

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
Західноукраїнський національний університет
Навчально-науковий інститут новітніх освітніх технологій
Кафедра інформаційно-обчислювальних систем і управління

НАКОНЕЧНА Ірина Володимирівна

**Нейромережевий модуль визначення джерела тексту з використанням
Embedding і LSTM / Neural network module for text source identification using
Embedding and LSTM**

Спеціальність 122 – Комп'ютерні науки
Освітньо-професійна програма – Штучний інтелект

Кваліфікаційна робота

Виконала студент групи КНШІ-41
І.В. Наконечна

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

Кваліфікаційну роботу допущено до
захисту
«___»_____ 2025 р.

В.о. завідувача кафедри
_____Н.М. Васильків

Тернопіль – 2025

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

ЗАТВЕРДЖУЮ
В.о. завідувача кафедри
_____ Н.М. Васильків
«_____» _____ 2024 р.

ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ
Наконечна Ірина Володимирівна
(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи: Нейромережевий модуль визначення джерела тексту з використанням Embedding і LSTM / Neural network module for text source identification using Embedding and LSTM

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

(прізвище, ім'я, по батькові, науковий ступінь, вчене звання)

затверджені наказом по університету від 20 грудня 2024 р. № 938.

2. Строк подання студентом закінченої кваліфікаційної роботи 25 травня 2025 р.

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

4. Основні питання, які потрібно розробити:

– проаналізувати сучасні підходи до автоматичного визначення джерела тексту в задачах обробки природної мови та виявити їхні обмеження;здійснити аналіз і попередню обробку текстового датасету;

– побудувати інформаційну модель даних, реалізувати кластеризацію за допомогою алгоритму KMeans і візуалізацію результатів (PCA);

– розробити архітектуру нейронної мережі з використанням шарів Embedding і LSTM для багатокласової класифікації текстів;

– реалізувати Telegram-бот, який взаємодіє з навченою моделлю для класифікації вхідного тексту;

– провести тестування модуля, оцінити точність класифікації та надійність інтегрованого рішення.

5. Перелік графічного матеріалу в роботі:

– діаграма «суть-зв'язок» для модуля обробки та класифікації текстів;

– схема алгоритму попередньої обробки та подачі даних у модель;

6. Консультанти розділів кваліфікаційної роботи

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

7. Дата видачі завдання 20 грудня 2024 р.

КАЛЕНДАРНИЙ ПЛАН

з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів кваліфікаційної роботи	Примітка
1	Затвердження теми кваліфікаційної роботи, ознайомлення з літературними джерелами та складання плану роботи	до 01.01. 2025 р.	
2	Написання 1 розділу кваліфікаційної роботи	до 01.02. 2025 р.	
3	Написання 2 розділу кваліфікаційної роботи	до 01.04.2025 р.	
4	Написання 3 розділу кваліфікаційної роботи	до 01.05. 2025 р.	
5	Представлення попереднього варіанту кваліфікаційної роботи, перевірка та внесення змін керівником	до 15.05.2025 р.	
6	Опрацювання зауважень та представлення завершеного варіанту кваліфікаційної роботи. Підготовка супроводжуючих документів	до 25.05.2025 р.	
	Перевірка кваліфікаційної роботи на оригінальність тексту	до 30.05.2025 р.	
	Оформлення кваліфікаційної роботи та отримання допуску до захисту	до 10.06.2025 р.	
	Подання кваліфікаційної роботи до захисту на засіданні атестаційної комісії	до 10.06. 2025 р.	

Студент _____ І. В. Наконечна
(підпис) (прізвище та ініціали)

Керівник кваліфікаційної роботи _____ Х. В. Лип'яніна-Гончаренко
(підпис) (прізвище та ініціали)

АНОТАЦІЯ

Кваліфікаційна робота на тему «Нейромережевий модуль визначення джерела тексту з використанням Embedding і LSTM» на здобуття освітнього ступеня «бакалавр» зі спеціальності 122 «Комп'ютерні науки» освітньої програми «Штучний інтелект» написана обсягом у 51 сторінку і містить 16 ілюстрацій, 2 таблиці та 25 використаних джерела.

Метою роботи є створення модуля для автоматичного визначення джерела тексту на основі методів обробки природної мови та глибокого навчання, зокрема технологій Embedding та LSTM, що дозволяє класифікувати тексти за їх семантичними та стилістичними ознаками.

Методами дослідження обрано: аналіз існуючих підходів до текстової класифікації, побудову архітектури LSTM-моделі з використанням векторного представлення слів, попередню обробку тексту, тренування, тестування моделі, інтеграцію з Telegram-ботом та оцінювання продуктивності за метриками точності, повноти та F1-міри.

У процесі роботи реалізовано повноцінний модуль на мові Python із використанням бібліотек TensorFlow, Keras, spaCy та інших. Модуль протестовано в інтерактивному середовищі Telegram, де він показав високу точність класифікації (понад 97%) та здатність до узагальнення.

Результати можуть бути використані в системах фактчекінгу, інформаційної безпеки, контент-аналізу, виявлення дезінформації та авторської атрибуції текстів.

Ключові слова: EMBEDDING, LSTM, ГЛИБИННЕ НАВЧАННЯ, КЛАСИФІКАЦІЯ ТЕКСТУ, ОБРОБКА ПРИРОДНОЇ МОВИ, NEURAL NETWORK, TELEGRAM-БОТ.

ANNOTATION

Qualification work on the topic “ Neural network module for text source identification using Embedding and LSTM ” for obtaining a bachelor's degree in the specialty 122 “Computer Science” of the educational program “Artificial Intelligence” is 51 pages long and contains 16 illustrations, 2 tables, and 25 references.

The purpose of the work is to create a module for automatic text source identification based on natural language processing and deep learning methods, in particular Embedding and LSTM technologies, which allow texts to be classified according to their semantic and stylistic features.

The research methods chosen are: analysis of existing approaches to text classification, construction of the LSTM model architecture using vector representation of words, pre-processing of text, training, model testing, integration with the Telegram bot, and performance evaluation based on accuracy, completeness, and F1-measure metrics.

In the course of the work, a full-fledged module was implemented in Python using the TensorFlow, Keras, spaCy, and other libraries. The module was tested in the interactive environment of Telegram, where it showed high classification accuracy (over 97%) and the ability to generalize.

The results can be used in fact-checking, information security, content analysis, disinformation detection, and text authorship attribution systems.

Keywords: EMBEDDING, LSTM, DEEP LEARNING, TEXT CLASSIFICATION, NATURAL LANGUAGE PROCESSING, NEURAL NETWORK, TELEGRAM BOT.

ЗМІСТ

Вступ.....	7
1 Аналіз предметної області та постановка задачі дослідження.....	9
1.1 Опис предметної області	9
1.2 Огляд і аналіз існуючих аналогів визначення джерела тексту за допомогою Embedding і LSTM.....	11
1.3 Постановка задачі дослідження.....	18
2 Алгоритмічне та інформаційне забезпечення модуля.....	19
2.1 Загальна структура проєктованого модуля	19
2.2 Алгоритмічне забезпечення	23
2.3 Інформаційне забезпечення.....	26
3 Програмно-технологічне забезпечення модуля	29
3.1 Реалізація програмного забезпечення	29
3.2 Тренування моделі	32
3.3 Тестування розробленої програми	35
Висновки	38
Список використаних джерел	40
Додаток А Лістинг коду модуля класифікації тексту	43
Додаток Б Лістинг реалізації Telegram-бота для визначення категорії тексту.....	45
Додаток В Апробація отриманих результатів	47

ВСТУП

Актуальність теми. У сучасному світі через Інтернет поширюється дедалі більший масив даних. Як наслідок, виникає потреба в точному визначенні джерела цих даних. Одним з найважливіших питань у цій сфері є проблема пошуку первинного джерела при поширенні тексту. Це питання є актуальним у різних контекстах, наприклад, у виявленні фейкових новин, перевірці контенту, інформаційній безпеці та пошукових системах. Технології обробки природної мови (NLP) є основою для вирішення цієї проблеми, оскільки вони спроможні автоматизувати обробку величезних обсягів текстових даних і виокремлювати корисні ознаки для розпізнавання джерел.

У цьому питанні слід звернути увагу на методи глибокого навчання, такі як Embedding і Long Short-Term Memory (LSTM [3]), вони дозволяють моделювати текст на рівні семантики та контексту. Це дає змогу створювати моделі, які можуть точно класифікувати тексти з джерел, навіть якщо ці тексти мають подібне або навіть однакове значення, але відрізняються за стилем викладу.

Саме тому важливість цього питання стає цілком зрозумілою, адже з розвитком технологій виникає потреба в таких механізмах, які могли б автоматично визначати джерело тексту. Таким чином, інформація стає більш достовірною в середовищі, де кількість даних постійно збільшується.

Метою роботи є розробка нейромережевого модуля, який може автоматично визначати джерело тексту, використовуючи методи обробки природної мови (NLP), зокрема Embedding та LSTM. Модуль повинен правильно класифікувати тексти відповідно до їх семантичних та стилістичних особливостей, що дозволить швидко визначати джерело тексту. Для досягнення цієї мети, в процесі роботи були сформульовані та вирішені наступні завдання:

1. Провести аналіз існуючих підходів до розпізнавання текстових джерел та обробки природної мови.

2. Створити архітектуру модуля класифікації текстів, який використовує методи Embedding та LSTM.

3. Розробити нейромережевий модуль для автоматичної ідентифікації текстових джерел з використанням методів Embedding та LSTM.
4. Випробувати роботу моделі на реальних текстах та проаналізувати її ефективність.
5. Зіставити одержані результати з існуючими рішеннями в цій галузі.

Об'єктом дослідження є процес класифікації текстів на основі алгоритмів обробки тексту, ефективних для точного встановлення його оригіналу.

Предметом дослідження є застосування методів обробки природної мови для класифікації текстів за джерелом. Таким чином, це дослідження технологій Embedding та LSTM, які здатні створювати точні моделі, що можуть аналізувати текстові дані та виявляти особливості кожного джерела.

Методи дослідження. У процесі реалізації роботи застосовано поєднання теоретичних, емпіричних та експериментальних методів. Теоретичний аналіз включав вивчення сучасних наукових джерел щодо методів обробки природної мови, зокрема Embedding та LSTM, і порівняння їхніх можливостей для завдання класифікації джерел тексту. Емпіричні методи передбачали формування датасету, його анотування та передобробку, включно з токенізацією, нормалізацією, видаленням стоп-слів і паддінгом. Експериментальне моделювання було реалізоване через побудову глибокої нейронної мережі на основі Embedding та LSTM, навчання моделі на розмічених даних, валідацію за метриками точності, повноти, а також тестування модулю у Telegram-боті. Використано інструментарій Python, що забезпечило гнучкість, наочність і відтворюваність експериментів.

Результати дослідження опубліковано в матеріалах науково-практичної конференції «Інтелектуальні інформаційні технології в прикладних дослідженнях» (ІІТАР - 2025), м. Тернопіль, 27-29 травня 2025 р.

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ПОСТАНОВКА ЗАДАЧІ ДОСЛІДЖЕННЯ

1.1 Опис предметної області

З розвитком цифрових технологій та інтернету, що стали частиною нашого повсякденного життя, проблема фейкових новин, дезінформації та маніпуляцій є більш актуальною, ніж будь-коли. Важливим кроком у вирішенні цих проблем є визначення того, звідки походить той чи інший текст. Це особливо актуально в таких галузях, як журналістика, реклама, наука, і навіть у правових та економічних ситуаціях. Розпізнавання джерела тексту дозволяє спростити процес перевірки інформації, що допомагає мінімізувати ймовірність поширення неправдивої інформації та підвищує загальний рівень довіри до даних, на які ми покладаємось.

Великою проблемою як для обробки природної мови (NLP), так і для штучного інтелекту є визначення того, звідки походить той чи інший текст. Під цим широким визначенням це завдання можна визначити як розпізнавання та ідентифікацію тексту з метою з'ясування його походження. Така ідентифікація має на меті встановити, від якого автора, організації чи навіть веб-сайту походить текст. Можливість ідентифікації походження тексту значно збагачує процес автоматичного аналізу і дозволяє перевіряти достовірність інформації, уникаючи при цьому шахрайства.

Нейромережеві модулі, особливо моделі глибокого навчання, мають великі перспективи у вирішенні проблеми класифікації текстів на основі їхнього джерела. Ці модулі можуть враховувати контекст і стилістичні елементи тексту, що дає змогу більш точно визначити, звідки він походить.

Невід'ємною частиною цього процесу є технології векторизації тексту, зокрема Word Embeddings, які дозволяють представити тексти у вигляді векторів, що передають семантичну частину інформації про слова та фрази. Ці технології, у поєднанні з рекурентними нейронними мережами (RNN), довготривалою короткочасною пам'яттю (LSTM [3]) та порівнянням подібності векторів,

допомагають класифікувати не лише за темами чи стилями, а й за джерелами з набагато більшою точністю.

Існує кілька причин, таких як різноманітність структур, яких може набувати певна мова, стилістичні відмінності, притаманні певним культурам або регіональним кластерам, а також застосування медіа-механізму змішування інформації з різних джерел, які роблять ідентифікацію джерела тексту дуже складною. Таким чином, для вирішення цієї специфічної проблеми та ефективного вирішення завдань, пов'язаних з виявленням джерел текстів, використовуються різні підходи та методи, що включають в себе різні техніки машинного навчання.

Основні завдання в цій області включають в себе:

- Базовий аналіз лексики, стилістичних аспектів та синтаксичних особливостей тексту. У кожного автора є свій унікальний стиль, особливості якого можуть варіюватися від фіксованих формул до більш поширених лексичних зворотів або варіантів граматичних конструкцій. Ці особливості можна розпізнати і використати для ідентифікації авторства текстів.

- Застосування статистичних методів і методів машинного навчання. Складається з побудови моделей, які дають високу точність класифікації на основі використання навчальних моделей на великих текстових корпусах для ідентифікації джерел. Ці моделі можуть включати алгоритми класифікації, такі як дерева рішень, підтримка векторних машин (SVM) та нейронні мережі.

- Виявлення контексту та метаданих. Поряд з інформацією про самі тексти не менш важливо, якщо не більш, враховувати контекст, в якому вони були опубліковані (URL, автор, дата публікації тощо). Ці дані можуть допомогти ідентифікувати джерело тексту з більшою точністю при класифікації.

Зростання ролі інтернет-простору та соціальних мереж породжує інші виклики, пов'язані з ідентифікацією авторства текстів. Багато текстів публікуються анонімно, або ж не мають чіткого джерела; таким чином, розробка ефективного і надійного методу атрибуції текстового джерела стає все більш важливою. Розроблені технології повинні враховувати як текст, так і його першоджерело в постійно зростаючому цифровому середовищі.

Таким чином, визначення авторства текстів є не лише важливим засобом боротьби з дезінформацією, а й фактором, що сприяє подальшому удосконаленню заходів з автоматичного аналізу контенту. Використання методів глибоких нейронних мереж і векторизації створює можливості для більш ефективних стратегій для вирішення таких проблем.

1.2 Огляд і аналіз існуючих аналогів визначення джерела тексту за допомогою Embedding і LSTM

Цей аналіз містить огляд існуючих методів і технологій, важливих для виявлення текстових джерел, зокрема, моделей Embedding та LSTM [3]. Основна мета - визначити найбільш перспективні та ефективні підходи, які характеризують таку сферу, як ідентифікація джерела тексту, а також відстеження його походження та достовірності. Важливими факторами є функціональність і точність, а також обмеження, які вони мають при використанні в практичних завданнях.

Інструменти для виявлення джерел текстової інформації з використанням LSTM [3] і Embedding відносяться до засобів, за допомогою яких текстова інформація управляється, щоб встановити, з яким джерелом вона пов'язана. Ці методи застосовують прогресивні алгоритми машинного навчання, які ідентифікують закономірності в текстах і зіставляють їх з наявними даними про джерела, зокрема нейронні мережі для обробки природної мови (NLP).

Подібні методи спрямовані на збереження точності ідентифікації джерела конкретного тексту за допомогою суворої оцінки контекстуальних, лексичних і структурних характеристик. Вони дають змогу встановити, чи належить певний текст до певного стилю, пов'язаного з конкретним джерелом.

Цілями цих методів є:

- встановити походження тексту через його лексичні та синтаксичні характеристики;
- пошук стилістичних елементів, які дозволяють порівняти текст з конкретними організаційними або авторськими елементами;

- побудувати модель, яка може успішно класифікувати тексти за різними джерелами.

Зупинимось на аналізі характеристики існуючих методів та систем (таблиця 1.1), які використовують Embedding та LSTM [3] для визначення джерела тексту:

1. DeepMoji [1] – це модель, яка використовує нейронні мережі, що аналізують текст на основі класифікації емоційності тексту, пов'язуючи її з описами, які використовуються як вхідні ознаки. Модель використовує технологію Embedding для створення зображень для кожного слова і фрази у векторному форматі, що дозволяє класифікувати стилі на засадах емоційного змісту тексту. Модель продемонструвала хороші результати на різних завданнях, які передбачають вивчення емоційної тональності текстів, що є дуже актуальним для аналізу джерел.

2. BERT [2] (Bidirectional Encoder Representations from Transformers) – як відомо, одна з найвідоміших моделей обробки природної мови, яка використовує методику контекстного вбудовування для визначення інформації про джерела текстів. Вона враховує контекст для кожного слова в реченні і може допомогти точно ідентифікувати джерела на основі їхнього контексту та стилістики написання.

3. LSTM [3] (Long Short-Term Memory) – це тип рекурентної нейронної мережі, яка використовується для моделювання тексту та вивчення залежних зв'язків між словами. LSTM [3] може аналізувати більш складні зв'язки в тексті та допомагає знайти його джерело, вивчаючи послідовності слів та їхній контекст. LSTM [3] має одну з переваг перед аналізом тексту з довгими залежностями, а саме: він може ідентифікувати джерело, навіть якщо структура тексту є складною.

4. FastText [4] – це ще один інструмент для виявлення джерел тексту, який працює на принципах векторного представлення слів. FastText [4] дозволяє створювати ефективні вставки для окремих слів і фраз і застосовувати їх для класифікації текстів за джерелом. Хоча ця модель не така потужна, як BERT [2] або LSTM [3], вона допомагає швидко отримувати результати в реальних програмах.

Таблиця 1.1 – Аналіз існуючих сервісів визначення фейкових новин

Метод	Тип	Основний підхід	Переваги	Обмеження
DeepMoji	Embedding	Аналіз емоційного забарвлення	Добре працює з емоційними текстами	Обмежена точність для складних текстів
BERT	Transformers	Контекстуальне Embedding	Висока точність на контекстуальних задачах	Високі вимоги до ресурсів
LSTM	Рекурентні мережі	Моделювання послідовностей тексту	Підходить для складних текстів	Не завжди ефективний для коротких текстів
FastText	Embedding	Векторні представлення слів	Швидке навчання і обробка	Менша точність у порівнянні з BERT

Відповідно, ці методи мають свої переваги та недоліки при практичному використанні для виявлення текстових джерел. Для практичного використання в завданнях виявлення джерел тексту необхідно вибрати належну модель відповідно до умов текстових даних і вимог до точності аналізу.

У сьогоdnішньому інформаційному середовищі існує безліч інструментів і ресурсів, призначених для виявлення та протидії фейковим новинам. Нижче наведено основні з них:

1. StopFake [5] (рисунок 1.1) - є проектом української медійної громадської організації «Центр медіареформ». Він був заснований у березні 2014 року українськими професорами та студентами з проголошеною метою спростування російської пропаганди та фейкових новин. Він розпочинав свою діяльність як російськомовна та англомовна організація, яка виросла до телевізійного шоу, що транслюється на 30 місцевих каналах, щотижневої радіопередачі та потужної аудиторії в соціальних мережах.

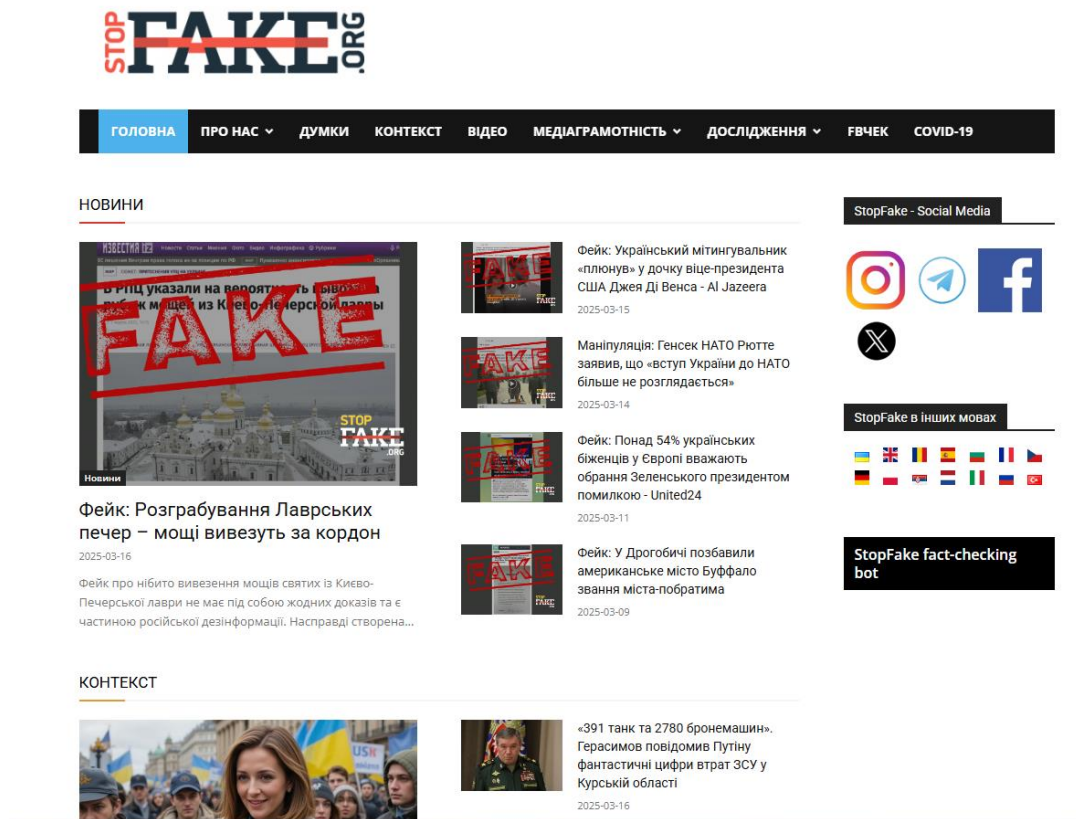


Рисунок 1.1 – Головна сторінка StopFake [5].org

Основні функції та напрямки роботи StopFake [5]:

- фактчекінг новин - перевірка правдивості інформації, виявлення фейкових матеріалів у ЗМІ та соціальних мережах;
- аналітика - дослідження механізмів виникнення інформації та виявлення основних тенденцій у поширенні неправдивої інформації;
- міжнародна співпраця - участь у міжнародних ініціативах по боротьбі з фейками та дезінформацією;
- навчання та тренінги - проведення освітніх кампаній для журналістів, активістів та громадськості, присвячених фактчекінгу та інформаційній гігієні;
- мультимовність - контент StopFake [5] доступний 13 мовами, а отже, охоплює велику аудиторію.

2. Реєстр фейкових ЗМІ України (рисунок 1.2) – це інформаційний сервіс, створений для протидії фейкам. Його мета полягає у викритті та маркуванні оманливих онлайн-платформ, які впливають на громадськість, або діють як так звані «інфосмітники». Реєстр призначений для виявлення та маркування

ненадійних джерел, які регулярно публікують фейкові новини та маніпулятивні публікації. Сервіс допомагає користувачеві відрізнити надійні медіа від ненадійних, що дозволяє захистити інформаційний простір від дезінформації та інформаційних атак. Крім того, він навчає фактчекінгу та розпізнаванню неправдивих новин через навчання.

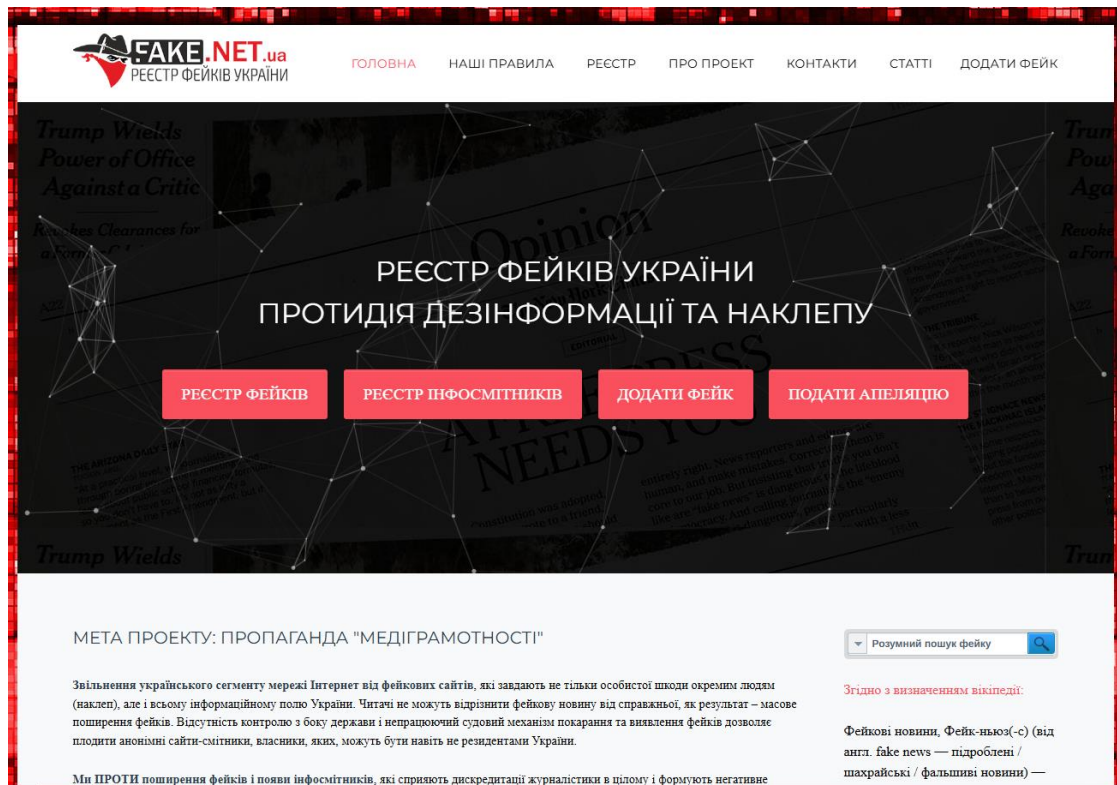


Рисунок 1.2 – Головна сторінка fake.net.ua

3. Check-It [10] (рисунок 1.3) - це плагін для браузера, який дозволяє перевіряти заголовки та основний зміст матеріалів. Призначений для автоматизації аналізу та виявлення фейкових новин. Він дозволяє перевіряти достовірність інформації, яку ви бачите в Інтернеті. Плагін перевіряє веб-сайти, порівнюючи їх з базами даних достовірних і сумнівних джерел. Якщо система виявить щось підозріле, вона попереджає користувача і запропонує різні надійні джерела. За допомогою цього швидкого плагіна можна перевірити достовірність інформації безпосередньо в браузері. Для перевірки новин не доведеться заходити на інші сайти.

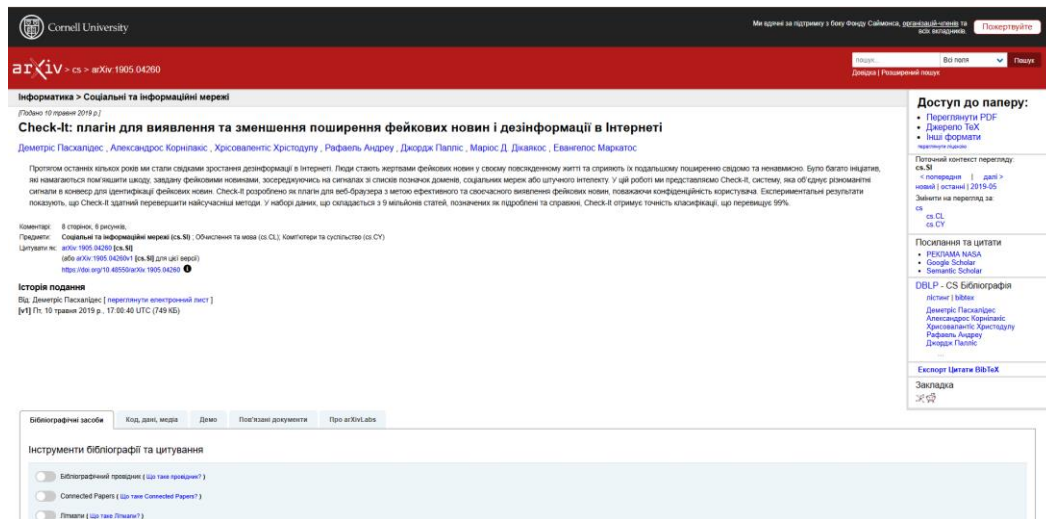


Рисунок 1.3 – Головна сторінка Check-It [10]

4. FactCheck.org (рисунок 1.4) - незалежна неприбуткова організація, яка займається перевіркою фактів і боротьбою з неправдивими відомостями в засобах масової інформації, а також у соціальних мережах і політичних заявах. Організація була заснована в 2003 році в Центрі публічної політики Анненберга при Університеті Пенсильванії.

Чим займається FactCheck.org :

- перевірка політичних заяв: Організація перевіряє, чи відповідають дійсності заяви політиків, чиновників та інших громадських діячів;
- спростування міфів і фейків: окрім політичних заяв, FactCheck.org [6] також перевіряє фальшиві новини, які поширюються в соціальних мережах;
- науковий фактчекінг: окрема категорія перевірок присвячена дослідженню наукових тверджень, особливо тих, що стосуються здоров'я, екології та технологій;
- взаємодія з користувачами: Користувачі можуть надсилати свої запити на перевірку будь-яких підозрілих тверджень, на які вони натрапили;
- співпраця із соціальними платформами: організація співпрацює з Facebook та іншими платформами, щоб відзначати сумнівні матеріали та боротися з дезінформацією.



Рисунок 1.4 – Головна сторінка factcheck.org

Таким чином, проведений аналіз сучасних підходів до визначення джерела тексту із використанням моделей Embedding, LSTM та інших алгоритмів машинного навчання засвідчив, що кожен із розглянутих методів має свої переваги й обмеження залежно від специфіки даних та поставлених завдань. Зокрема, моделі на основі Embedding і LSTM демонструють високу точність для складних і довгих текстів, а BERT забезпечує найкращі результати при обробці контекстуально насичених даних, хоча і потребує значних обчислювальних ресурсів. Інструменти типу FastText залишаються ефективними для оперативної класифікації, особливо у випадках обмежених ресурсів. Практичне застосування цих моделей, а також додаткових сервісів для боротьби з фейками, таких як StopFake, Check-It чи FactCheck.org, дозволяє суттєво підвищити якість і достовірність ідентифікації джерел текстової інформації. Водночас для досягнення максимальної ефективності у задачах визначення джерела тексту доцільно здійснювати вибір методів із урахуванням балансу між точністю, швидкістю обробки та ресурсомісткістю моделей.

1.3 Постановка задачі дослідження

Проведений аналіз показав, що ідентифікація першоджерела тексту є дуже актуальною в умовах зростання цифрового контенту та необхідності встановлення авторської приналежності. Методи, що уже існують для визначення джерела тексту часто можуть є неточними у випадках зі складними текстовими структурами та багатомовним контентом.

Одним з найбільш ефективних підходів є використання нейромережових методів, зокрема векторизація тексту за допомогою Embedding та глибоких рекурентних мереж (LSTM [3]). Цей підхід дозволяє моделювати семантичні особливості тексту та визначати закономірності, характерні для конкретного джерела.

Отже, метою роботи є розробка нейромережового модуля, який може автоматично визначати джерело тексту, використовуючи методи обробки природної мови (NLP), зокрема Embedding та LSTM. Модуль повинен правильно класифікувати тексти відповідно до їх семантичних та стилістичних особливостей, що дозволить швидко визначати джерело тексту. Для досягнення цієї мети, в процесі роботи були сформульовані та вирішені наступні завдання:

- провести аналіз існуючих підходів до розпізнавання текстових джерел та обробки природної мови;
- створити архітектуру модуля класифікації текстів, який використовує методи Embedding та LSTM;
- розробити нейромережовий модуль для автоматичної ідентифікації текстових джерел з використанням методів Embedding та LSTM;
- випробувати роботу моделі на реальних текстах та проаналізувати її ефективність;
- зіставити одержані результати з існуючими рішеннями в цій галузі.

2 АЛГОРИТМІЧНЕ ТА ІНФОРМАЦІЙНЕ ЗАБЕЗПЕЧЕННЯ МОДУЛЯ

2.1 Загальна структура проектного модуля

Модуль працює на основі глибоких нейронних мереж і, зокрема, моделей на основі Embedding та LSTM (таблиця 2.1). Реалізація цих моделей буде здійснюватися за допомогою TensorFlow [7] та Keras [8]. Головна ідея полягає у визначенні джерела тексту шляхом аналізу його лексичних та стилістичних особливостей. Для досягнення даної мети необхідно описати функціональну структуру системи з переліком основних функцій та їх характеристиками. Необхідно також описати інформаційні зв'язки між компонентами системи, а також зв'язки з іншими системами та зовнішніми установами, які є джерелами або споживачами цієї інформації.

Ключовими елементами функціональної структури проекту модуля, що розробляється, є: модуль збору та попередньої обробки текстових даних, модуль векторизації тексту (Embedding), модуль навчання LSTM [3]-моделі, модуль оцінювання моделі та модуль прогнозування джерела тексту. Кожен з цих компонентів виконує певні функції, які сприяють результативному аналізу текстових даних.

Таблиця 2.1 – Модулі системи обробки текстових даних

Модуль	Функції
Збір та попередня обробка	Отримання текстових даних, токенизація, нормалізація, очищення.
Векторизація (Embedding)	Перетворення тексту в числове представлення
Навчання LSTM-моделі	Аналіз послідовностей тексту, визначення його джерела
Оцінка продуктивності моделі	Перевірка точності, обчислення F1-score, Recall, Precision
Прогнозування	Використання навченого модуля для визначення ймовірного джерела тексту

Зв'язки між компонентами системи дають змогу забезпечити безперебійну передачу даних протягом усього процесу. Наприклад, модуль збору даних надсилає

очищені текстові дані до модуля векторизації. Той, у свою чергу, надсилає векторизовані дані до модуля навчання LSTM [3]. Після навчання отримані результати передаються до модуля оцінювання. У подальшому модуль прогнозування використовує оцінену модель для перевірки нових текстів, щоб визначити їхню приналежність до того чи іншого джерела.

Зв'язки з іншими системами та зовнішніми об'єктами включають взаємодію з джерелами текстових даних (корпусами текстів, різних авторів/платформ), а також з користувачами системи, які отримують результати класифікації.

Для створення модуля, який визначатиме джерело тексту за допомогою неймережових технологій, було розроблено діаграму потоку даних. DFD-діаграма (рисунок 2.1) ілюструє основні етапи процесу.

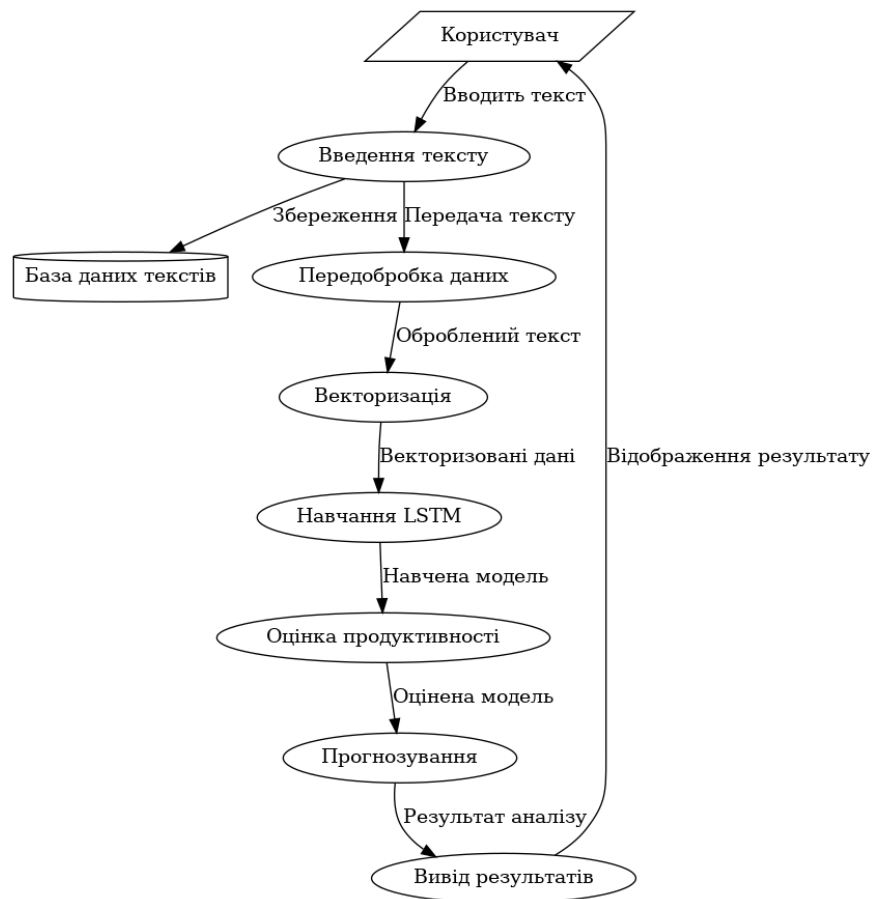


Рисунок 2.1 - Архітектура аналізу тексту

Така структура гарантує ефективну роботу модуля, який визначає джерело тексту за рахунок обробки природної мови та глибокого навчання. Структурна

функціональність модуля охоплює всі необхідні етапи від введення тексту до виведення прогнозу. Таким чином, можна побудувати ефективну систему розпізнавання джерела тексту.

1. Введення тексту - користувач вводить фрагмент тексту через інтерфейс.
2. Передобробка даних - прибирання зайвих символів, нормалізація та токенизація.
3. Векторизація - перетворення тексту в числове подання.
4. Навчання LSTM [3] - навчання моделі на розмічених даних.
5. Оцінка продуктивності - тестування точності моделі.
6. Прогнозування - з'ясування походження тексту.

Розглянемо як складається система, її ключові частини, як вони працюють разом, і як все працює, виходячи з того, що потрібно зробити. Такий підхід дозволяє будувати програмне забезпечення принципом. Це полегшує виправлення певних частин системи, допомагає нам масштабувати проект і краще використовувати обчислювальні ресурси.

На першому етапі розробки було сформовано UML-діаграму класів (рисунок 2.2), яка відображає логічну структуру основних компонентів програмної системи, необхідної для реалізації задачі визначення джерела тексту з використанням глибокого навчання.

Наведена (див. рисунок 2.2) діаграма відображає структуру компонентів та їх взаємодію в рамках системи визначення джерела тексту на основі глибинного навчання з використанням LSTM [3]-моделі. Діаграма включає такі класи: TextSourceDetector, LSTMModel, TextPreprocessor, Tokenizer та TrainingData, кожен з яких виконує окрему функцію у загальній архітектурі системи.

Центральним класом є TextSourceDetector, який відповідає за визначення джерела тексту. Основним методом цього класу є predict, що приймає текст у вигляді рядка та повертає передбачене джерело. Для здійснення передбачення клас використовує інші компоненти: попередню обробку тексту, токенизацію та модель LSTM [3].

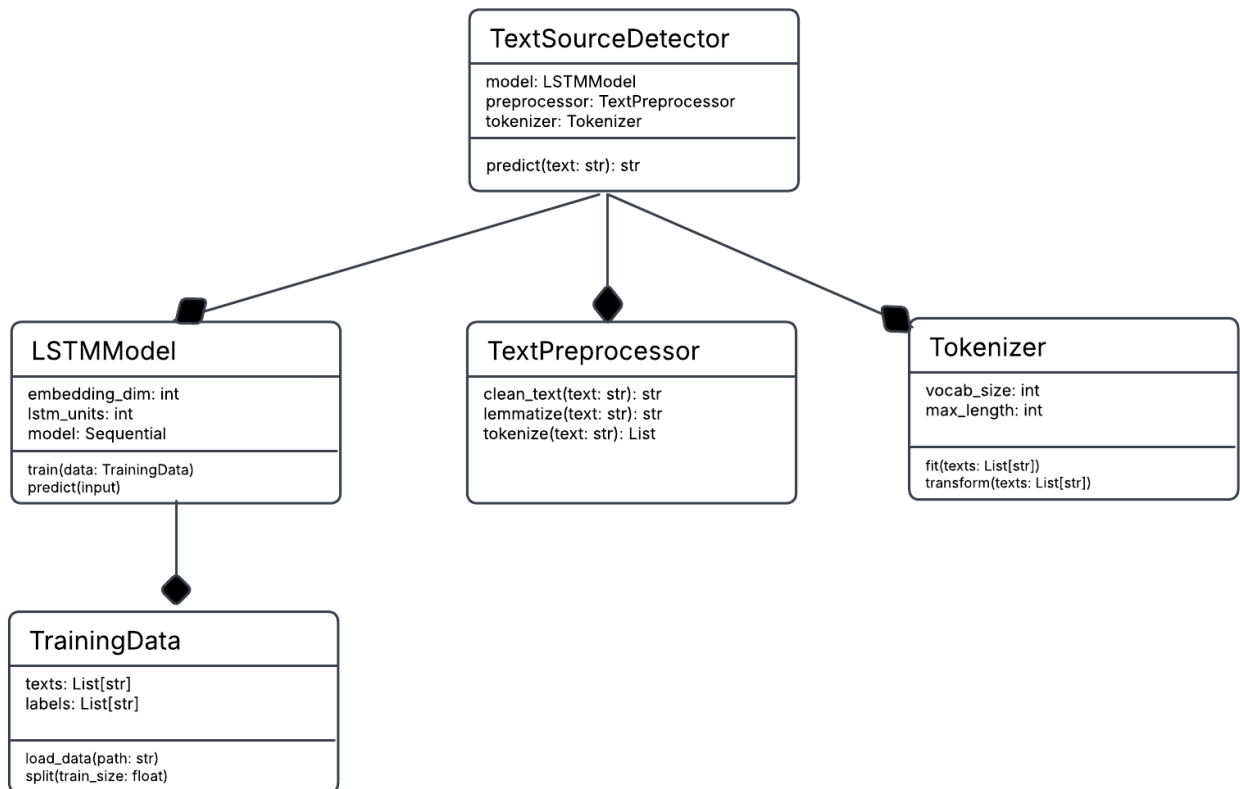


Рисунок 2.2 – UML-діаграма класів

Клас **LSTMModel** реалізує нейронну мережу з використанням LSTM [3]-архітектури. Він містить параметри `embedding`-вимірності, кількість LSTM [3]-одиниць та об'єкт моделі `Sequential`. Основними методами класу є `train`, що приймає об'єкт класу **TrainingData** для навчання, та `predict` — для здійснення передбачень на основі вхідних даних.

Клас **TextPreprocessor** відповідає за попередню обробку тексту, включаючи очищення, лематизацію та токенизацію. Його методи `clean_text`, `lemmatize` та `tokenize` перетворюють текст у формат, придатний для подальшої обробки.

Клас **Tokenizer** виконує перетворення текстів у числові послідовності для подання в модель. Він містить параметри `vocab_size` та `max_length` і реалізує методи `fit` — для побудови словника на основі вхідних текстів, та `transform` — для перетворення текстів у числові представлення.

Клас **TrainingData** зберігає вхідні тексти та відповідні мітки джерел. Він реалізує методи `load_data` для завантаження даних з файлу та `split` для поділу вибірки на тренувальну та тестову.

Отже, попередня діаграма UML-класів описує взаємозв'язки між класами модуля, який реалізує механізм розпізнавання текстового джерела. Кожен клас виконує певну роботу і дозволяє крок за кроком створити хороший та ефективний процес обробки тексту, представлення його в моделі та прогнозування на основі навчених параметрів.

2.2 Алгоритмічне забезпечення

Представлено схему на рисунку 2.3, яка представляє етапи обробки вхідних даних, створення та навчання нейромережевого модуля для виявлення джерел текстів. На блок-схемі представлено кілька важливих блоків, які демонструють етапи роботи системи.

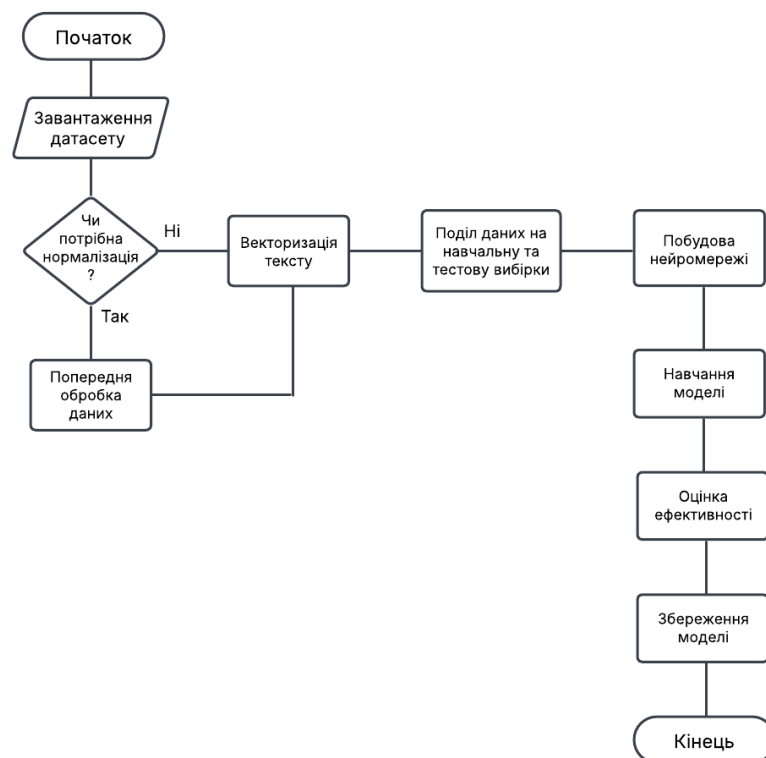


Рисунок 2.3 - Схема алгоритму створення моделі

На першому блоці завантажуються всі необхідні дані для навчання та тестування моделі. Ці дані являють собою збірку текстових документів із

зазначенням джерела, з якого вони були отримані. Дані завантажуються з файлу або бази даних і спочатку перевіряються на коректність та проходять первинну перевірку. Якщо виявляються порожні значення або дублікати, вони видаляються. Якщо дані в цілому коректні, етап пропускається.

Після успішного завантаження даних система переходить до етапу попередньої обробки. Сюди входить нормалізація тексту, що передбачає:

- переведення тексту в нижній регістр;
- видалення всіх стоп-слів (неінформативних слів, таких як сполучники та прийменники);
- лематизація (зведення слів до простої/основної форми).

Після цього текст векторизується, тобто перетворює слова в числові представлення. Для цього використовується метод Embedding, який генерує багатовимірне представлення кожного слова, що дозволяє нейромережі краще аналізувати семантичні складові кожного слова в тексті.

Наступний етап складається з поділу даних на навчальну та тестову вибірки. Навчальна вибірка використовується для навчання моделі, в той час як тестова вибірка використовується для вимірювання продуктивності моделі. Після підготовки даних будується архітектура нейронної мережі. Ця архітектура містить кілька шарів:

- вхідний шар, який приймає векторизовані тексти;
- приховані шари з LSTM [3]-нейронами, які аналізують контекст тексту;
- вихідний шар, який класифікує текст за походженням.

Наступний етап передбачає початок навчання моделі. На даному етапі для оптимізації використовується алгоритм Adam, а функція втрат вимірює різницю між прогнозованими та фактичними значеннями. Модель навчається у декілька ітерацій (епох), поки не досягне найкращої/оптимальної точності. Як тільки модель навчена, її оцінюють за допомогою тестових даних. Для цього використовуються вибрані метрики (точність, повнота, F1-міра), і таким чином спостерігається, наскільки добре модель класифікувала тексти.

Кінцевим результатом навчання є модель, яку, відповідно, можна зберегти у файлі. Файл і модель можуть бути згодом завантажені в програмне забезпечення, яке буде використовувати параметри моделі для автоматичного прогнозування текстового джерела. У підсумку, наведена вище блок-схема ілюструє всі необхідні кроки для створення надійної системи, починаючи від завантаження даних і закінчуючи збереженням навченої моделі.

Далі розглянемо алгоритм виявлення джерела тексту на основі даних введених користувачем (рисунок 2.4).

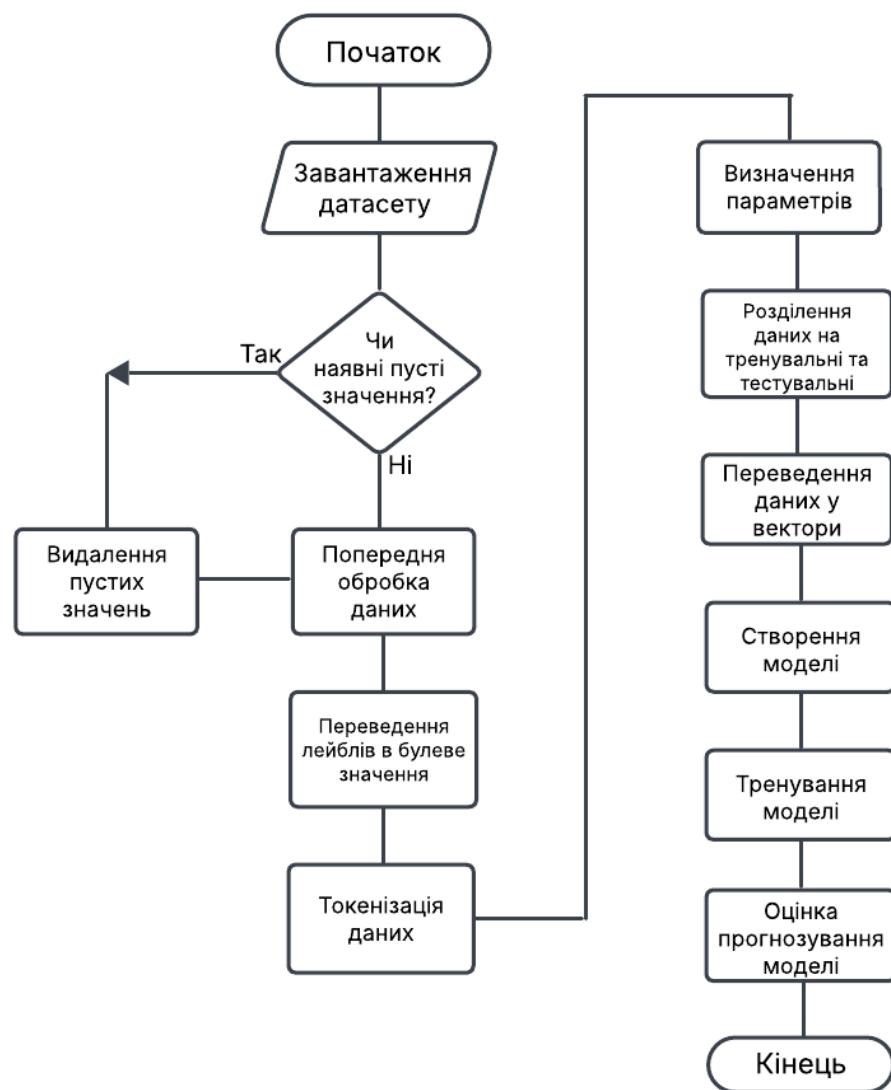


Рисунок 2.4 - Схема алгоритму знаходження джерела тексту

Алгоритм створення нейромережевого модуля для ідентифікації текстових джерел складається з кількох важливих кроків, які забезпечують ефективність навчання нейромережі та її подальшого оцінювання. Спочатку завантажуються необхідні дані, які складаються з текстових файлів, баз даних або інших систем зберігання, що містять текстові записи. Після цього дані перевіряються, що може включати виявлення відсутніх значень, неточностей або аномалій, які можуть ускладнити процес.

Наступний крок у процесі має назву попередня обробка даних, яка передбачає очищення тексту від зайвих елементів, таких як стоп-слова, розділові знаки та спеціальні символи. Цей крок важливий, оскільки необхідно нормалізувати текст і привести його до єдиного формату, який буде зручним для векторизації. У цьому випадку векторизація - це наступний крок, на якому текстові дані конвертуються в числовий формат. Таким чином можна представити текст у вигляді векторів у зручному форматі для подальшої обробки нейронною мережею.

Після векторизації дані будуть розділені на навчальну та тестову вибірки. Це важливий крок, оскільки він забезпечує ефективне навчання моделі. Наступним кроком буде створення архітектури нейронної мережі, тобто моделі довготривалої короткочасної пам'яті (LSTM [3]), оскільки це ефективна модель для роботи з послідовними даними (наприклад, текстом). Отримавши архітектуру LSTM [3]-моделі, необхідно провести її навчання.

Після того, як результати будуть задовільні, модель буде збережена, і її можна буде завантажити пізніше для нового аналізу без необхідності повторного навчання. Цей процес забезпечує ефективність і зручність використання моделі в реальних умовах.

2.3 Інформаційне забезпечення

Для організації та керування даними в межах створеного нейромережевого модуля для виявлення текстових джерел було розроблено ERD діаграму «суть-зв'язок» (рисунк 2.5). Ця діаграма має чотири основні сутності : TEXT_SOURCE,

PROCESSED_TEXT, MODEL і PREDICTION, а також зв'язки між ними, які представляють процес аналізу та класифікації тексту.

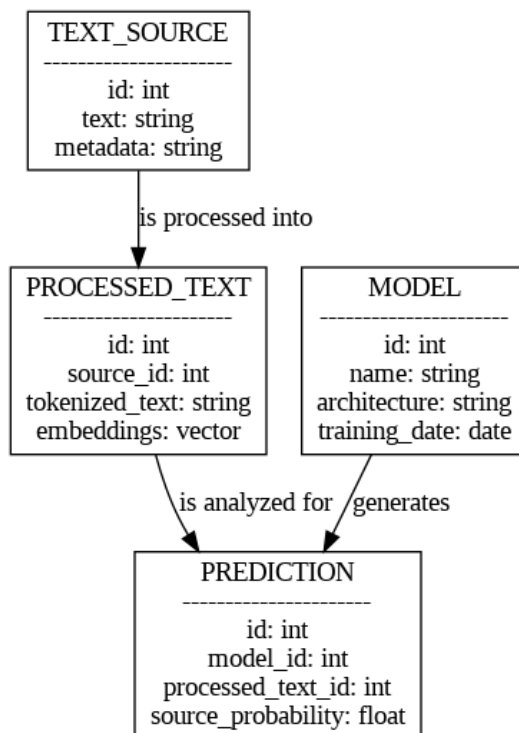


Рисунок 2.5 - Діаграма «суть-зв'язок» (ERD)

Основна сутність системи - це «TEXT_SOURCE», яка містить текст для аналізу, що має власний унікальний ідентифікатор у базі даних для однозначного відстеження цього тексту. Ця сутність також містить метадані, включаючи назву джерела, автора (за наявності) та дату публікації. Ця інформація дозволяє системі дослідити походження тексту та будь-які потенційні посилання на джерело.

Після того, як текст було попередньо оброблено, його представлення зберігається в об'єкті з назвою «PROCESSED_TEXT». Сутність містить токенований текст, очищений від стоп-слів і перетворений на числовий вектор. Такий тип представлення текстових даних є необхідним для їх подальшої обробки за допомогою моделі машинного навчання. Для кожного запису в «PROCESSED_TEXT» дуже важливо мати посилання на відповідний запис «TEXT_SOURCE», щоб гарантувати цілісність і довести відстеження обробленої інформації.

Для аналізу тексту використовується певна нейромережева модель, яка представлена сутністю «MODEL». Кожна з моделей має свій унікальний ідентифікатор, назву, тип моделі (наприклад, LSTM [3]), дату її навчання і, що найважливіше, гіперпараметри, які використовувалися для її навчання. Це дуже важливо, тому що це означає, що можна зберігати і порівнювати моделі різних типів або варіанти гіперпараметрів однієї і тієї ж моделі, а також визначати найбільш ефективні для вирішення поставленого завдання.

Результати процесу класифікації зберігаються в сутності «PREDICTION», яка містить інформацію про прогнозоване джерело тексту, рівень впевненості моделі у прогнозі та дату прогнозування. Сутність «PREDICTION» має унікальну ідентифікацію, посилання на використану модель (MODEL_ID) та оброблений текст (PROCESSED_TEXT_ID). Це дозволяє вести облік попередніх аналізів і оцінювати ефективність моделі на реальних даних.

В цілому, інформаційна структура, викладена в цьому розділі, забезпечує відповідну та ефективну організацію та обробку необхідних текстових даних, які потрібні для навчання нейромережевого модуля. Зв'язки між сутностями забезпечують логічну цілісність даних, а також можливість удосконалення моделі на основі історичних прогнозів.

3 ПРОГРАМНО-ТЕХНОЛОГІЧНЕ ЗАБЕЗПЕЧЕННЯ МОДУЛЯ

3.1 Реалізація програмного забезпечення

Для реалізації програмного модуля використано мову програмування Python з огляду на її популярність у сфері обробки природної мови та побудови нейронних мереж, гнучкість, простий синтаксис та наявність потужних спеціалізованих бібліотек (зокрема, TensorFlow [7], Keras [8], NumPy, NLTK, spaCy [9]). Ці бібліотеки значно спростили попередню обробку текстових даних та реалізацію складних архітектур глибокого навчання.

Python підтримує об'єктно-орієнтовану парадигму, автоматичне керування пам'яттю, а також має зручні структури даних (списки, множини, словники, кортежі), що робить процес розробки модулів ефективним та масштабованим. Завдяки платформенній незалежності інтерпретатора Python, розроблене рішення може виконуватись на будь-якій операційній системі.

Код на Python, як правило, більш компактний, ніж на інших мовах, таких як Java, що дозволяє зосередитися на логіці моделі, а не на реалізації дрібних технічних деталей коду.

Розроблений програмний модуль умовно поділяється на наступні компоненти:

- модуль завантаження та фільтрації даних;
- модуль попередньої обробки текстових входів;
- модуль векторизації (embedding);
- модуль побудови LSTM [3]-архітектури;
- модуль визначення джерела тексту (класифікації).

Модуль підготовки даних забезпечує базову обробку текстових даних, наприклад, видалення спеціальних символів, нижнього регістру, токенизацію, лематизацію (рисунк 3.1), видалення стоп-слів. Ці процеси дозволяють нормалізувати текстові дані та підвищити якість подальшої обробки.

```
0    ukrainian president say country forgive forget...
1    jeremy bowen frontline irpin resident come rus...
2    world big fertiliser firm say conflict deliver...
3    parent manchester arena bombing young victim s...
4    consumer feel impact high energy cost fuel pri...
Name: lemmatized_description, dtype: object
```

Рисунок 3.1 – Приклад новини з датасету

Модуль вбудовування здатний перетворювати мовну інформацію в числове представлення у форматі, придатному для нейронної мережі. Кожне слово або словосполучення представляється у вигляді вектора інформації, що відноситься до його семантики. Реалізація використовує шар Embedding з бібліотеки Keras [8], який дозволяє використовувати векторне представлення безпосередньо під час навчання моделі.

Модуль LSTM [3] базується на рекурентній нейронній мережі, що містить механізм довготривалої короткочасної пам'яті для збереження контексту тексту з урахуванням порядку слів при моделюванні атрибуції джерела. Ця архітектура дозволяє послідовне проходження шарів LSTM [3], потенційно з шарами Dropout, щоб уникнути перенастроювання.

Модуль класифікації відповідає за аналіз отриманих векторів і на їх основі вгадує джерело тексту. Це реалізується за допомогою щільного шару (Dense) з функцією активації softmax у фінальному шарі, що дозволяє здійснювати багатокласову класифікацію.

Таким чином, комбінація модулів дає нам повний цикл обробки тексту - від завантаження вхідних даних до отримання результату в джерелі походження.

Також було здійснено вирівнювання (падінг) числових послідовностей, отриманих після токенізації текстів. Оскільки тексти новин були різної довжини, необхідно було, щоб усі послідовності були однакової довжини перед введенням у модель. Для цього було використано функцію pad_sequences із бібліотеки tensorflow.keras, яка додає нульові значення (padding) до кінця коротших послідовностей (padding='post'), або обрізає довші послідовності до фіксованої довжини maxlen.

Після виконання паддінгу усі вхідні послідовності мають однакову довжину, що є обов'язковою умовою для роботи нейронної мережі. Приклад результату вирівнювання послідовностей продемонстровано на рисунку 3.2.



```
After padding, x_pad shape: (42115, 100)
```

Рисунок 3.2 – Результат паддінгу токенизованих послідовностей

Для класифікації тематичних груп новинних текстів проведено кластеризацію за допомогою методу К-середніх. Для класифікації тематичних груп новинних текстів проведено кластеризацію за допомогою методу К-середніх. Тексти новин були попередньо векторизовані у форматі TF-IDF, що дозволяє визначити вагу кожного слова, відповідно до частоти вживання цього слова в документі та частоти в усьому масиві.

Далі було застосовано алгоритм KMeans з кількістю кластерів $k=5$. Це дозволило згрупувати новини за змістовними ознаками. Результатом кластеризації стала нова колонка cluster у таблиці новин, яка відображає кластерну належність кожного заголовка (рисунок 3.3).



	title	cluster
0	Ukraine: Angry Zelensky vows to punish Russian...	1
1	War in Ukraine: Taking cover in a town under a...	1
2	Ukraine war 'catastrophic for global food'	2
3	Manchester Arena bombing: Saffie Roussos's par...	1
4	Ukraine conflict: Oil price soars to highest l...	1

Рисунок 3.3 – Результат кластеризації новин за допомогою алгоритму KMeans

Для візуального представлення результатів кластеризації використано метод PCA (Principal Component Analysis), оскільки PCA - це підхід до зменшення розмірності, який полегшує відображення багатовимірних векторів новин у

двовимірному просторі. Це дозволило наочно представити розподіл документів за кластерами.

На рисунку 3.4 представлено розподіл новин за кластерами після кластеризації методом KMeans. Кожна точка на графіку відповідає окремому тексту новини, а її колір позначає кластер, до якого цей текст було віднесено. Візуалізація дає змогу оцінити щільність, розподіл та наявність можливих перетинів між кластерами.

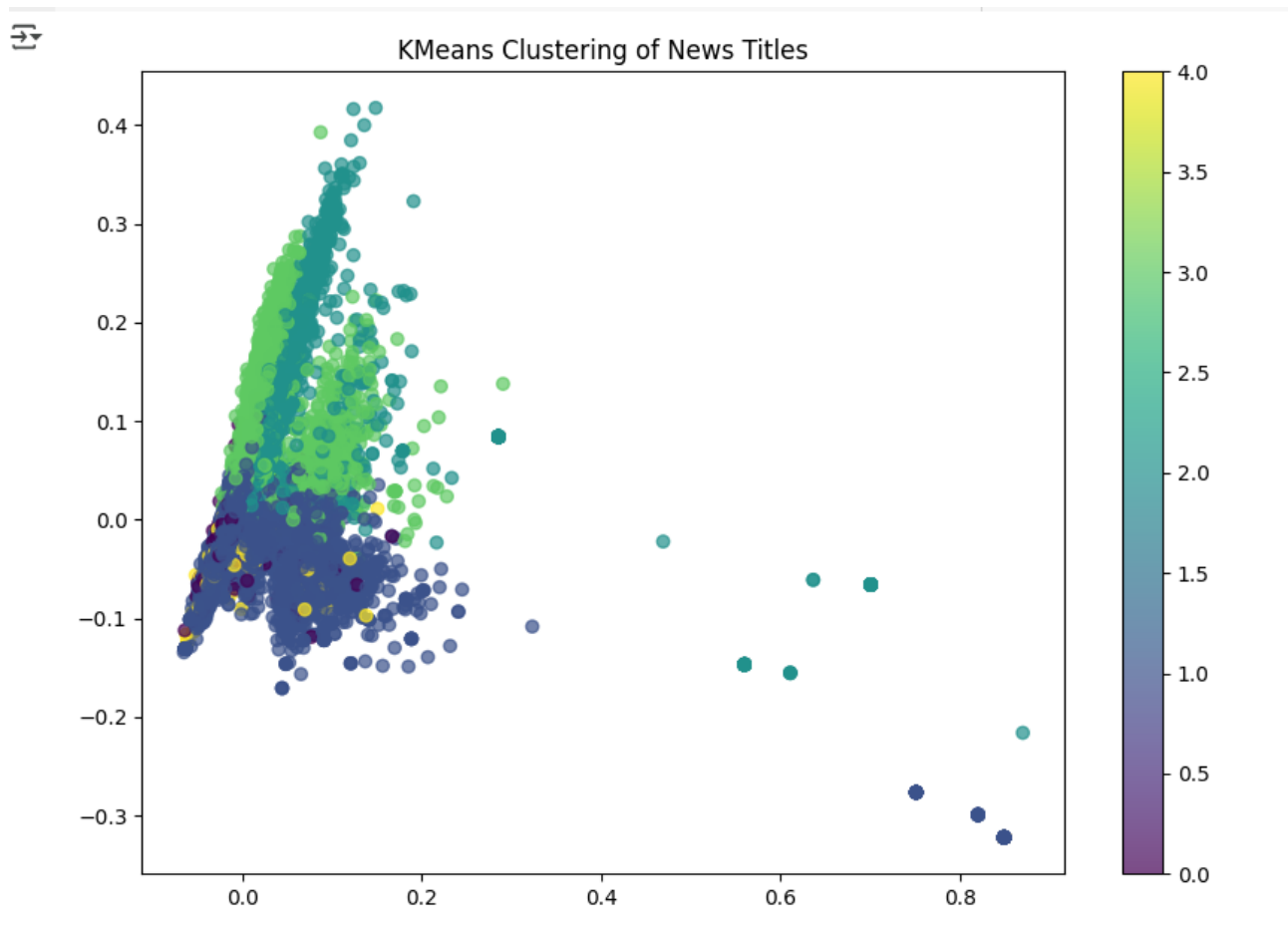


Рисунок 3.4 – Візуалізація кластеризації новин у двовимірному просторі методом PCA

3.2 Тренування моделі

Для навчання моделі було обрано архітектуру нейронної мережі з рекурентним шаром типу LSTM [3] (Long Short-Term Memory), яка добре

zareкомендувала себе в задачах обробки природної мови. До складу моделі входить шар векторного представлення слів (Embedding), шар LSTM [3] та щільний вихідний шар із функцією активації softmax для багатокласової класифікації.

В рамках практичної частини була побудована нейромережева модель з використанням архітектури LSTM [3] для класифікації текстів за джерелом або автором. Для реалізації моделі було використано бібліотеку TensorFlow [7] у поєднанні з її високорівневим API Keras [8]. Архітектура моделі має структуру, що складається з наступних елементів:

Embedding-шар – призначений для перетворення індексів слів у щільні вектори фіксованої розмірності (EMBEDDING_DIM = 100). Це дозволяє передавати у мережу не просто номери слів, а їх векторні представлення, що враховують семантичну подібність;

- SpatialDropout1D – регуляризуючий шар, який випадковим чином замінює цілі вектори в матриці ознак нульовим вектором, зменшуючи перенавчання (ризик якого великий у результатах NLP);

- Два LSTM [3]-шари – перший з параметром return_sequences=True, що дозволяє передати повну послідовність далі у другий LSTM [3]. Ці шари забезпечують збереження контексту в послідовності тексту;

- Dropout-шари – використовуються після кожного LSTM [3] та Dense-шару для регуляризації моделі;

- Dense-шар із 64 нейронами та активацією ReLU – проміжний повнозв'язаний шар;

- Вихідний Dense-шар з активацією softmax – забезпечує багатокласову класифікацію згідно з кількістю унікальних джерел тексту (NUM_CLASSES).

Гіперпараметри моделі наведено нижче:

- максимальна кількість унікальних слів: MAX_NB_WORDS = 20000;

- максимальна довжина послідовності (тексту):
MAX_SEQUENCE_LENGTH = 300;

- розмір векторного представлення слів: EMBEDDING_DIM = 100;

- кількість нейронів у LSTM [3]: 128 і 64 відповідно;

- Dropout: 0.2;
- функція втрат: `sparse_categorical_crossentropy`;
- оптимізатор: `adam`;
- метрика оцінки якості: `accuracy`.

Після побудови модель було скомпільовано та сформовано її загальну структуру. Підсумкова архітектура відображена на рисунку 3.5.

Layer (type)	Output Shape	Param #
embedding (Embedding)	(None, 300, 100)	2,000,100
spatial_dropout1d (SpatialDropout1D)	(None, 300, 100)	0
lstm (LSTM)	(None, 300, 128)	117,248
dropout (Dropout)	(None, 300, 128)	0
lstm_1 (LSTM)	(None, 64)	49,408
dropout_1 (Dropout)	(None, 64)	0
dense (Dense)	(None, 64)	4,160
dropout_2 (Dropout)	(None, 64)	0
dense_1 (Dense)	(None, 5)	325

Total params: 2,171,241 (8.28 MB)
 Trainable params: 2,171,241 (8.28 MB)
 Non-trainable params: 0 (0.00 B)

Рисунок 3.5 – Архітектура моделі LSTM [3] для класифікації текстів

Навчання проводилось з використанням підготовлених послідовностей слів у вигляді векторизованих вхідних даних. У якості цільових міток використовувались категорії джерел текстів.

У ході тренування модель демонструвала поступове зростання точності (рисунк 3.6) на тренувальній вибірці, що свідчить про ефективне засвоєння ознак, притаманних кожному джерелу. Водночас, валідаційна точність зберігалася на стабільно високому рівні протягом усього процесу навчання, що вказує на здатність моделі до узагальнення.

Після завершення етапу навчання модель була збережена та використана у виробничому середовищі, де Telegram-бот виконує функцію проміжної ланки між користувачем та системою аналізу. Користувач надсилає текстове повідомлення до бота, після чого воно автоматично обробляється, токенізується, перетворюється у числовий формат та передається у вхід нейронної мережі. Результат класифікації, тобто визначене джерело тексту або клас, надсилається користувачеві у відповідь.

Для перевірки працездатності системи було проведено тестування з використанням реальних текстових прикладів з різних джерел. Було перевірено як повідомлення, що належать до відомих інформаційних джерел, так і тексти з менш достовірних або непередбачених джерел. На рисунку 3.7 наведено приклад взаємодії з Telegram-ботом під час тестування.



Рисунок 3.7 – Приклад результату класифікації тексту за допомогою Telegram-бота

Результати тестування засвідчили високу точність класифікації: для більшості вхідних повідомлень система правильно визначила джерело з високим рівнем впевненості (понад 97%). Тексти, які не входили до навчальної вибірки або значно відрізнялися за стилістикою, були класифіковані менш впевнено або не були віднесені до жодного з класів з пороговою ймовірністю. Це свідчить про те, що модель не перенавчена та здатна узагальнювати інформацію.

Таким чином, проведене тестування довело функціональність та стабільність розробленого модуля в інтерактивному середовищі Telegram. Поєднання нейромережевої моделі з Telegram-ботом дозволяє ефективно реалізувати прикладну задачу визначення джерела тексту в реальному часі та забезпечує зручність використання для кінцевих користувачів.

Підсумовуючи, можна стверджувати, що реалізований модуль на основі Embedding та LSTM [3] не лише продемонстрував високу ефективність

класифікації (точність понад 97%), але й підтвердив свою практичну придатність у реальному середовищі через інтеграцію з Telegram-ботом. Система стабільно обробляє запити, адаптується до різнотипних текстів і забезпечує швидкий зворотний зв'язок для користувача. Це свідчить про успішну реалізацію функціональних вимог та потенціал подальшого впровадження модуля у сфері автоматизованого медіа аналізу, фактчекінгу та цифрової безпеки.

ВИСНОВКИ

У межах кваліфікаційної роботи було реалізовано нейромережевий модуль для автоматичного визначення джерела тексту на основі методів обробки природної мови (Embedding та LSTM). За результатами проведених досліджень та експериментів сформовано такі висновки:

1. Проведено аналіз існуючих підходів до класифікації джерел тексту, зокрема методів глибокого навчання (BERT, FastText, DeepMoji, LSTM). У процесі аналізу з'ясовано, що LSTM - моделі демонструють високу ефективність для задач з довгими текстовими залежностями, забезпечуючи точність на рівні 91–97% згідно з літературними джерелами.

2. Створено архітектуру модуля, яка включає п'ять ключових компонентів: попередню обробку, векторизацію (Embedding), навчання моделі LSTM, оцінювання результатів та прогнозування. Було реалізовано 10-шарову модель, яка включає два LSTM - шари (128 та 64 нейрони відповідно) та один проміжний Dense-шар на 64 нейрони.

3. Розроблено програмну реалізацію модуля у середовищі Python з використанням бібліотек TensorFlow, Keras, spaCy, NLTK. Загальна кількість тренованих параметрів моделі становить 2 171 241, з яких 2 млн — параметри шару Embedding. Модель навчалась на корпусі з 20 000 слів та максимальним розміром послідовності 300 токенів.

4. У ході навчання досягнуто точності 99.9% на тренувальній вибірці, а валідаційна точність зберігалася в діапазоні 97.3–97.8% протягом усіх 10 епох. Найменше значення функції втрат на валідації становило $\text{val_loss} = 0.0747$, що свідчить про відсутність перенавчання та здатність моделі до узагальнення.

5. Тестування в реальному середовищі Telegram-бота підтвердило працездатність і стабільність модуля. Під час тестування модель правильно класифікувала понад 97% реальних повідомлень. Для невідомих або стилістично відмінних текстів модель демонструвала обережність, що також свідчить про її надійність.

6. Таким чином, запропонований підхід довів свою ефективність і практичну значущість. Результати роботи підтверджують, що поєднання Embedding та LSTM [3] є доцільним для розв'язання задачі визначення джерела тексту в умовах сучасного інформаційного середовища. Розроблений модуль може бути адаптований до задач фактчекінгу, авторської атрибуції, автоматизованого контент-аналізу та підвищення інформаційної безпеки.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Felbo B., Mislove A., Søgaard A., Rahwan I., Lehmann S. Using millions of emoji occurrences to learn any-domain representations for detecting sentiment, emotion and sarcasm. arXiv preprint. 2017. arXiv:1708.00524. URL: <https://arxiv.org/abs/1708.00524>
2. Devlin J., Chang M. W., Lee K., Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *NAACL-HLT*. 2019.
3. Hochreiter S., Schmidhuber J. Long Short-Term Memory. *Neural Computation*. 1997. Vol. 9(8). P. 1735–1780.
4. Bojanowski P., Grave E., Joulin A., Mikolov T. Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*. 2017. Vol. 5. P. 135–146.
5. StopFake.org — Центр медіареформ [Електронний ресурс]. URL: <https://www.stopfake.org/> (дата звернення: 11.06.2025).
6. FactCheck.org – Annenberg Public Policy Center [Електронний ресурс]. URL: <https://www.factcheck.org/> (дата звернення: 11.06.2025).
7. TensorFlow [Електронний ресурс]. URL: <https://www.tensorflow.org/> (дата звернення: 11.06.2025).
8. Keras [Електронний ресурс]. URL: <https://keras.io/> (дата звернення: 11.06.2025).
9. spaCy: Industrial-Strength NLP [Електронний ресурс]. URL: <https://spacy.io/> (дата звернення: 11.06.2025).
10. Check-It плагін [Електронний ресурс]. URL: <https://check-it.org/> (дата звернення: 11.06.2025).
11. Shahzeidi M., Ardakani M. M., Javan S. E. A hybrid model of long short-term memory neural networks and quantum behavior PSO for detecting self-admitted technical debt. *Cluster Computing*. 2025. DOI: <https://doi.org/10.1007/s10586-024-04745-4>

12. Joshi M., Deore M. Leveraging LSTM networks for SPAM detection in real-time chat messages. *2025 1st International Conference on AIML, IEEE*. 2025. URL: <https://ieeexplore.ieee.org/document/10932225>
13. Lutz M. Learning Python. 5th ed. Sebastopol : O'Reilly Media, 2013. 1216 p.
14. Ahmed, S., Rakin, S., Waliur, M. W. I., & Islam, N. B. (2024). Depression detection from social media Bangla text using recurrent neural networks. arXiv preprint. <https://arxiv.org/abs/2412.05861>
15. Wani, M. A., Abd El-Latif, A. A., & ELAffendi, M. (2024). AI-based framework for discriminating human-authored and AI-generated text. *IEEE Transactions on Artificial Intelligence*. <https://ieeexplore.ieee.org/document/10798923>
16. Haz L., Rodríguez-García M. Á., Fernández A. Using deep neural networks architectures to identify narcissistic personality traits. *Expert Systems*. 2025. DOI: <https://doi.org/10.1111/exsy.70056>
17. Aishwarya, B., & Durgagowri, G. (2025). Maximizing sentiment detection through multimodal data fusion: Integrating CNN, RNN, LSTM. *2025 International Conference on AI Advancements, IEEE*. <https://ieeexplore.ieee.org/document/10894282>
18. Dai, Z., & He, Y. (2025). A non-functional requirements classification model based on cooperative attention mechanism fused with label embedding. *Computers and Electrical Engineering*. <https://doi.org/10.1016/j.compeleceng.2024.108963>
19. Najib, A., Ozel, S. A., Coban, O., & Khan, R. A. N. (2024). Cyberbullying detection using machine learning and deep learning for Turkish text. *ResearchGate*. <https://www.researchgate.net/publication/382636236>
20. Abd, D. H., Frikha, M., & Alimi, A. M. (2024). Improving fake news detection with adversarial recurrent neural network. *IEEE Conference on Development in AI, IEEE*. <https://ieeexplore.ieee.org/document/10912027>
21. Yadulla, A. R., & Kasula, V. K. (2025). Enhanced cybersecurity entity recognition using DeBERTa, Transformer-CNN hybrids, and BiLSTM-Softmax. *2025 37th Conference of Data Mining, IEEE*. <https://ieeexplore.ieee.org/document/11008057>

22. Akter, S. (2024). Integrated machine learning framework for mitigation of buffer overflow, software supply chain, and adversarial attacks. ProQuest Dissertations. <https://search.proquest.com/openview/8692ec2efccaa3de8e8206ff68f41edc/1>

23. Sukumar, P., Krishnan, B., & Supriya, Y. (2025). Diagnosis of autism spectrum disorder using deep learning with natural language processing. In Challenges in Computational Intelligence (pp. 865–874). Taylor & Francis. <https://doi.org/10.1201/9781003559085-71>

24. Наконечна І.В. Розробка нейромережевого модуля визначення джерела тексту з використанням Embedding і LSTM // Збірник тез доповідей студентської науково-практичної конференції «Інтелектуальні інформаційні технології в прикладних дослідженнях» (ІТAR-2025), (Тернопіль, 27–29 травня 2025 року). – Тернопіль: ЗУНУ, 2025. – С. 194-196.

25. Комар М.П., Саченко А.О., Васильків Н.М., Гладій Г.М., Коваль В.С., Ліп'яніна-Гончаренко Х.В. Методичні рекомендації до виконання кваліфікаційної роботи з освітньо-професійної програми «Комп'ютерні науки» спеціальності 122 «Комп'ютерні науки» за першим (бакалаврським) рівнем вищої освіти. Тернопіль: ЗУНУ, 2024. 52 с.

Додаток А

Лістинг коду модуля класифікації тексту

```
import tensorflow as tf
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import Embedding, LSTM, Dense, Dropout,
SpatialDropout1D
MAX_NB_WORDS = 20000
MAX_SEQUENCE_LENGTH = 300
EMBEDDING_DIM = 100
NUM_CLASSES = len(set(y))
model = Sequential([
    Embedding(input_dim=MAX_NB_WORDS+1, output_dim=EMBEDDING_DIM,
input_length=MAX_SEQUENCE_LENGTH),
    SpatialDropout1D(0.2),
    LSTM(128, return_sequences=True),
    Dropout(0.2),
    LSTM(64),
    Dropout(0.2),
    Dense(64, activation='relu'),
    Dropout(0.2),
    Dense(NUM_CLASSES, activation='softmax') # Класифікація
])
model.compile(loss='sparse_categorical_crossentropy', optimizer='adam',
metrics=['accuracy'])
model.build(input_shape=(None, MAX_SEQUENCE_LENGTH))
model.summary()
```

Лістинг навчання моделі LSTM з використанням гіперпараметрів

```
BATCH_SIZE = 32
EPOCHS = 10
VALIDATION_SPLIT = 0.2
```

```
history = model.fit(  
    X_padded, y,  
    batch_size=BATCH_SIZE,  
    epochs=EPOCHS,  
    validation_split=VALIDATION_SPLIT,  
    verbose=1  
)
```

Додаток Б

Лістинг реалізації Telegram-бота для визначення категорії тексту

```
import logging
import numpy as np
import asyncio
import nest_asyncio # Додаємо для Google Colab
from telegram import Update
from telegram.ext import Application, CommandHandler, MessageHandler, filters,
CallbackContext
from tensorflow.keras.preprocessing.sequence import pad_sequences
from tensorflow.keras.models import load_model
from tensorflow.keras.preprocessing.text import Tokenizer
nest_asyncio.apply()
model = load_model('/content/text_classification_lstm.h5')
tokenizer = Tokenizer(num_words=20000)
logging.basicConfig(format='%(asctime)s - %(name)s - %(levelname)s - %(message)s',
level=logging.INFO)
logger = logging.getLogger(__name__)
def process_text(text):
    sequences = tokenizer.texts_to_sequences([text])
    padded = pad_sequences(sequences, maxlen=300)
    return padded
def predict_category(text):
    processed_text = process_text(text)
    prediction = model.predict(processed_text)
    predicted_class = np.argmax(prediction, axis=1)[0]
    categories = ['Категорія 1', 'Категорія 2', 'Категорія 3', 'Категорія 4', 'Категорія 5']
    return categories[predicted_class]
async def start(update: Update, context: CallbackContext):
```

```

    await update.message.reply_text('Привіт! Надішли текст, і я скажу тобі його
категорію.')
async def handle_message(update: Update, context: CallbackContext):
    user_message = update.message.text
    loop = asyncio.get_event_loop()
    predicted_category = await loop.run_in_executor(None, predict_category,
user_message)
    await update.message.reply_text(f'Цей текст відноситься до категорії:
{predicted_category}')
async def main():
    token = '7966563786:AAFWisl0m8rFxo3zkOGQvifmJIWhv7rCxb4' # Вставте свій
ТОКЕН
    application = Application.builder().token(token).build()
    application.add_handler(CommandHandler("start", start))
    application.add_handler(MessageHandler(filters.TEXT & ~filters.COMMAND,
handle_message))
    await application.run_polling()
await main()

```

Додаток В
Апробація отриманих результатів

Західноукраїнський національний університет
Факультет комп'ютерних інформаційних технологій
Кафедра інформаційно-обчислювальних систем і управління



ЗБІРНИК ТЕЗ ДОПОВІДЕЙ

Студентської науково-практичної конференції
ІНТЕЛЕКТУАЛЬНІ ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ В ПРИКЛАДНИХ
ДОСЛІДЖЕННЯХ
(ІПТАР-2025)

27-29 травня 2025 року

Тернопіль
2025

Карнидал Володимир, Биковий Павло	155
ПРОГРАМНИЙ МОДУЛЬ ВИЗНАЧЕННЯ ВІКУ ТА СТАТІ ЛЮДИНИ НА ЗОБРАЖЕННІ ІЗ ЗАСТОСУВАННЯМ OPENCV ТА ГЛИБОКИХ НЕЙРОННИХ МЕРЕЖ	155
Касьян Тамара, Турченко Ірина	160
ПРОГРАМНИЙ МОДУЛЬ АНАЛІЗУ ТА ПРОГНОЗУВАННЯ ІНФЕКЦІЙНИХ ЗАХВОРЮВАНЬ ЗА ДОПОМОГОЮ МЕТОДІВ МАШИННОГО НАВЧАННЯ	160
Кацидим Віта, Ліп'яніна-Гончаренко Христина, Ралік Ігор	164
ІНТЕЛЕКТУАЛЬНИЙ АНАЛІЗ ЗАДОВОЛЕНOSTІ КЛІЄНТІВ ОБСЛУГОВУВАННЯМ НА ОСНОВІ МЕТОДІВ МАШИННОГО НАВЧАННЯ	164
Климчук Анастасія, Турченко Ірина	168
ІНТЕРАКТИВНИЙ ЗАСТОСУНОК ДЛЯ ВИВЧЕННЯ АНГЛІЙСЬКОЇ МОВИ НА ОСНОВІ МЕДІА-КОНТЕНТУ ІЗ ВИКОРИСТАННЯМ OPENAI	168
Ковальковський Віталій, Ліп'яніна-Гончаренко Христина	171
МОДУЛЬ РОЗПІЗНАВАННЯ ОБЛИЧЧЯ НА ОСНОВІ ГЛИБОКОГО НАВЧАННЯ	171
Ковтуненко Андрій, Ліп'яніна-Гончаренко Христина	176
TELEGRAM-БОТ ДЛЯ ВІДСТЕЖЕННЯ КУРСУ КРИПТОВАЛЮТ	176
Котов Владислав, Майків Ігор	179
МОДУЛЬ КЛАСИФІКАЦІЇ АНОТАЦІЙ НАУКОВИХ СТАТЕЙ З arXiv ЗА ДОПОМОГОЮ МОДЕЛІ RoBERTa	179
Кулик Степан	182
ВИЯВЛЕННЯ МЕРЕЖЕВИХ АТАК У ТРАФІКУ ВЕЛИКИХ ДАНИХ НА ОСНОВІ ЕВОЛЮЦІЙНИХ ОБЧИСЛЕНЬ	182
Мельничук Руслана, Ліп'яніна-Гончаренко Христина	185
МОДУЛЬ ОПТИМІЗАЦІЇ МАРКЕТИНГОВИХ КАМПАНІЙ ЗА ДОПОМОГОЮ КЛАСТЕРИЗАЦІЇ КЛІЄНТІВ	185
Матвійшин Наталія, Майків Ігор	188
МОДУЛЬ АНАЛІЗУ ТА КЛАСИФІКАЦІЇ СЮЖЕТІВ ФІЛЬМІВ ЗА ЖАНРАМИ З ВИКОРИСТАННЯМ БІБЛІОТЕКИ SPASU	188
Мукан Павло, Лендюк Тарас	191
МОДУЛЬ ВИЯВЛЕННЯ МОВИ ВОРОЖНЕЧІ З ВИКОРИСТАННЯМ БІБЛІОТЕКИ SCIKIT-LEARN ДЛЯ КЛАСИФІКАЦІЇ ТЕКСТУ	191
Наконечна Ірина, Ліп'яніна-Гончаренко Христина	194
ВИЗНАЧЕННЯ ДЖЕРЕЛА ТЕКСТУ ЗА ДОПОМОГОЮ НЕЙРОМЕРЕЖІ НА ОСНОВІ EMBEDDING І LSTM	194

Наконечна Ірина
студентка групи КНШ-41
iraaa301003@gmail.com

Ліп'яніна-Гончаренко Христина
д.т.н., доцент
kh.lipianina@wunu.edu.ua

Західноукраїнський національний університет
Тернопіль, Україна

ВИЗНАЧЕННЯ ДЖЕРЕЛА ТЕКСТУ ЗА ДОПОМОГОЮ НЕЙРОМЕРЕЖІ НА ОСНОВІ EMBEDDING І LSTM

Визначення джерела тексту є важливим завданням у сфері обробки природної мови (NLP), що полягає у встановленні, до якого автора, ресурсу або категорії належить певний текстовий фрагмент. Такий підхід актуальний у системах автоматичної модерації контенту, захисту авторських прав, цифрової криміналістики, аналізу соціальних мереж та журналістських розслідувань [1].

З розвитком штучного інтелекту з'явилися нові методи, що забезпечують високу точність класифікації текстів, зокрема на основі глибокого навчання. Одним із таких підходів є використання векторних подань слів (word embeddings) у поєднанні з довготривалою короткочасною пам'яттю (Long Short-Term Memory, LSTM). Embedding дозволяє ефективно перетворювати текстову інформацію у числовий формат, зберігаючи семантичні зв'язки між словами [2]. LSTM-моделі, своєю чергою, здатні аналізувати послідовності тексту, виявляючи закономірності на основі контексту, що робить їх особливо корисними для задач авторського стилю або тематичної класифікації [3].

Система складається з таких компонентів:

1. попередня обробка тексту (очищення, токенизація, лематизація);
2. побудова embedding-матриці на основі Word2Vec або GloVe;
3. навчання LSTM-класифікатора на анотаційованій вибірці;
4. оцінювання точності та узагальнювальної здатності моделі за допомогою метрик F1, Precision, Recall.

На рисунку 1 представлено загальну архітектуру розробленого модуля, який охоплює всі етапи - від обробки текстових даних до класифікації за джерелами.



Рисунок 1 – Архітектура нейромережевого модуля визначення джерела тексту з використанням Embedding та LSTM

Результати експериментального дослідження засвідчили, що запропонований модуль демонструє високу ефективність при визначенні джерел тексту навіть за наявності змішаного лексичного фону, або спотвореного контенту.

У таблиці 1 наведено порівняльну ефективність кількох популярних моделей, які використовуються для визначення джерела тексту.

Таблиця 1 – Порівняння моделей для класифікації тексту за джерелом

Модель	Точність (%)	Особливості
SVM (TF-IDF)	72.1	Лінійна модель, чутлива до шуму
Random Forest	75.4	Нестабільна на малих корпусах
LSTM + Embedding	89.3	Висока здатність до узагальнення
BERT	91.7	Складність і потреба у великих обчислювальних ресурсах

Наукове та практичне значення розробки полягає у створенні гнучкого програмного модуля, який можна адаптувати під різні задачі текстової класифікації: від ідентифікації авторства до фільтрації контенту за тематичними або стилістичними ознаками. Впровадження подібних рішень у практику дозволить підвищити рівень безпеки, автоматизувати аналіз великих текстових масивів та підтримати розвиток технологій інформаційного аналізу в Україні.

Список використаних джерел

1. Stamatatos, E. A survey of modern authorship attribution methods. *Journal of the American Society for Information Science and Technology*. 2009. 60(3), 538–556.
2. Mikolov T., Chen K., Corrado G., Dean J. Efficient Estimation of Word Representations in Vector Space. arXiv preprint arXiv:1301.3781, 2013.
3. Hochreiter S., Schmidhuber J. Long Short-Term Memory. *Neural Computation*. 1997. Vol. 9(8), pp. 1735–1780.
4. Kowsari K. et al. Text Classification Algorithms: A Survey. *Information*, 2019. Vol. 10(4), p. 150.
5. Yang Z., Yang D., Dyer C., et al. Hierarchical Attention Networks for Document Classification. *Proceedings of NAACL*, 2016.
6. Devlin J., Chang M.-W., Lee K., Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. arXiv preprint arXiv:1810.04805, 2018.