
    """
    # Створення часових ознак з timestamps
    if 'timestamps' in self.data.columns:
        timestamps = self.data['timestamps'].apply(
            lambda x: json.loads(x) if isinstance(x, str) else x
        )

        # Витягування статистичних ознак
        self.data['posting_frequency'] = timestamps.apply(
            lambda x: len(x) / max(1, (max(x) - min(x)) / 3600) if len(x) > 1 else 0
        )

        self.data['posting_regularity'] = timestamps.apply(
            self._calculate_regularity
        )

        self.data['night_activity'] = timestamps.apply(
            lambda x: sum(1 for t in x if 22 <= (t % 86400) // 3600 <= 6) / len(x) if x else 0
        )

        # Додаткові часові ознаки
        self.data['burst_activity'] = timestamps.apply(self._detect_bursts)
        self.data['weekday_activity'] = timestamps.apply(self._weekday_pattern)

def _calculate_regularity(self, timestamps: List[int]) -> float:
    """
    Обчислення регулярності активності
    """
    if len(timestamps) < 2:
        return 0.0

    intervals = np.diff(sorted(timestamps))
    if len(intervals) == 0:

```

```

        return 0.0

    # Коефіцієнт варіації інтервалів
    mean_interval = np.mean(intervals)
    std_interval = np.std(intervals)

    if mean_interval == 0:
        return 0.0

    cv = std_interval / mean_interval
    regularity = 1 / (1 + cv) # Нормалізація до [0, 1]

    return regularity

def _detect_bursts(self, timestamps: List[int]) -> float:
    """
    Виявлення бурстової активності
    """
    if len(timestamps) < 3:
        return 0.0

    # Розділення на вікна по 1 годині
    hour_counts = {}
    for ts in timestamps:
        hour = ts // 3600
        hour_counts[hour] = hour_counts.get(hour, 0) + 1

    counts = list(hour_counts.values())
    if not counts:
        return 0.0

    # Виявлення аномально високої активності
    mean_count = np.mean(counts)
    std_count = np.std(counts)

    if std_count == 0:
        return 0.0

    # Частка годин з активністю > mean + 2*std
    threshold = mean_count + 2 * std_count
    burst_ratio = sum(1 for c in counts if c > threshold) / len(counts)

    return burst_ratio

def _weekday_pattern(self, timestamps: List[int]) -> float:
    """
    Аналіз патернів активності по днях тижня
    """
    if not timestamps:
        return 0.0

    # Розподіл по днях тижня
    weekday_counts = [0] * 7
    for ts in timestamps:
        weekday = (ts // 86400) % 7
        weekday_counts[weekday] += 1

```

```

total_posts = sum(weekday_counts)
if total_posts == 0:
    return 0.0

# Рівномірність розподілу (ентропія)
probabilities = [c / total_posts for c in weekday_counts if c > 0]
entropy = -sum(p * np.log2(p) for p in probabilities)

# Нормалізація ентропії
max_entropy = np.log2(7) # Максимальна ентропія для 7 днів
normalized_entropy = entropy / max_entropy if max_entropy > 0 else 0

return normalized_entropy

def _process_metadata(self):
    """
    Обробка метаданих профілю
    """
    # Заповнення відсутніх значень
    metadata_columns = [
        'followers_count', 'friends_count', 'statuses_count',
        'account_age_days', 'profile_has_image', 'verified',
        'description_length', 'location_specified'
    ]

    for col in metadata_columns:
        if col in self.data.columns:
            if col in ['profile_has_image', 'verified', 'location_specified']:
                self.data[col] = self.data[col].fillna(0).astype(int)
            else:
                self.data[col] = self.data[col].fillna(0).astype(float)

    # Створення додаткових ознак
    if 'followers_count' in self.data.columns and 'friends_count' in self.data.columns:
        self.data['followers_friends_ratio'] = self.data.apply(
            lambda row: row['followers_count'] / max(1, row['friends_count']), axis=1
        )

    if 'statuses_count' in self.data.columns and 'account_age_days' in self.data.columns:
        self.data['posts_per_day'] = self.data.apply(
            lambda row: row['statuses_count'] / max(1, row['account_age_days']), axis=1
        )

    # Нормалізація числових ознак
    if self.is_training:
        self.metadata_scaler = StandardScaler()
        numeric_columns = self.data.select_dtypes(include=[np.number]).columns
        metadata_numeric = [col for col in numeric_columns if col not in ['label']]

        if metadata_numeric:
            self.data[metadata_numeric] = self.metadata_scaler.fit_transform(
                self.data[metadata_numeric]
            )
    else:
        if hasattr(self, 'metadata_scaler'):
            numeric_columns = self.data.select_dtypes(include=[np.number]).columns
            metadata_numeric = [col for col in numeric_columns if col not in ['label']]

```

```

        if metadata_numeric:
            self.data[metadata_numeric] = self.metadata_scaler.transform(
                self.data[metadata_numeric]
            )

def _save_preprocessors(self):
    """
    Збереження препроцесорів для використання на тестових даних
    """
    if hasattr(self, 'metadata_scaler'):
        with open('metadata_scaler.pkl', 'wb') as f:
            pickle.dump(self.metadata_scaler, f)

def _load_preprocessors(self):
    """
    Завантаження збережених препроцесорів
    """
    try:
        with open('metadata_scaler.pkl', 'rb') as f:
            self.metadata_scaler = pickle.load(f)
    except FileNotFoundError:
        print("Попередження: Препроцесори не знайдені")

def __len__(self) -> int:
    return len(self.data)

def __getitem__(self, idx: int) -> Dict[str, torch.Tensor]:
    """
    Отримання одного зразка даних
    """
    row = self.data.iloc[idx]

    # Токенізація тексту
    text = str(row['text'])
    encoding = self.tokenizer(
        text,
        add_special_tokens=True,
        max_length=self.max_length,
        padding='max_length',
        truncation=True,
        return_tensors='pt'
    )

    # Часові ознаки
    temporal_features = self._get_temporal_features(row)

    # Метадані
    metadata_features = self._get_metadata_features(row)

    sample = {
        'text_ids': encoding['input_ids'].squeeze(0),
        'attention_mask': encoding['attention_mask'].squeeze(0),
        'temporal_features': torch.tensor(temporal_features, dtype=torch.float32),
        'metadata_features': torch.tensor(metadata_features, dtype=torch.float32),
    }

```

```

# Додавання міток для навчальних даних
if 'label' in row and not pd.isna(row['label']):
    sample['labels'] = torch.tensor(float(row['label']), dtype=torch.float32)

return sample

def _get_temporal_features(self, row: pd.Series) -> List[float]:
    """
    Витягування часових ознак для зразка
    """
    temporal_cols = [
        'posting_frequency', 'posting_regularity', 'night_activity',
        'burst_activity', 'weekday_activity'
    ]

    features = []
    for col in temporal_cols:
        if col in row:
            features.append(float(row[col]) if not pd.isna(row[col]) else 0.0)
        else:
            features.append(0.0)

    # Доповнення до потрібної розмірності (48 ознак)
    while len(features) < 48:
        features.append(0.0)

    return features[:48]

def _get_metadata_features(self, row: pd.Series) -> List[float]:
    """
    Витягування метаданих для зразка
    """
    metadata_cols = [
        'followers_count', 'friends_count', 'statuses_count',
        'account_age_days', 'profile_has_image', 'verified',
        'description_length', 'location_specified',
        'followers_friends_ratio', 'posts_per_day'
    ]

    features = []
    for col in metadata_cols:
        if col in row:
            features.append(float(row[col]) if not pd.isna(row[col]) else 0.0)
        else:
            features.append(0.0)

    # Доповнення до потрібної розмірності (98 ознак)
    while len(features) < 98:
        features.append(0.0)

    return features[:98]

```

А.6. Використання моделі та інференс

А.6.1. Інференс та оцінка

```

def predict_batch(model: MultimodalBotDetector,
                  dataloader: DataLoader,
                  device: torch.device) -> Tuple[np.ndarray, np.ndarray]:
    """

```

Передбачення для батчу даних

Returns:

predictions: Ймовірності належності до класу "бот"

features: Вилучені ознаки для аналізу

"""

model.eval()

all_predictions = []

all_features = []

with torch.no_grad():

for batch in tqdm(dataloader, desc="Inference"):

Переміщення на GPU

batch = {k: v.to(device) if torch.is_tensor(v) else v
for k, v in batch.items()}

Видалення міток якщо є

if 'labels' in batch:

batch.pop('labels')

Прямий прохід

outputs = model(batch)

Збір передбачень та ознак

predictions = outputs['probabilities'].cpu().numpy()

features = outputs['fused_features'].cpu().numpy()

all_predictions.extend(predictions.flatten())

all_features.extend(features)

return np.array(all_predictions), np.array(all_features)

def analyze_predictions(model: MultimodalBotDetector,

dataloader: DataLoader,

device: torch.device,

threshold: float = 0.5) -> Dict:

"""

Аналіз передбачень моделі з поясненнями

"""

import shap

from sklearn.metrics import classification_report, confusion_matrix

Отримання передбачень

predictions, features = predict_batch(model, dataloader, device)

Збір справжніх міток якщо доступні

true_labels = []

for batch in dataloader:

if 'labels' in batch:

true_labels.extend(batch['labels'].numpy().flatten())

results = {

'predictions': predictions,

'features': features,

'binary_predictions': (predictions >= threshold).astype(int)

```

}

if true_labels:
    true_labels = np.array(true_labels)
    results['true_labels'] = true_labels

    # Метрики класифікації
    results['classification_report'] = classification_report(
        true_labels, results['binary_predictions']
    )

    results['confusion_matrix'] = confusion_matrix(
        true_labels, results['binary_predictions']
    )

    # Детальні метрики
    results['metrics'] = compute_metrics(
        torch.tensor(predictions), torch.tensor(true_labels), threshold
    )

return results

```

```

def explain_prediction(model: MultimodalBotDetector,
                      sample: Dict[str, torch.Tensor],
                      device: torch.device) -> Dict:
    """
    Пояснення окремого передбачення
    """
    model.eval()

    # Переміщення на GPU
    sample = {k: v.to(device) if torch.is_tensor(v) else v
              for k, v in sample.items()}

    # Додавання batch dimension
    for key in sample:
        if torch.is_tensor(sample[key]) and sample[key].dim() == 1:
            sample[key] = sample[key].unsqueeze(0)

    with torch.no_grad():
        outputs = model(sample)

    explanation = {
        'prediction': outputs['probabilities'].item(),
        'confidence': abs(outputs['probabilities'].item() - 0.5) * 2,
        'text_features': outputs['text_features'].cpu().numpy(),
        'temporal_features': outputs['temporal_features'].cpu().numpy(),
        'metadata_features': outputs['metadata_features'].cpu().numpy(),
        'attention_weights': outputs.get('attention_weights', None)
    }

    # Класифікація
    if explanation['prediction'] >= 0.5:
        explanation['classification'] = 'BOT'
        explanation['confidence_label'] = f"Бот з впевненістю {explanation['confidence']:.2%}"
    else:

```

```

        explanation['classification'] = 'HUMAN'
        explanation['confidence_label'] = f"Людина"
        {explanation['confidence']:.2%}"
    """
    return explanation

```

3

впевненістю

```

def evaluate_model_comprehensive(model: MultimodalBotDetector,
                                test_loader: DataLoader,
                                device: torch.device) -> Dict:
    """
    Комплексна оцінка моделі
    """
    from sklearn.metrics import (
        precision_recall_curve, roc_curve, average_precision_score,
        precision_recall_fscore_support
    )

    # Отримання всіх передбачень
    all_predictions = []
    all_targets = []
    all_features = []

    model.eval()
    with torch.no_grad():
        for batch in tqdm(test_loader, desc="Evaluation"):
            batch = {k: v.to(device) if torch.is_tensor(v) else v
                     for k, v in batch.items()}

            targets = batch.pop('labels') if 'labels' in batch else None
            outputs = model(batch)

            all_predictions.extend(outputs['probabilities'].cpu().numpy().flatten())
            all_features.extend(outputs['fused_features'].cpu().numpy())

            if targets is not None:
                all_targets.extend(targets.cpu().numpy().flatten())

    predictions = np.array(all_predictions)
    features = np.array(all_features)

    evaluation_results = {
        'predictions': predictions,
        'features': features,
        'num_samples': len(predictions)
    }

    if all_targets:
        targets = np.array(all_targets)
        evaluation_results['targets'] = targets

    # ROC кривая
    fpr, tpr, roc_thresholds = roc_curve(targets, predictions)
    evaluation_results['roc_curve'] = {'fpr': fpr, 'tpr': tpr, 'thresholds': roc_thresholds}
    evaluation_results['auc_roc'] = roc_auc_score(targets, predictions)

    # Precision-Recall кривая

```

```

precision, recall, pr_thresholds = precision_recall_curve(targets, predictions)
evaluation_results['pr_curve'] = {
    'precision': precision, 'recall': recall, 'thresholds': pr_thresholds
}
evaluation_results['auc_pr'] = average_precision_score(targets, predictions)

# Метрики для різних порогів
thresholds = np.arange(0.1, 1.0, 0.1)
threshold_metrics = []

for threshold in thresholds:
    binary_preds = (predictions >= threshold).astype(int)
    metrics = compute_metrics(torch.tensor(predictions), torch.tensor(targets), threshold)
    metrics['threshold'] = threshold
    threshold_metrics.append(metrics)

evaluation_results['threshold_analysis'] = threshold_metrics

# Оптимальний поріг (максимізація F1)
best_threshold = max(threshold_metrics, key=lambda x: x['f1_score'])
evaluation_results['optimal_threshold'] = best_threshold

# Матриця плутанини для оптимального порогу
optimal_binary_preds = (predictions >= best_threshold['threshold']).astype(int)
evaluation_results['confusion_matrix'] = confusion_matrix(targets, optimal_binary_preds)

# Детальний звіт
evaluation_results['classification_report'] = classification_report(
    targets, optimal_binary_preds, output_dict=True
)

return evaluation_results

```

А.7. Візуалізація та аналіз результатів

А.7.1. Функції для візуалізації

```

import matplotlib.pyplot as plt
import seaborn as sns
from sklearn.manifold import TSNE
from sklearn.decomposition import PCA
import plotly.graph_objects as go
import plotly.express as px
from plotly.subplots import make_subplots

def plot_training_history(trainer: BotDetectorTrainer, save_path: str = None):
    """
    Візуалізація історії навчання
    """
    train_history = trainer.train_history
    val_history = trainer.val_history

    epochs = range(1, len(train_history) + 1)

    fig, ((ax1, ax2), (ax3, ax4)) = plt.subplots(2, 2, figsize=(15, 10))

    # Loss
    ax1.plot(epochs, [h['loss'] for h in train_history], 'b-', label='Навчання')
    ax1.plot(epochs, [h['loss'] for h in val_history], 'r-', label='Валідація')
    ax1.set_title('Втрати')

```

```

ax1.set_xlabel('Епоха')
ax1.set_ylabel('Loss')
ax1.legend()
ax1.grid(True)

# F1 Score
ax2.plot(epochs, [h['f1_score'] for h in train_history], 'b-', label='Навчання')
ax2.plot(epochs, [h['f1_score'] for h in val_history], 'r-', label='Валідація')
ax2.set_title('F1-міра')
ax2.set_xlabel('Епоха')
ax2.set_ylabel('F1 Score')
ax2.legend()
ax2.grid(True)

# Precision
ax3.plot(epochs, [h['precision'] for h in train_history], 'b-', label='Навчання')
ax3.plot(epochs, [h['precision'] for h in val_history], 'r-', label='Валідація')
ax3.set_title('Точність')
ax3.set_xlabel('Епоха')
ax3.set_ylabel('Precision')
ax3.legend()
ax3.grid(True)

# Recall
ax4.plot(epochs, [h['recall'] for h in train_history], 'b-', label='Навчання')
ax4.plot(epochs, [h['recall'] for h in val_history], 'r-', label='Валідація')
ax4.set_title('Повнота')
ax4.set_xlabel('Епоха')
ax4.set_ylabel('Recall')
ax4.legend()
ax4.grid(True)

plt.tight_layout()

if save_path:
    plt.savefig(save_path, dpi=300, bbox_inches='tight')

plt.show()

```

```

def plot_roc_pr_curves(evaluation_results: Dict, save_path: str = None):

```

```

    """
    Візуалізація ROC та PR кривих
    """
    fig, (ax1, ax2) = plt.subplots(1, 2, figsize=(15, 6))

    # ROC Curve
    roc_data = evaluation_results['roc_curve']
    auc_roc = evaluation_results['auc_roc']

    ax1.plot(roc_data['fpr'], roc_data['tpr'], 'b-',
             label=f'ROC кривая (AUC = {auc_roc:.3f})')
    ax1.plot([0, 1], [0, 1], 'r--', label='Випадковий класифікатор')
    ax1.set_xlabel('False Positive Rate')
    ax1.set_ylabel('True Positive Rate')
    ax1.set_title('ROC Кривая')
    ax1.legend()

```

```

ax1.grid(True)

# PR Curve
pr_data = evaluation_results['pr_curve']
auc_pr = evaluation_results['auc_pr']

ax2.plot(pr_data['recall'], pr_data['precision'], 'b-',
         label=f'PR кривая (AUC = {auc_pr:.3f})')
ax2.set_xlabel('Recall')
ax2.set_ylabel('Precision')
ax2.set_title('Precision-Recall Кривая')
ax2.legend()
ax2.grid(True)

plt.tight_layout()

if save_path:
    plt.savefig(save_path, dpi=300, bbox_inches='tight')

plt.show()

def plot_confusion_matrix(confusion_matrix: np.ndarray,
                          class_names: List[str] = ['Людина', 'Бот'],
                          save_path: str = None):
    """
    Візуалізація матриці плутанини
    """
    plt.figure(figsize=(8, 6))

    # Нормалізована матриця плутанини
    cm_normalized = confusion_matrix.astype('float') / confusion_matrix.sum(axis=1)[:,
np.newaxis]

    sns.heatmap(cm_normalized, annot=True, fmt='.2%', cmap='Blues',
                xticklabels=class_names, yticklabels=class_names)

    plt.title('Матриця плутанини (нормалізована)')
    plt.xlabel('Передбачені мітки')
    plt.ylabel('Справжні мітки')

    # Додавання абсолютних значень
    for i in range(len(class_names)):
        for j in range(len(class_names)):
            plt.text(j + 0.5, i + 0.7, f'{confusion_matrix[i, j]}',
                    ha='center', va='center', fontsize=10, color='red')

    plt.tight_layout()

    if save_path:
        plt.savefig(save_path, dpi=300, bbox_inches='tight')

    plt.show()

def plot_feature_importance(model: MultimodalBotDetector,
                           sample_data: torch.Tensor,

```

```

        device: torch.device,
        save_path: str = None):
    """
    Візуалізація важливості ознак за допомогою SHAP
    """
    import shap

    # Підготовка даних для SHAP
    model.eval()

    def model_predict(x):
        with torch.no_grad():
            # Конвертація в правильний формат
            batch = {
                'text_ids': torch.tensor(x[:, :512], dtype=torch.long).to(device),
                'attention_mask': torch.ones((x.shape[0], 512), dtype=torch.long).to(device),
                'temporal_features': torch.tensor(x[:, 512:560], dtype=torch.float32).to(device),
                'metadata_features': torch.tensor(x[:, 560:], dtype=torch.float32).to(device)
            }
            outputs = model(batch)
            return outputs['probabilities'].cpu().numpy()

    # SHAP аналіз
    explainer = shap.KernelExplainer(model_predict, sample_data[:100])
    shap_values = explainer.shap_values(sample_data[:20])

    # Візуалізація
    plt.figure(figsize=(12, 8))
    shap.summary_plot(shap_values, sample_data[:20],
                      feature_names=['Text', 'Temporal', 'Metadata'],
                      show=False)

    if save_path:
        plt.savefig(save_path, dpi=300, bbox_inches='tight')

    plt.show()

def plot_embeddings_tsne(features: np.ndarray, labels: np.ndarray = None,
                        save_path: str = None):
    """
    Візуалізація вбудовувань за допомогою t-SNE
    """
    # t-SNE зменшення розмірності
    tsne = TSNE(n_components=2, random_state=42, perplexity=30)
    embeddings_2d = tsne.fit_transform(features)

    plt.figure(figsize=(12, 8))

    if labels is not None:
        # Кольорове маркування за мітками
        colors = ['blue' if l == 0 else 'red' for l in labels]
        scatter = plt.scatter(embeddings_2d[:, 0], embeddings_2d[:, 1],
                              c=colors, alpha=0.6, s=20)

    # Легенда
    blue_patch = plt.Line2D([0], [0], marker='o', color='w',

```

```

        markerfacecolor='blue', markersize=8, label='Люди')
    red_patch = plt.Line2D([0], [0], marker='o', color='w',
        markerfacecolor='red', markersize=8, label='Боти')
    plt.legend(handles=[blue_patch, red_patch])
else:
    plt.scatter(embeddings_2d[:, 0], embeddings_2d[:, 1], alpha=0.6, s=20)

plt.title('t-SNE візуалізація вбудовувань')
plt.xlabel('t-SNE компонент 1')
plt.ylabel('t-SNE компонент 2')
plt.grid(True, alpha=0.3)

if save_path:
    plt.savefig(save_path, dpi=300, bbox_inches='tight')

plt.show()

def plot_threshold_analysis(evaluation_results: Dict, save_path: str = None):
    """
    Аналіз впливу порогу класифікації
    """
    threshold_data = evaluation_results['threshold_analysis']

    thresholds = [d['threshold'] for d in threshold_data]
    f1_scores = [d['f1_score'] for d in threshold_data]
    precisions = [d['precision'] for d in threshold_data]
    recalls = [d['recall'] for d in threshold_data]

    plt.figure(figsize=(12, 8))

    plt.plot(thresholds, f1_scores, 'b-o', label='F1-міра', linewidth=2)
    plt.plot(thresholds, precisions, 'r-s', label='Точність', linewidth=2)
    plt.plot(thresholds, recalls, 'g-^', label='Повнота', linewidth=2)

    # Позначення оптимального порогу
    optimal_threshold = evaluation_results['optimal_threshold']['threshold']
    optimal_f1 = evaluation_results['optimal_threshold']['f1_score']

    plt.axvline(x=optimal_threshold, color='orange', linestyle='--',
        label=f'Оптимальний поріг = {optimal_threshold:.2f}')
    plt.plot(optimal_threshold, optimal_f1, 'ro', markersize=10,
        label=f'Макс. F1 = {optimal_f1:.3f}')

    plt.xlabel('Поріг класифікації')
    plt.ylabel('Значення метрики')
    plt.title('Аналіз впливу порогу класифікації')
    plt.legend()
    plt.grid(True, alpha=0.3)
    plt.xlim(0.0, 1.0)
    plt.ylim(0.0, 1.0)

    if save_path:
        plt.savefig(save_path, dpi=300, bbox_inches='tight')

    plt.show()

```

A.8. Основний скрипт для запуску

A.8.1. Конфігурація та запуск навчання

```
import argparse
import yaml
import os
from pathlib import Path

def load_config(config_path: str) -> Dict:
    """
    Завантаження конфігурації з YAML файлу
    """
    with open(config_path, 'r', encoding='utf-8') as f:
        config = yaml.safe_load(f)
    return config

def setup_logging():
    """
    Налаштування логуювання
    """
    import logging

    logging.basicConfig(
        level=logging.INFO,
        format='%(asctime)s - %(name)s - %(levelname)s - %(message)s',
        handlers=[
            logging.FileHandler('bot_detection.log'),
            logging.StreamHandler()
        ]
    )

    return logging.getLogger(__name__)

def main():
    """
    Основна функція для запуску навчання або інференсу
    """
    parser = argparse.ArgumentParser(description='Навчання моделі виявлення ботів')
    parser.add_argument('--config', type=str, required=True,
                        help='Шлях до конфігураційного файлу')
    parser.add_argument('--mode', type=str, choices=['train', 'evaluate', 'predict'],
                        default='train', help='Режим роботи')
    parser.add_argument('--model_path', type=str,
                        help='Шлях до збереженої моделі')
    parser.add_argument('--data_path', type=str,
                        help='Шлях до даних')
    parser.add_argument('--output_dir', type=str, default='./outputs',
                        help='Директорія для збереження результатів')

    args = parser.parse_args()

    # Налаштування логуювання
    logger = setup_logging()
    logger.info(f"Запуск у режимі: {args.mode}")

    # Завантаження конфігурації
    config = load_config(args.config)
```

```

# Створення директорії для результатів
output_dir = Path(args.output_dir)
output_dir.mkdir(parents=True, exist_ok=True)

# Ініціалізація токенизатора
tokenizer = RobertaTokenizer.from_pretrained('roberta-base')

# Пристрій для обчислень
device = torch.device('cuda' if torch.cuda.is_available() else 'cpu')
logger.info(f"Використання пристрою: {device}")

if args.mode == 'train':
    # Навчання моделі
    logger.info("Початок навчання моделі")

    # Завантаження даних
    train_dataset = BotDetectionDataset(
        data_path=config['data']['train_path'],
        tokenizer=tokenizer,
        max_length=config['model']['max_length'],
        is_training=True
    )

    val_dataset = BotDetectionDataset(
        data_path=config['data']['val_path'],
        tokenizer=tokenizer,
        max_length=config['model']['max_length'],
        is_training=False
    )

    # DataLoader
    train_loader = DataLoader(
        train_dataset,
        batch_size=config['training']['batch_size'],
        shuffle=True,
        num_workers=4
    )

    val_loader = DataLoader(
        val_dataset,
        batch_size=config['training']['batch_size'],
        shuffle=False,
        num_workers=4
    )

    # Ініціалізація моделі
    model = MultimodalBotDetector(config['model'])
    logger.info(f"Модель ініціалізована з {sum(p.numel() for p in model.parameters())}
параметрів")

    # Тренер
    trainer = BotDetectorTrainer(model, config['training'])

    # Навчання
    trainer.train(
        train_loader=train_loader,

```

```

        val_loader=val_loader,
        num_epochs=config['training']['num_epochs'],
        early_stopping_patience=config['training']['early_stopping_patience']
    )

    # Візуалізація історії навчання
    plot_training_history(trainer, save_path=output_dir / 'training_history.png')

    logger.info("Навчання завершено")

elif args.mode == 'evaluate':
    # Оцінка моделі
    logger.info("Початок оцінки моделі")

    # Завантаження моделі
    model = MultimodalBotDetector(config['model'])
    checkpoint = torch.load(args.model_path, map_location=device)
    model.load_state_dict(checkpoint['model_state_dict'])
    model.to(device)

    # Завантаження тестових даних
    test_dataset = BotDetectionDataset(
        data_path=args.data_path or config['data']['test_path'],
        tokenizer=tokenizer,
        max_length=config['model']['max_length'],
        is_training=False
    )

    test_loader = DataLoader(
        test_dataset,
        batch_size=config['training']['batch_size'],
        shuffle=False,
        num_workers=4
    )

    # Комплексна оцінка
    evaluation_results = evaluate_model_comprehensive(model, test_loader, device)

    # Збереження результатів
    with open(output_dir / 'evaluation_results.json', 'w') as f:
        # Конвертація numpy arrays для JSON
        results_serializable = {}
        for key, value in evaluation_results.items():
            if isinstance(value, np.ndarray):
                results_serializable[key] = value.tolist()
            else:
                results_serializable[key] = value

        json.dump(results_serializable, f, indent=2)

    # Візуалізації
    plot_roc_pr_curves(evaluation_results, save_path=output_dir / 'roc_pr_curves.png')
    plot_confusion_matrix(evaluation_results['confusion_matrix'],
                          save_path=output_dir / 'confusion_matrix.png')
    plot_threshold_analysis(evaluation_results,
                           save_path=output_dir / 'threshold_analysis.png')

```

```

# Візуалізація вбудовувань
if 'targets' in evaluation_results:
    plot_embeddings_tsne(evaluation_results['features'],
                        evaluation_results['targets'],
                        save_path=output_dir / 'embeddings_tsne.png')

logger.info("Оцінка завершена")

elif args.mode == 'predict':
    # Передбачення для нових даних
    logger.info("Початок передбачення")

    # Завантаження моделі
    model = MultimodalBotDetector(config['model'])
    checkpoint = torch.load(args.model_path, map_location=device)
    model.load_state_dict(checkpoint['model_state_dict'])
    model.to(device)

    # Завантаження даних
    dataset = BotDetectionDataset(
        data_path=args.data_path,
        tokenizer=tokenizer,
        max_length=config['model']['max_length'],
        is_training=False
    )

    dataloader = DataLoader(
        dataset,
        batch_size=config['training']['batch_size'],
        shuffle=False,
        num_workers=4
    )

    # Передбачення
    predictions, features = predict_batch(model, dataloader, device)

    # Збереження результатів
    results_df = pd.DataFrame({
        'prediction_probability': predictions,
        'prediction_binary': (predictions >= 0.5).astype(int),
        'confidence': np.abs(predictions - 0.5) * 2
    })

    results_df.to_csv(output_dir / 'predictions.csv', index=False)

    logger.info(f"Передбачення збережено в {output_dir / 'predictions.csv'}")
    logger.info(f"Виявлено {sum(predictions >= 0.5)} ботів з {len(predictions)} облікових
записів")

if __name__ == '__main__':
    main()

```

A.8.2. Приклад конфігураційного файлу (config.yaml)

Конфігурація моделі виявлення ботів

```
model:
  hidden_dim: 768
  num_attention_heads: 12
  temporal_dim: 128
  metadata_dim: 64
  dropout_rate: 0.1
  max_length: 512
  temporal_features: 48
  metadata_features: 98

data:
  train_path: "data/train.csv"
  val_path: "data/val.csv"
  test_path: "data/test.csv"

training:
  batch_size: 32
  learning_rate: 2e-5
  weight_decay: 0.01
  num_epochs: 12
  early_stopping_patience: 3

# Функція втрат
focal_weight: 0.8
l2_weight: 0.01
focal_alpha: 0.25
focal_gamma: 2.0

# Планувальник
warmup_ratio: 0.1

# Збереження моделі
save_dir: "models/"
save_every_n_epochs: 1

evaluation:
  threshold: 0.5
  metrics:
    - precision
    - recall
    - f1_score
    - auc_roc
    - auc_pr
    - mcc

logging:
  level: INFO
  log_file: "bot_detection.log"

wandb:
  project: "bot-detection"
  entity: "your-username"
  enabled: false
```

Додаток Б. Результати додаткових експериментів

Б.1. Порівняльний аналіз різних архітектур трансформерів

Для визначення оптимальної базової архітектури було проведено порівняльне дослідження ефективності різних трансформер-моделей на наборі даних TwiBot-20. Експерименти проводились з однаковими гіперпараметрами для забезпечення справедливого порівняння.

Таблиця Б.1. Порівняння різних трансформер-архітектур

Модель	Параметри (М)	F1-Score	AUC-ROC	AUC-PR	Час навчання (год)	Пам'ять (ГБ)
BERT-base	110	0.895 ± 0.006	0.951 ± 0.005	0.876 ± 0.011	42.3	8.2
RoBERTa-base	125	0.903 ± 0.007	0.957 ± 0.006	0.891 ± 0.012	45.1	8.9
DistilBERT	66	0.878 ± 0.009	0.932 ± 0.008	0.851 ± 0.015	28.7	5.1
ALBERT-base	12	0.863 ± 0.012	0.925 ± 0.010	0.823 ± 0.018	35.6	4.3
DeBERTa-base	140	0.908 ± 0.005	0.962 ± 0.004	0.897 ± 0.009	52.4	9.7
Electra-base	110	0.891 ± 0.008	0.948 ± 0.007	0.869 ± 0.013	38.9	8.1

Результати показують, що RoBERTa-base забезпечує найкращий баланс між ефективністю та обчислювальною складністю, тому була обрана як базова архітектура для запропонованої моделі.

Б.2. Дослідження впливу розміру навчального набору

Проведено експерименти для визначення мінімального розміру навчального набору, необхідного для досягнення стабільної ефективності моделі. Використовувались підмножини різного розміру з повного навчального набору TwiBot-20.

Таблиця Б.2. Вплив розміру навчального набору на ефективність

Розмір набору	% від повного	F1-Score	AUC-ROC	Стандартне відхилення F1	Час навчання (хв)
5,000	3.1%	0.762 ± 0.023	0.834 ± 0.019	0.031	12.3
10,000	6.2%	0.821 ± 0.018	0.889 ± 0.015	0.024	18.7
25,000	15.5%	0.874 ± 0.012	0.932 ± 0.010	0.018	35.2
50,000	31.0%	0.912 ± 0.008	0.963 ± 0.006	0.012	58.4
100,000	62.1%	0.936 ± 0.006	0.975 ± 0.004	0.008	98.1
161,306 (повний)	100%	0.944 ± 0.004	0.981 ± 0.003	0.006	142.5

Аналіз показує, що для досягнення F1-міри вище 0.90 необхідно мінімум 50,000 зразків. Повний набір забезпечує найкращу стабільність результатів з найменшим стандартним відхиленням.

Б.3. Експерименти з різними стратегіями аугментації даних

Досліджено ефективність різних методів аугментації даних для подолання проблеми дисбалансу класів та покращення генералізації моделі.

Таблиця Б.3. Порівняння стратегій аугментації даних

Стратегія аугментації	Кількість синт. зразків	F1-Score	Precision	Recall	Покращення
Без аугментації	0	0.863 ± 0.012	0.892 ± 0.010	0.836 ± 0.014	-
Синонімічна заміна	15,000	0.881 ± 0.009	0.903 ± 0.008	0.860 ± 0.011	+0.018

Перефразування (PEGASUS)	20,000	0.896 ± 0.007	0.916 ± 0.006	0.877 ± 0.009	+0.033
Зворотний переклад	18,000	0.889 ± 0.008	0.909 ± 0.007	0.870 ± 0.010	+0.026
SMOTE для метаданих	12,000	0.878 ± 0.010	0.898 ± 0.009	0.859 ± 0.012	+0.015
Часова аугментація	25,000	0.885 ± 0.008	0.907 ± 0.007	0.864 ± 0.010	+0.022
Комбінована стратегія	45,000	0.912 ± 0.006	0.929 ± 0.005	0.896 ± 0.007	+0.049
Мультимодальна (запропонована)	50,000	0.937 ± 0.004	0.947 ± 0.005	0.927 ± 0.006	+0.074

Запропонована мультимодальна стратегія аугментації показала найкращі результати, забезпечивши покращення F1-міри на 7.4% порівняно з базовою моделлю.

Б.4. Аналіз ефективності різних механізмів уваги

Проведено порівняльне дослідження різних типів механізмів уваги для аналізу часових патернів активності користувачів.

Таблиця Б.4. Порівняння механізмів уваги для часових даних

Тип уваги	Параметри	F1-Score	Часова складність	Пам'ять (МБ)	Інтерпретованість
Без уваги (LSTM)	2M	0.894 ± 0.008	O(n)	145	Низька
Стандартна увага	4M	0.918 ± 0.006	O(n ²)	289	Середня
Multi-head attention	8M	0.925 ± 0.005	O(n ²)	312	Середня
Sparse attention	6M	0.921 ± 0.006	O(n log n)	234	Середня
Відносна часова увага	7M	0.938 ± 0.005	O(n ²)	298	Висока
Linformer	5M	0.915 ± 0.007	O(n)	187	Низька
Performer	6M	0.919 ± 0.006	O(n log n)	203	Низька

Відносна часова увага, використана в запропонованій моделі, забезпечує найкращу ефективність та високу інтерпретованість результатів.

Б.5. Дослідження стійкості до адверсаріальних атак

Проведено тестування стійкості запропонованої моделі до цілеспрямованих спроб обходу виявлення.

Таблиця Б.5. Стійкість до адверсаріальних атак

Тип атаки	Опис модифікації	F1-Score (до)	F1-Score (після)	Зниження
Базова модель	Без атак	0.944	0.944	0%
Додавання шуму	Випадкові символи в тексті (5%)	0.944	0.931	-1.4%
Синонімічна заміна	Заміна 10% слів синонімами	0.944	0.926	-1.9%
Часова маскування	Видалення 20% часових точок	0.944	0.918	-2.8%
Метадата спотінг	Модифікація метаданих профілю	0.944	0.912	-3.4%
Комбінована атака	Всі вищезазначені модифікації	0.944	0.897	-5.0%
FGSM атака	Gradient-based adversarial	0.944	0.889	-5.8%
PGD атака	Projected gradient descent	0.944	0.876	-7.2%

Моделю демонструє відносну стійкість до більшості типів атак, зберігаючи F1-міру вище 0.87 навіть при складних адверсаріальних модифікаціях.

Б.6. Аналіз ефективності на різних типах ботів

Детальний аналіз ефективності моделі для виявлення різних категорій ботів на наборі CRESCI-2017.

Таблиця Б.6. Ефективність виявлення різних типів ботів

Тип бота	Кількість зразків	Precision	Recall	F1-Score	Особливості
Спам-боти	1,610	0.967 ± 0.008	0.951 ± 0.012	0.959 ± 0.007	Легко виявляються
Боти-підсилювачі	3,474	0.934 ± 0.011	0.928 ± 0.014	0.931 ± 0.009	Регулярні патерни
Фальшиві підписники	19,276	0.912 ± 0.006	0.897 ± 0.009	0.904 ± 0.005	Варіабельна активність
Боти-агрегатори	892	0.889 ± 0.015	0.876 ± 0.018	0.882 ± 0.013	Автоматичний контент
Імітатори людей	2,156	0.847 ± 0.012	0.823 ± 0.016	0.835 ± 0.011	Складні для виявлення
GPT-подібні боти	1,234	0.793 ± 0.018	0.765 ± 0.022	0.779 ± 0.016	Найскладніші
Середнє значення	-	0.890 ± 0.012	0.873 ± 0.015	0.882 ± 0.010	-

Найскладнішими для виявлення виявились GPT-подібні боти, які використовують генеративні моделі для створення природного контенту.

Б.7. Експерименти з різними розмірами контекстного вікна

Дослідження впливу максимальної довжини вхідної послідовності на ефективність моделі.

Таблиця Б.7. Вплив розміру контекстного вікна

Max Length	F1-Score	AUC-ROC	Час інференсу (мс)	Пам'ять (МБ)	Покриття тексту (%)
128	0.887 ± 0.012	0.943 ± 0.009	23	187	67%
256	0.921 ± 0.008	0.967 ± 0.006	35	234	84%
512	0.944 ± 0.004	0.981 ± 0.003	58	298	96%
768	0.946 ± 0.005	0.982 ± 0.003	89	412	98%
1024	0.947 ± 0.004	0.983 ± 0.003	134	567	99%

Оптимальним виявився розмір 512 токенів, що забезпечує покриття 96% текстового контенту при прийнятній обчислювальній складності.

Б.8. Аналіз впливу дисбалансу класів

Дослідження поведінки моделі при різних співвідношеннях класів у навчальних даних.

Таблиця Б.8. Вплив дисбалансу класів на ефективність

Співвідношення (Люди:Боти)	F1-Score	Precision	Recall	Стратегія балансування
50:50	0.967 ± 0.003	0.972 ± 0.004	0.962 ± 0.005	Штучне балансування
60:40	0.958 ± 0.004	0.965 ± 0.005	0.951 ± 0.006	Незначний дисбаланс
70:30	0.948 ± 0.005	0.958 ± 0.006	0.938 ± 0.007	Помірний дисбаланс
78:22 (реальний)	0.944 ± 0.004	0.947 ± 0.005	0.942 ± 0.007	Focal Loss + аугментація
85:15	0.925 ± 0.008	0.943 ± 0.007	0.908 ± 0.011	Сильний дисбаланс
90:10	0.891 ± 0.012	0.934 ± 0.009	0.853 ± 0.016	Критичний дисбаланс

95:5	0.832 ± 0.018	0.918 ± 0.012	0.761 ± 0.023	Екстремальний дисбаланс
------	---------------	---------------	---------------	-------------------------

Результати підтверджують ефективність запропонованої стратегії обробки дисбалансу класів навіть при реальних співвідношеннях.

Б.9. Крос-доменна генералізація

Тестування здатності моделі, навченої на одному наборі даних, ефективно працювати на інших платформах.

Таблиця Б.9. Крос-доменна ефективність моделі

Навчання Тестування →	F1-Score	AUC-ROC	Падіння ефективності	Примітки
Twibot-20 → Twibot-20	0.944 ± 0.004	0.981 ± 0.003	-	Базовий результат
Twibot-20 → CRESCI-2017	0.892 ± 0.009	0.951 ± 0.007	-5.5%	Добра генералізація
Twibot-20 → BotometerLite	0.876 ± 0.012	0.943 ± 0.009	-7.2%	Середня генералізація
Twibot-20 → Власний набір	0.823 ± 0.015	0.898 ± 0.013	-12.8%	GPT-боти складніші
CRESCI-2017 → Twibot-20	0.857 ± 0.011	0.924 ± 0.008	-9.2%	Менший навчальний набір
Комбінований → Всі	0.918 ± 0.006	0.963 ± 0.005	-2.8%	Найкраща генералізація

Навчання на комбінованому наборі даних забезпечує найкращу крос-доменну генералізацію.

Б.10. Аналіз обчислювальної ефективності на різному обладнанні

Порівняння швидкості роботи моделі на різних типах обладнання.

Таблиця Б.10. Обчислювальна ефективність на різному обладнанні

Обладнання	Batch Size	Час інференсу (мс/зразок)	Пропускна здатність (зразків/с)	Пам'ять (ГБ)
CPU (Intel i7-10700K)	1	180	5.6	2.1
CPU (Intel i7-10700K)	32	98	326	2.8
GPU (GTX 1080 Ti)	1	67	14.9	3.2
GPU (GTX 1080 Ti)	32	28	1,143	8.9
GPU (RTX 3080)	1	45	22.2	3.1
GPU (RTX 3080)	32	18	1,778	8.7
GPU (V100)	1	32	31.3	3.0
GPU (V100)	32	12	2,667	8.5
GPU (A100)	1	23	43.5	2.9
GPU (A100)	32	8	4,000	8.3

Для продакшн-розгортання рекомендується використання GPU з батчевою обробкою для досягнення оптимальної пропускної здатності.

Б.11. Порівняння з комерційними рішеннями

Порівняння запропонованої моделі з доступними комерційними системами виявлення ботів.

Таблиця Б.11. Порівняння з комерційними системами

Система	F1-Score	AUC-ROC	Швидкість	Вартість	Доступність API
Botometer	0.823 ± 0.015	0.897 ± 0.012	Швидка	Безкоштовна (ліміт)	Так
BotSlayer	0.798 ± 0.018	0.883 ± 0.014	Повільна	\$50/місяць	Так
TweetDeck (Twitter)	0.756 ± 0.021	0.841 ± 0.017	Швидка	Інтегровано	Обмежено

Microsoft Defender	0.834 ± 0.012	0.912 ± 0.009	Середня	\$200/місяць	Так
AWS Fraud Detector	0.847 ± 0.011	0.923 ± 0.008	Швидка	\$0.10/1000 запитів	Так
Запропонована модель	0.944 ± 0.004	0.981 ± 0.003	Налаштовується	Безкоштовна	Так

Запропонована модель демонструє значні переваги за точністю виявлення при збереженні гнучкості розгортання.

Б.12. Аналіз помилок за географічними регіонами

Дослідження ефективності моделі для користувачів з різних географічних регіонів.

Таблиця Б.12. Ефективність за географічними регіонами

Регіон	Кількість зразків	F1-Score	Особливості	Основні помилки
Північна Америка	89,234	0.951 ± 0.005	Високоякісні дані	Мало помилок
Європа	67,891	0.946 ± 0.006	Мультимовність	Мовні варіації
Азія	34,567	0.923 ± 0.009	Різні часові пояси	Часові патерни
Південна Америка	12,456	0.934 ± 0.008	Ієрогліфи, емодзі	Кодування тексту
Африка	5,789	0.912 ± 0.012	Обмежені дані	Недостатнє представлення
Океанія	2,134	0.897 ± 0.015	Малий набір	Статистична похибка

Модель показує стабільну ефективність у різних регіонах з незначними варіаціями, пов'язаними з локальними особливостями.

Б.13. Довгострокова стабільність моделі

Аналіз деградації ефективності моделі з часом без перенавчання.

Таблиця Б.13. Деградація ефективності з часом

Період після навчання	F1-Score	Зниження	Основні причини
0 місяців (базовий)	0.944 ± 0.004	0%	-
1 місяць	0.941 ± 0.005	-0.3%	Незначні зміни
3 місяці	0.932 ± 0.007	-1.3%	Еволюція ботів
6 місяців	0.918 ± 0.009	-2.8%	Нові стратегії ботів
12 місяців	0.897 ± 0.012	-5.0%	Значні зміни середовища
18 місяців	0.876 ± 0.015	-7.2%	Потреба в оновленні
24 місяці	0.852 ± 0.018	-9.7%	Критична потреба оновлення

Рекомендується перенавчання моделі кожні 6-12 місяців для збереження високої ефективності.

Б.14. Аналіз інтерпретованості рішень

Результати експертної оцінки якості пояснень, що надаються моделлю.

Таблиця Б.14. Оцінка якості пояснень (експертна оцінка 1-10)

Аспект пояснення	Оцінка експертів	Стандартне відхилення	Коментарі
Релевантність текстових ознак	8.7 ± 0.9	0.9	Добре виділяє підозрілі фрази
Зрозумілість часових патернів	8.3 ± 1.1	1.1	Наочні графіки активності
Важливість метаданих	7.9 ± 1.0	1.0	Потрібні додаткові пояснення
Загальна логічність	8.5 ± 0.8	0.8	Логічні та послідовні рішення
Практична корисність	9.1 ± 0.7	0.7	Допомагає в прийнятті рішень
Довіра до системи	8.6 ± 0.9	0.9	Високий рівень довіри

Пояснення моделі отримали високі оцінки експертів, особливо за практичну корисність та загальну логічність.

Б.15. Ресурси та рекомендації для реплікації

Детальна інформація про необхідні ресурси для відтворення результатів дослідження.

Таблиця Б.15. Мінімальні системні вимоги для реплікації

Компонент	Мінімальні вимоги	Рекомендовані вимоги	Оптимальні вимоги
CPU	Intel i5-8400 / AMD Ryzen 5 2600	Intel i7-10700K / AMD Ryzen 7 3700X	Intel i9-11900K / AMD Ryzen 9 5900X
RAM	16 ГБ	32 ГБ	64 ГБ
GPU	GTX 1660 (6 ГБ)	RTX 3070 (8 ГБ)	RTX 4090 (24 ГБ)
Диск	100 ГБ SSD	500 ГБ NVMe SSD	1 ТБ NVMe SSD
Час навчання	~8 годин	~4 години	~2 години
Споживання енергії	~300 Вт	~450 Вт	~700 Вт

Програмне забезпечення:

- Python 3.8+
- PyTorch 1.12+
- Transformers 4.20+
- CUDA 11.6+ (для GPU)
- Ubuntu 20.04+ / Windows 10+ / macOS 12+

Набори даних:

- TwiBot-20: <https://github.com/BunsenFeng/TwiBot-20>
- CRESCI-2017: <https://github.com/MibexSoftware/cresci-2017>
- BotometerLite: <https://github.com/IUNetSci/botometer-python>