

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

Вітрук Іван Петрович

Модуль визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі / Module for sentiment polarity detection in Yelp reviews using an LSTM network

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

Кваліфікаційна робота

Виконав студент групи КН-42
Вітрук І. П.

Науковий керівник:
к.т.н., доцент
Лендюк Т.В.

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 – Комп'ютерні науки

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

ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ

Вітрук Іван Петрович

(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи: Модуль визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі / Module for sentiment polarity detection in Yelp reviews using an LSTM network

керівник роботи к.т.н., доцент, Лендюк Т.В.

(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:

– Провести аналіз предметної області та існуючих методів аналізу тональності текстів.

– Дослідити архітектуру LSTM-мереж та їхню застосовність до задач аналізу відгуків.

– Розробити архітектуру модуля аналізу полярності відгуків.

– Реалізувати алгоритм збору, підготовки та обробки даних для навчання моделі.

– Налаштувати параметри LSTM-мережі та оцінити її продуктивність за допомогою обраних метрик.

– Порівняти результати запропонованого підходу з іншими методами аналізу тональності.

5. Перелік графічного матеріалу в роботі:

– Архітектура модуля аналізу полярності відгуків.

– Алгоритм збору та підготовки даних.

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.	
7	Перевірка кваліфікаційної роботи на оригінальність тексту	до 30.05.2025 р.	
8	Оформлення кваліфікаційної роботи та отримання допуску до захисту	до 10.06.2025 р.	
9	Подання кваліфікаційної роботи до захисту на засіданні атестаційної комісії	до 10.06. 2025 р.	

Студент _____ Вітрук І. П.
(підпис) (прізвище та ініціали)

Керівник кваліфікаційної роботи _____ Лендюк Т.В
(підпис) (прізвище та ініціали)

АНОТАЦІЯ

Кваліфікаційна робота на тему «Модуль визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі» на здобуття освітнього ступеня «бакалавр» зі спеціальності 122 «Комп'ютерні науки» виконана на 69 сторінках, містить 12 рисунків, 2 таблиці, 3 додатки та 24 використані джерела.

Метою роботи є розробка програмного модуля для автоматичного визначення позитивної чи негативної полярності текстових відгуків Yelp на основі архітектури LSTM.

У дослідженні використано методи попередньої обробки текстів, word embedding, рекурентні нейронні мережі типу LSTM, а також метрики Accuracy, Precision, Recall, F1 та візуалізацію результатів через confusion matrix.

Створений модуль продемонстрував точність 90,28 %, F1-міру 0,91 на тестовій вибірці 38 000 відгуків; при додатковій перевірці на 30 реальних рецензіях точність склала 93,33 %. Отримані результати підтверджують ефективність LSTM-підходу для аналізу споживчих настроїв.

Модуль може бути інтегрований у системи моніторингу клієнтського досвіду, автоматизовані CRM-платформи та бізнес-аналітику задля оперативного реагування на зміну настроїв споживачів.

Ключові слова: АНАЛІЗ ТОНАЛЬНОСТІ, LSTM, NLP, ВІДГУКИ YELP, DEEP LEARNING.

ANNOTATION

The bachelor's thesis titled "Module for Sentiment Polarity Detection in Yelp Reviews Using an LSTM Network" was prepared for the Bachelor's degree in Computer Science. The manuscript comprises 69 pages, 12 figures, 2 tables, 3 appendices, and 24 references.

The aim of the work is to develop a software module that automatically classifies Yelp reviews as positive or negative using a Long Short-Term Memory (LSTM) neural-network architecture.

The study employs text-pre-processing, word embeddings, LSTM networks, and evaluation metrics such as accuracy, precision, recall, F1-score, alongside confusion-matrix visualisation.

The implemented module achieved 90.28 % accuracy and an F1-score of 0.91 on a 38 000-review test set; an additional real-world sample of 30 reviews yielded 93.33 % accuracy. These results confirm the effectiveness of the LSTM approach for consumer-sentiment analysis.

The module can be integrated into customer-experience monitoring systems, automated CRM platforms, and business-intelligence pipelines to enable timely responses to shifts in consumer sentiment.

Keywords: SENTIMENT ANALYSIS, LSTM, NLP, YELP REVIEWS, DEEP LEARNING.

ЗМІСТ

Вступ	7
1 Аналіз предметної області та постановка задачі.....	9
1.1 Огляд платформи Yelp та її значення для аналізу відгуків.....	9
1.2 Методи аналізу тональності текстів: огляд існуючих рішень	11
1.3 Глибокі нейронні мережі для аналізу тексту.....	13
1.4 Основи LSTM-мереж та їх застосування у текстовій аналітиці	15
1.5 Постановка задачі	17
2 Проектування модуля аналізу полярності відгуків.....	20
2.1 Архітектура модуля	20
2.2 Алгоритм процесу збору та підготовки даних	23
2.3 Вибір параметрів моделі LSTM	25
2.4 Метрики оцінки ефективності моделі	28
2.5 Алгоритм визначення популярності відгуків за допомогою LSTM-мережі	30
3 Реалізація модуля полярності відгуків	33
3.1 Встановлення середовища розробки та підключення бібліотек	33
3.2 Завантаження та попередня обробка даних	34
3.4 Побудова та навчання LSTM-моделі	38
3.5 Оцінка якості класифікації	40
Висновки	45
Список використаних джерел	47
Додаток А Програмний код.....	50
Додаток Б Приклади реалізації.....	60
Додаток В Апробація отриманих результатів	65

ВСТУП

У сучасному цифровому світі онлайн-відгуки відіграють ключову роль у прийнятті рішень споживачами та формуванні бізнес-стратегій. Платформи, такі як Yelp, містять величезні обсяги текстових рецензій, які є цінним джерелом інформації про якість товарів і послуг. Аналіз цих відгуків дає можливість оцінити рівень задоволеності клієнтів, виявити основні проблеми та визначити тенденції у сфері споживчих уподобань.

Однак традиційні методи аналізу тональності тексту мають певні обмеження, пов'язані з розпізнаванням складних мовних конструкцій, таких як сарказм, багатозначність слів і контекстуальні зміни значень. Методи машинного навчання дозволяють частково вирішити ці проблеми, але їхня ефективність залежить від якості вибору ознак та обсягу тренувальних даних. Використання глибоких нейронних мереж, зокрема LSTM (Long Short-Term Memory), дозволяє зберігати довготривалі залежності між словами, що значно покращує результати аналізу.

Автоматизація аналізу відгуків є актуальною задачею для підприємств, які прагнуть оперативно реагувати на зміну споживчих настроїв. Запропонований у цій роботі модуль аналізу полярності відгуків на основі LSTM-мережі дозволить ефективно класифікувати відгуки як позитивні чи негативні, що сприятиме глибшому розумінню клієнтських уподобань і покращенню якості обслуговування.

Таким чином, дослідження спрямоване на розробку ефективного підходу до аналізу тональності текстів, що поєднує передові методи обробки природної мови з можливостями глибокого навчання. Це забезпечить високоточну класифікацію відгуків та можливість масштабованого застосування для аналізу великих обсягів даних.

Мета роботи – розробка та оцінка ефективності модуля визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі.

Основні завдання:

1. Провести аналіз предметної області та існуючих методів аналізу тональності текстів.

2. Дослідити архітектуру LSTM-мереж та їхню застосовність до задач аналізу відгуків.

3. Розробити архітектуру модуля аналізу полярності відгуків.

4. Реалізувати алгоритм збору, підготовки та обробки даних для навчання моделі.

5. Налаштувати параметри LSTM-мережі та оцінити її продуктивність за допомогою обраних метрик.

6. Порівняти результати запропонованого підходу з іншими методами аналізу тональності.

Предметом дослідження є методи автоматизованого аналізу полярності текстових відгуків за допомогою нейронних мереж.

Об'єктом дослідження є процес обробки текстових даних, їхньої підготовки та класифікації за допомогою LSTM-мережі.

Методи дослідження включають аналіз існуючих підходів до аналізу тональності текстів, застосування методів машинного навчання та глибоких нейронних мереж, експериментальне дослідження ефективності моделі, а також оцінку її продуктивності на основі стандартних метрик класифікації. Реалізація та тестування моделі проводяться з використанням програмних засобів Python та бібліотек TensorFlow і Keras.

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ПОСТАНОВКА ЗАДАЧІ

1.1 Огляд платформи Yelp та її значення для аналізу відгуків

Платформа Yelp є однією з найпопулярніших у світі систем для збору та аналізу відгуків про підприємства різних категорій, таких як ресторани, готелі, салони краси, магазини та інші сервіси. Вона була заснована у 2004 році і з того часу набула величезної популярності, зібравши мільйони відгуків від користувачів з усього світу.

Однією з основних функцій Yelp є можливість залишати текстові відгуки та оцінювати підприємства за п'ятибальною шкалою. Ці оцінки є важливими для потенційних клієнтів, які шукають найкращі послуги та товари. Крім того, платформа дозволяє користувачам додавати фотографії, коментувати відгуки інших користувачів та взаємодіяти з підприємствами безпосередньо через месенджер.

Дані, що містяться на платформі Yelp, є цінним джерелом інформації для аналізу споживчих настроїв та оцінки якості обслуговування. Аналіз тональності відгуків дозволяє визначити основні тенденції у відгуках клієнтів, покращити якість послуг підприємств та формувати маркетингові стратегії на основі реальних відгуків.

Оскільки Yelp містить величезний обсяг текстових даних, їх обробка вимагає застосування методів обчислювальної лінгвістики, машинного навчання та штучного інтелекту. Сучасні технології дозволяють автоматизувати процес аналізу відгуків та виділяти корисну інформацію для прийняття бізнес-рішень.

Дані, що містяться на платформі, представлені у вигляді відкритого набору Yelp Dataset, який включає:

- Текстові відгуки користувачів із зазначенням рейтингу.
- Інформацію про підприємства (категорія, місцезнаходження, рейтинг, кількість відгуків тощо).
- Дані про користувачів, які залишають відгуки (ідентифікатор, дата реєстрації, кількість написаних відгуків).

– Взаємодію між користувачами, таку як вподобання або коментарі до відгуків.

Дані на платформі Yelp представлені у вигляді структурованої інформації, що містить різні атрибути для подальшого аналізу (Таблиця 1.1).

Таблиця 1.1 - Структура даних на платформі Yelp

Категорія даних	Опис
Відгуки	Текст відгуку, рейтинг (1-5), дата публікації, мітки корисності
Підприємства	Назва, категорія, адреса, загальний рейтинг, кількість відгуків
Користувачі	Унікальний ідентифікатор, дата реєстрації, загальна активність
Взаємодії	Лайки, коментарі, відповіді підприємств, спільноти

Ця інформація є основою для побудови моделей аналізу тональності тексту, зокрема за допомогою нейронних мереж, таких як LSTM. Дані можуть використовуватися для тренування моделей, які автоматично класифікують відгуки на позитивні та негативні, а також прогнозують рівень задоволеності клієнтів.

Ця інформація є основою для побудови моделей аналізу тональності тексту, зокрема за допомогою нейронних мереж, таких як LSTM. Дані можуть використовуватися для тренування моделей, які автоматично класифікують відгуки на позитивні та негативні, а також прогнозують рівень задоволеності клієнтів.

Завдяки аналізу відгуків на платформі Yelp підприємства можуть:

- Виявляти основні проблеми в обслуговуванні.
- Отримувати реальні відгуки про якість товарів і послуг.
- Оптимізувати маркетингові стратегії на основі клієнтських оцінок.
- Автоматично аналізувати настрої клієнтів за допомогою алгоритмів машинного навчання.

Отже, Yelp не лише виконує функцію бази відгуків, а й слугує важливим інструментом для бізнес-аналізу, прогнозування та покращення якості

обслуговування. У наступних підрозділах буде розглянуто основні методи аналізу тональності тексту, які застосовуються для автоматичної класифікації відгуків на цій платформі.

1.2 Методи аналізу тональності текстів: огляд існуючих рішень

Аналіз тональності тексту є однією з ключових задач у сфері обробки природної мови (NLP). Його мета полягає у визначенні емоційного забарвлення тексту, що може бути позитивним, негативним або нейтральним. Сучасні методи аналізу тональності можна умовно поділити на три великі групи: лексиконні методи, методи на основі правил та методи машинного навчання.

Лексиконні методи базуються на використанні попередньо створених словників тональних слів. У таких підходах кожне слово оцінюється відповідно до його позитивного або негативного забарвлення. Наприклад, якщо текст містить більше позитивних слів, його тональність визначається як позитивна. Однією з популярних систем цього класу є SentiWordNet, яка містить оцінки тональності для широкого набору слів. Основна перевага таких методів полягає у простоті реалізації, проте вони мають обмеження, оскільки не враховують контекст і не здатні розпізнавати складні мовні конструкції, такі як сарказм чи багатозначність слів.

Методи на основі правил покладаються на створення спеціальних логічних алгоритмів, які визначають тональність тексту на основі синтаксичних та морфологічних особливостей мови. Такі системи можуть використовувати аналіз граматичних структур, а також спеціальні правила, що змінюють тональність залежно від контексту. Наприклад, слово «добре» зазвичай має позитивне значення, проте якщо воно використовується в конструкції «не добре», тональність змінюється. Висока точність у граматично коректних текстах є однією з переваг цього підходу, проте розробка правил є трудомісткою, і такі методи слабо пристосовані до змін у мовних структурах.

Методи машинного навчання стали найбільш популярними для аналізу тональності тексту завдяки своїй гнучкості та здатності узагальнювати інформацію з великих обсягів даних. До традиційних методів цього класу належать наївний байєсівський класифікатор, метод опорних векторів (SVM) і деревоподібні моделі. Вони навчаються на великих наборах анотованих текстів і використовують статистичні підходи для прогнозування тональності. Хоча такі методи показують високу продуктивність, вони залежать від вибору вхідних ознак, що може впливати на якість їхньої роботи.

Глибокі нейронні мережі, зокрема рекурентні нейронні мережі (RNN) і моделі на основі довготривалої короткочасної пам'яті (LSTM), є найсучаснішими підходами до аналізу тональності тексту. Вони працюють із текстовими послідовностями, зберігаючи контекст попередніх слів для точнішого визначення настрою. LSTM-мережі, на відміну від звичайних RNN, здатні ефективно працювати з довготривалими залежностями, що робить їх ідеальними для обробки тексту.

Ще одним сучасним методом є трансформерні моделі, такі як BERT і GPT, які використовують механізм самоуваги для розуміння значень слів у контексті. Вони забезпечують найвищу точність у завданнях аналізу тексту, включаючи визначення тональності, але потребують значних обчислювальних ресурсів.

На рисунку 1.1 представлено основні підходи до аналізу тональності текстів.



Рисунок 1.1 - Основні методи аналізу тональності текстів

Порівняльний аналіз методів демонструє, що вибір конкретного підходу залежить від поставленої задачі та доступних ресурсів. Лексиконні методи є простими, проте їх точність є нижчою, оскільки вони не враховують контекст. Методи на основі правил можуть досягати високої точності у вузькоспеціалізованих задачах, проте їх складно адаптувати до загального використання. Методи машинного навчання, особливо глибокі нейронні мережі, забезпечують найвищу якість аналізу тональності, але потребують великих обсягів даних для навчання.

Таким чином, у сучасному аналізі тональності тексту найбільш ефективними є методи, що використовують глибоке навчання, зокрема моделі LSTM та трансформери. У наступному підрозділі буде розглянуто специфіку застосування глибоких нейронних мереж для аналізу тональності відгуків.

1.3 Глибокі нейронні мережі для аналізу тексту

Глибокі нейронні мережі є одним із найбільш ефективних методів аналізу тексту, включаючи визначення його тональності, категоризацію та узагальнення. Вони засновані на використанні багаторівневої архітектури штучних нейронів, що дає змогу моделі виявляти складні патерни в даних. Основною перевагою глибоких мереж у порівнянні з традиційними методами є їх здатність працювати без необхідності ручного створення правил або словників, що робить їх особливо корисними для обробки природної мови.

Одним із ключових підходів у глибокому навчанні є рекурентні нейронні мережі (RNN), які використовуються для обробки послідовних даних, таких як текстові рядки. RNN здатні зберігати контекст попередніх елементів послідовності, що робить їх придатними для аналізу зв'язності слів у реченнях. Однак класичні RNN мають проблему затухання градієнтів, що ускладнює навчання на довгих послідовностях.

Для вирішення цієї проблеми була запропонована архітектура довготривалої короткочасної пам'яті (LSTM). На відміну від стандартних RNN,

LSTM-мережі використовують спеціальні механізми збереження інформації, що дозволяє їм ефективно працювати з довготривалими залежностями в текстах. Це робить їх ідеальними для завдань аналізу тональності відгуків, машинного перекладу та автоматичного реферування текстів.

Ще одним важливим досягненням у сфері глибокого навчання є трансформерні моделі, зокрема BERT (Bidirectional Encoder Representations from Transformers) і GPT (Generative Pre-trained Transformer). На відміну від LSTM, вони використовують механізм самоуваги, що дозволяє їм обробляти всі слова в тексті одночасно, враховуючи їх контекст у всій послідовності. Завдяки цьому трансформери показують найкращі результати в задачах обробки природної мови та встановлюють нові стандарти в аналізі тексту.

На рисунку 1.2 представлено загальну архітектуру глибоких нейронних мереж для аналізу тексту.

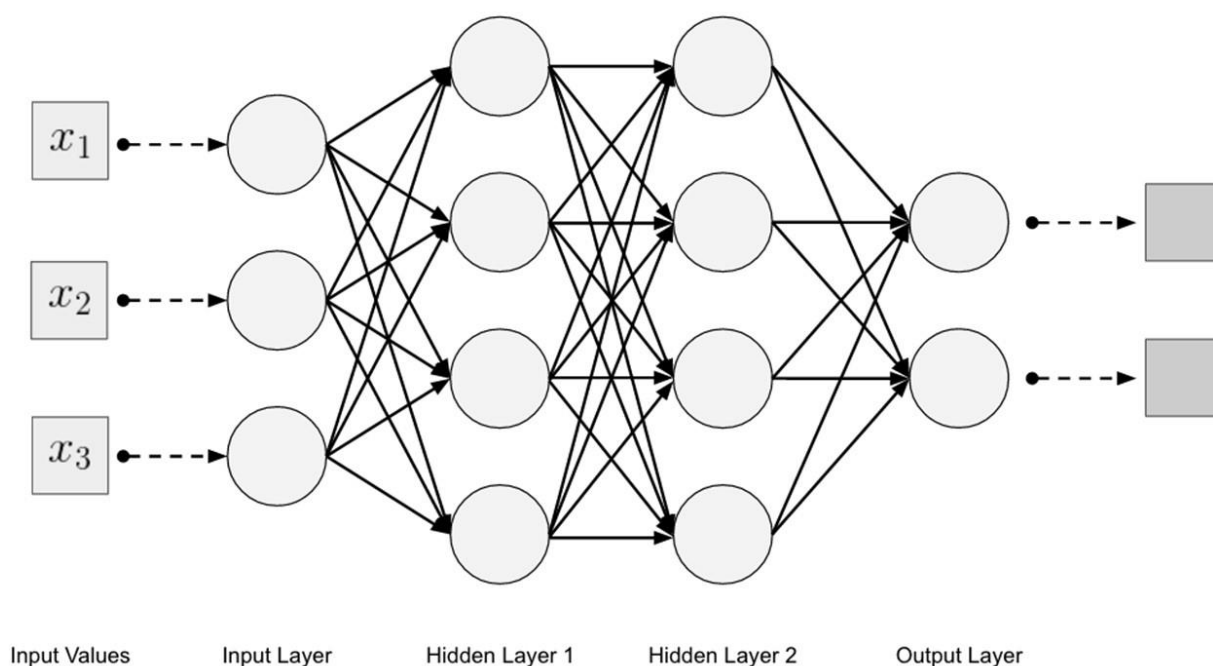


Рисунок 1.2 - Архітектура глибоких нейронних мереж для аналізу тексту.

Глибокі нейронні мережі використовують багаторівневу обробку інформації, що дозволяє їм аналізувати не лише окремі слова, а й розуміти загальний контекст речень і навіть абзаців. Це особливо важливо в завданнях

аналізу відгуків, де значення речення може змінюватися в залежності від попередніх і наступних фраз. У випадку LSTM, кожен осередок мережі містить вхідний, вихідний і забуттєвий гейти, що регулюють потік інформації через мережу.

Застосування глибоких мереж у аналізі тексту також передбачає використання попередньо навчених моделей, які були натреновані на великих корпусах тексту. Це дозволяє значно зменшити потребу у великих наборах навчальних даних та скоротити час розгортання нових систем. Моделі, такі як BERT, можуть бути адаптовані для конкретних задач шляхом додаткового донавчання на спеціалізованих наборах даних.

Глибокі нейронні мережі значно покращують автоматичний аналіз тексту, дозволяючи не лише класифікувати відгуки як позитивні чи негативні, а й розпізнавати складні мовні конструкції, такі як сарказм або контекстуальні зміни значень слів. Використання механізму самоуваги в трансформерних моделях робить їх надзвичайно ефективними, адже вони можуть аналізувати текстові дані у двох напрямках і визначати важливість кожного слова у всьому контексті.

Таким чином, глибокі нейронні мережі, зокрема LSTM та трансформерні архітектури, є найбільш ефективними підходами до аналізу тексту. Вони забезпечують високу точність, гнучкість у використанні та здатність адаптуватися до нових даних. У наступному підрозділі буде розглянуто специфіку їхнього застосування в контексті аналізу полярності відгуків.

1.4 Основи LSTM-мереж та їх застосування у текстовій аналітиці

LSTM (Long Short-Term Memory) – це тип рекурентних нейронних мереж (RNN), який спеціально розроблений для роботи з послідовними даними, такими як текст. Основна особливість LSTM-мереж – здатність зберігати довготривалі залежності завдяки унікальній структурі комірок пам'яті. Це робить їх особливо ефективними в задачах обробки природної мови, включаючи аналіз тональності текстів.

Класичні RNN страждають від проблеми затухання градієнтів, що обмежує їхню здатність запам'ятовувати далекі залежності в послідовності. LSTM вирішує цю проблему за допомогою спеціальної архітектури, що включає три основні компоненти: вхідний, вихідний і забуттєвий гейт. Забуттєвий гейт визначає, яку частину попередньої інформації слід зберегти, а яку – відкинути. Вхідний гейт регулює оновлення стану пам'яті, а вихідний – контролює, яка інформація передається далі.

LSTM-мережі широко використовуються для аналізу тексту, оскільки вони можуть ефективно працювати з контекстом і зберігати взаємозв'язки між словами в довгих реченнях. Це особливо важливо для аналізу відгуків, де тональність може змінюватися в залежності від розташування ключових слів і їхнього взаємозв'язку.

На рисунку 1.3 представлена загальна структура LSTM-комірки, що демонструє її основні компоненти та механізм обробки інформації.

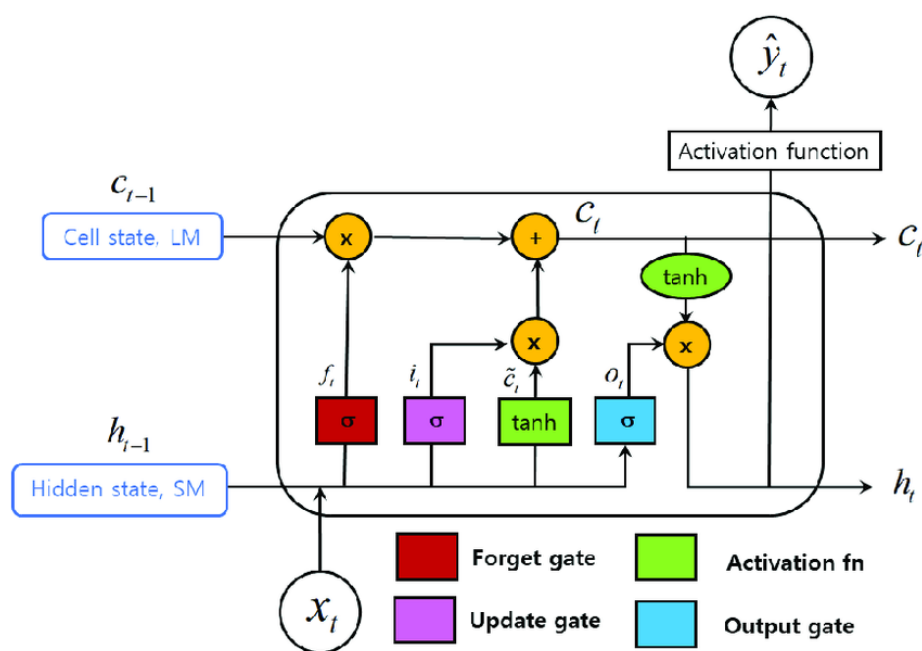


Рисунок 1.3 - Структура LSTM-комірки.

У текстовій аналітиці LSTM використовується для розв'язання низки задач, зокрема класифікації тексту, машинного перекладу, генерації тексту та розпізнавання мовних конструкцій. Для аналізу тональності текстів LSTM-

мережа навчається на великих наборах анотованих текстів, що дозволяє їй виявляти позитивні, негативні та нейтральні настрої.

Основний процес роботи LSTM у текстовій аналітиці складається з кількох етапів. Спочатку текст проходить через процес передобробки, включаючи токенизацію, нормалізацію та кодування слів у числові вектори. Потім ці вектори подаються у LSTM-мережу, яка аналізує їхню послідовність, зберігаючи контекст і залежності між словами. Після цього отримані ознаки передаються у вихідний шар, який виконує класифікацію або інше передбачення.

LSTM-мережі мають кілька переваг перед іншими підходами до аналізу тексту. Вони можуть обробляти довгі текстові послідовності, враховувати контекст і працювати з мовними структурами, які важко формалізувати традиційними методами. У той же час, їхнє навчання є обчислювально складним і потребує значних ресурсів, особливо при роботі з великими корпусами тексту.

Сучасні дослідження показують, що комбінація LSTM із іншими технологіями, такими як самоувага або трансформерні архітектури, може значно покращити результати аналізу тексту. Наприклад, гібридні моделі, що поєднують LSTM із механізмом уваги, здатні виділяти найважливіші елементи тексту для аналізу.

Таким чином, LSTM-мережі є одним із найбільш ефективних інструментів для обробки тексту, що забезпечують високу точність аналізу завдяки можливості роботи з довготривалими залежностями. У наступному підрозділі буде розглянуто специфіку їх застосування у задачі аналізу полярності відгуків.

1.5 Постановка задачі

У сучасному цифровому світі онлайн-відгуки є важливим джерелом інформації про якість товарів і послуг. Платформи, такі як Yelp, містять мільйони текстових рецензій, які впливають на прийняття рішень споживачами та стратегії бізнесів. Аналіз полярності відгуків дозволяє автоматично визначати їхній

емоційний зміст, що має велике значення для автоматизованого моніторингу споживчих настроїв.

Завдяки розвитку методів обробки природної мови (NLP) та глибокого навчання з'явилася можливість застосовувати нейронні мережі для автоматичного визначення тональності тексту. Однак традиційні підходи мають обмеження, пов'язані з розумінням контексту та обробкою довгих текстових послідовностей. Використання LSTM-мереж дозволяє подолати ці проблеми, зберігаючи важливі залежності між словами та покращуючи точність аналізу.

Застосування сучасних методів аналізу відгуків може допомогти компаніям ефективніше обробляти великі масиви даних, оцінювати рівень задоволеності клієнтів та оперативно реагувати на негативні коментарі. У цьому контексті актуальним є створення ефективного модуля аналізу полярності відгуків на основі LSTM-мережі, що забезпечить високу точність класифікації текстових даних.

LSTM-мережі є одними з найефективніших методів для роботи з текстовими послідовностями, оскільки вони дозволяють зберігати важливі семантичні зв'язки між словами навіть у довгих реченнях. На відміну від класичних моделей машинного навчання, таких як SVM або Naïve Bayes, вони не потребують ручного створення ознак і можуть самостійно виявляти важливі патерни в тексті. Це забезпечує високу точність аналізу та адаптацію до різних мовних конструкцій.

Мета роботи – розробка та оцінка ефективності модуля визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі.

Основні завдання:

- Провести аналіз предметної області та існуючих методів аналізу тональності текстів.
- Дослідити архітектуру LSTM-мереж та їхню застосовність до задач аналізу відгуків.
- Розробити архітектуру модуля аналізу полярності відгуків.

- Реалізувати алгоритм збору, підготовки та обробки даних для навчання моделі.
- Налаштувати параметри LSTM-мережі та оцінити її продуктивність за допомогою обраних метрик.
- Порівняти результати запропонованого підходу з іншими методами аналізу тональності.

2 ПРОЕКТУВАННЯ МОДУЛЯ АНАЛІЗУ ПОЛЯРНОСТІ ВІДГУКІВ

2.1 Архітектура модуля

Архітектура модуля аналізу полярності відгуків на платформі Yelp побудована з урахуванням особливостей роботи з текстовими даними та використання глибоких нейронних мереж. Основною метою проектування є створення ефективної та масштабованої системи, здатної обробляти великі обсяги текстової інформації та класифікувати відгуки за їх полярністю.

Архітектура модуля включає наступні основні компоненти:

1. Компонент збору та підготовки даних:

- Отримання відгуків із заданого джерела (файл CSV, база даних тощо).
- Видалення HTML-тегів, спеціальних символів та очищення тексту.
- Токенізація – розбиття тексту на окремі слова або фрази.
- Векторизація – перетворення текстових даних у числове представлення.
- Нормалізація та обмеження довжини послідовностей.
- Формування наборів даних для тренування, валідації та тестування.

2. Компонент побудови моделі:

- Використання попередньо підготовлених векторних представлень слів (word embeddings) або навчання власних векторних представлень.
- Вбудований шар (Embedding Layer) для перетворення токенозованого тексту у векторне представлення.
- Рекурентний шар (LSTM), який зберігає контекст у текстових послідовностях.
- Додаткові шари регуляризації, такі як Dropout, для запобігання перенавчанню.
- Щільний вихідний шар (Dense Layer), що класифікує відгук як позитивний або негативний.

- Визначення оптимальної архітектури шляхом експериментів з кількістю шарів та нейронів.

3. Компонент навчання моделі:

- Використання функції втрат для оцінки точності передбачень (наприклад, бінарна кросентропія).
- Оптимізація вагових коефіцієнтів за допомогою алгоритму Adam.
- Використання навчальної та валідаційної вибірок для коригування параметрів моделі.
- Вибір кількості епох та розміру міні-батчів для навчання.
- Використання технік ранньої зупинки (early stopping) для запобігання перенавчанню.

4. Компонент оцінки якості моделі:

- Використання тестової вибірки для оцінки точності моделі.
- Обчислення основних метрик: точність (Accuracy), F1-міра, матриця помилок.
- Візуалізація результатів роботи моделі у вигляді графіків.
- Аналіз помилкових передбачень для виявлення можливих покращень.
- Порівняння результатів з іншими моделями (наприклад, логістичною регресією, SVM).

5. Компонент інтеграції та розгортання:

- Збереження навченої моделі у форматі, придатному для подальшого використання.
- Розгортання моделі у вигляді веб-сервісу або інтеграція у програмні додатки.
- Надання API для взаємодії з іншими системами.
- Побудова інтерфейсу користувача для взаємодії з системою.
- Тестування продуктивності системи на реальних запитах.

Архітектуру модуля подана на рисунку 2.1.

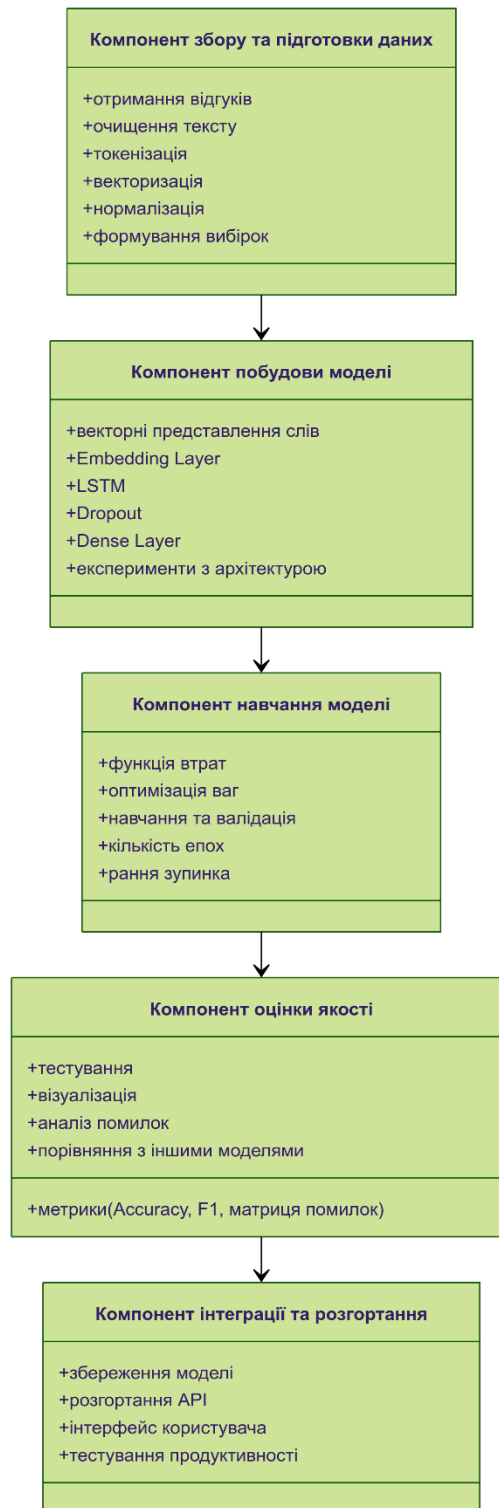


Рисунок 2.1 -Архітектура модуля аналізу полярності відгуків.

Запропонована архітектура забезпечує ефективну роботу з великими обсягами текстової інформації, адаптивність до змін у структурі даних та

можливість масштабування для обробки нових відгуків. У наступних підрозділах буде детально розглянуто кожен з етапів проектування та реалізації модуля.

2.2 Алгоритм процесу збору та підготовки даних

Процес збору та підготовки даних є ключовим етапом у побудові модуля аналізу полярності відгуків. Від якості цього етапу залежить точність роботи моделі, оскільки правильно підготовлені дані допомагають нейромережі ефективно навчатися.

На рисунку 2.2 представлено загальний алгоритм процесу збору та підготовки даних.

Процес збору даних включає отримання відгуків із заданого джерела, їх очищення та підготовку для подальшого аналізу. Вхідні текстові дані, отримані з Yelp, можуть містити HTML-розмітку, спеціальні символи, повторювані елементи або шумові дані. Тому на першому етапі проводиться попередня обробка, яка включає:

- Видалення HTML-тегів та непотрібних символів.
- Конвертацію тексту в нижній регістр для усунення відмінностей між однаковими словами.
- Видалення стоп-слів (наприклад, "the", "and", "in"), що не несуть значного смислового навантаження.
- Нормалізацію пробілів та пунктуації для уніфікації тексту.
- Формально, текстовий відгук після попередньої обробки перетворюється у множину токенів:

$$X = \{w_1, w_2, \dots, w_n\}, w_i \in V \quad (2.1)$$

де V – словник унікальних слів, що використовуються в моделі.

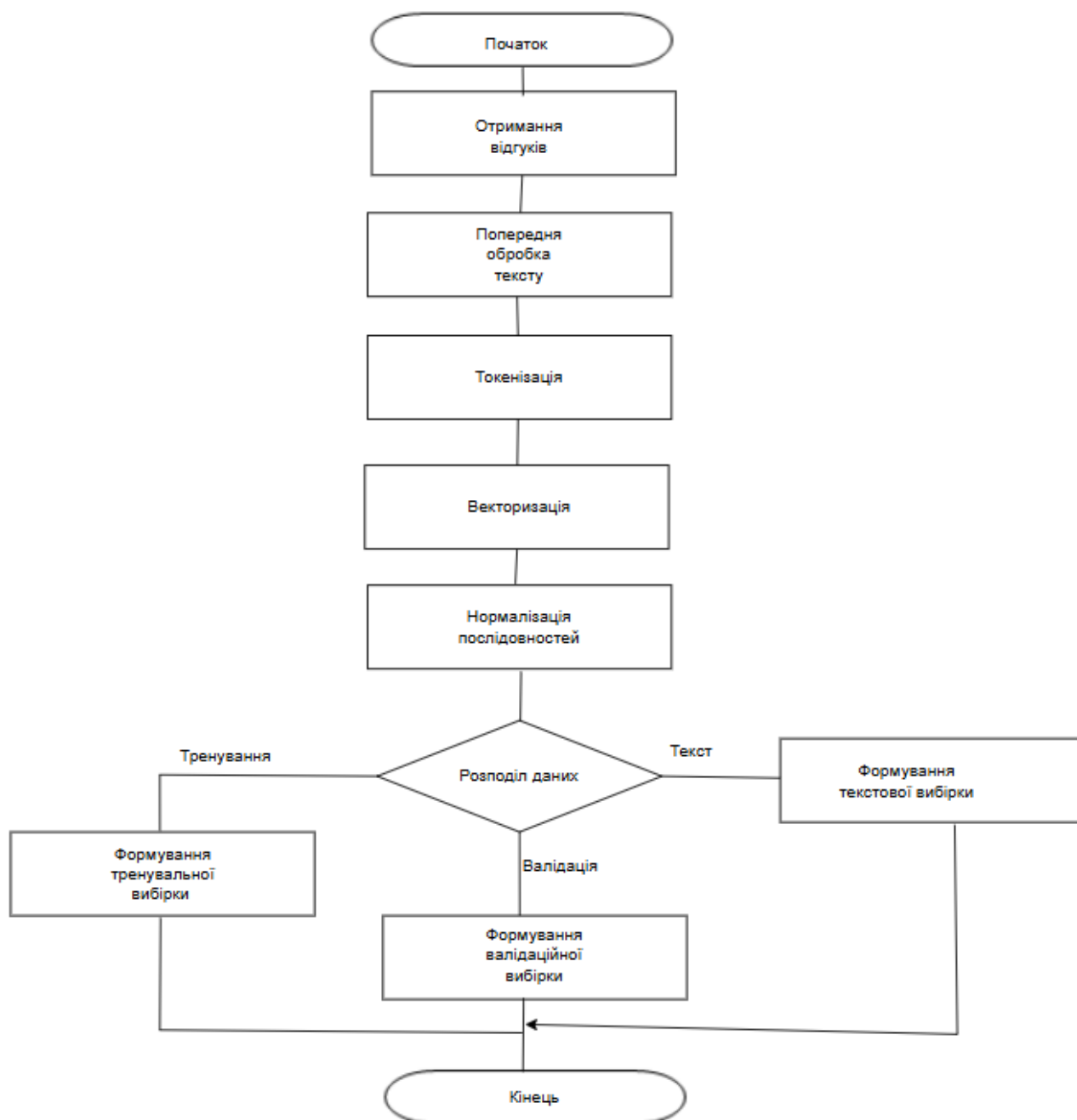


Рисунок 2.2 - Алгоритм збору та підготовки даних.

Наступним етапом є токенизація, яка передбачає розбиття тексту на окремі слова або фрази. Це дозволяє моделі працювати з текстами як із послідовностями числових векторів. Потім проводиться векторизація – кожне слово зі словника відображається у відповідний числовий вектор за допомогою техніки

вбудовування слів (word embedding). Для цього використовується матриця E , яка відображає кожне слово у вектор простору розмірності d : $E: w_i \rightarrow \mathbb{R}^d$

Для роботи з рекурентними неймережами потрібно вирівняти довжини текстів до єдиного розміру. Це досягається шляхом усічення або доповнення послідовностей до фіксованої довжини $L: X_{norm} = \{w_1, \dots, w_L\}$ де коротші послідовності доповнюються спеціальними заповнювачами (padding).

Остаточно, підготовлені дані розбиваються на навчальний, валідаційний та тестовий набори. Це забезпечує коректну оцінку якості моделі під час її навчання та перевірки. Розподіл даних здійснюється таким чином:

- Тренувальна вибірка – 80% даних, використовується для навчання моделі.
- Валідаційна вибірка – 10%, використовується для налаштування параметрів моделі.
- Тестова вибірка – 10%, застосовується для оцінки фінальної продуктивності моделі.

Запропонований алгоритм забезпечує ефективну обробку текстових даних для подальшого використання в LSTM-моделі. У наступному підрозділі буде розглянуто вибір параметрів моделі та методи оптимізації навчального процесу.

2.3 Вибір параметрів моделі LSTM

Процес вибору параметрів для моделі LSTM є критично важливим для забезпечення її ефективності та продуктивності. Визначення оптимальних гіперпараметрів впливає на здатність мережі правильно класифікувати відгуки та узагальнювати нові дані. Вибір параметрів здійснюється шляхом експериментального аналізу та теоретичних міркувань, базуючись на властивостях рекурентних нейронних мереж.

На рисунку 2.3 представлено загальну архітектуру вибору параметрів LSTM-моделі.

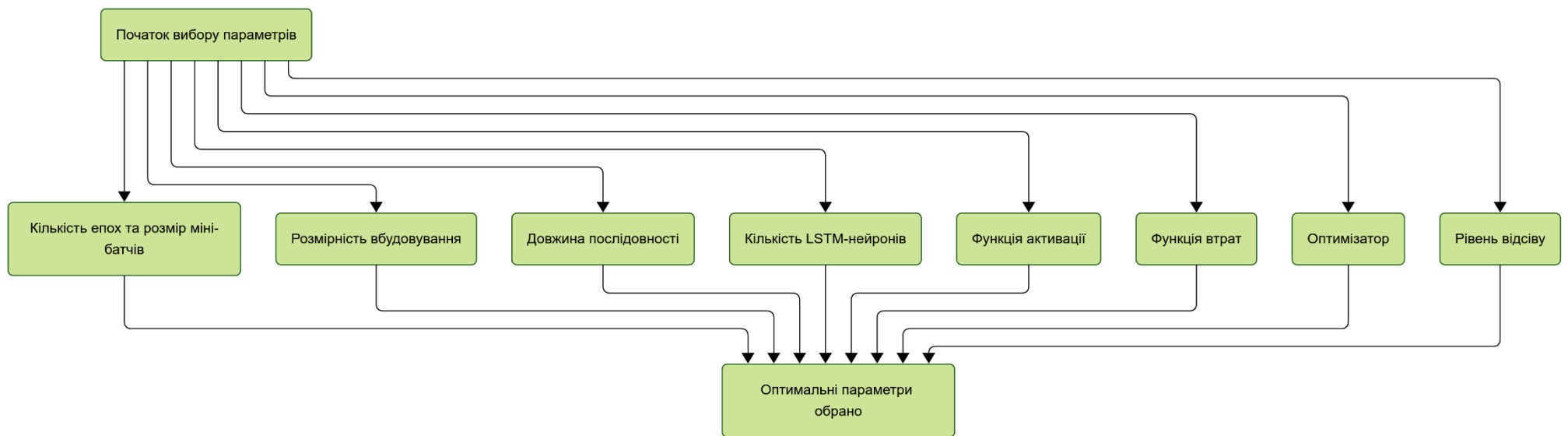


Рисунок 2.3 - Вибір параметрів для моделі LSTM.

Основні параметри, які визначають роботу моделі LSTM, включають:

1. Розмірність вбудовування (Embedding Dimension). Цей параметр визначає розмір векторного представлення слів. Більші значення дозволяють кодувати більше контекстної інформації, однак можуть призводити до перенавчання. Типові значення: 50, 100, 300.

2. Довжина послідовності (Sequence Length). Оскільки модель працює з текстовими послідовностями, необхідно встановити фіксовану максимальну довжину. Занадто короткі значення можуть втратити важливий контекст, тоді як надто довгі збільшують витрати на обчислення. Оптимальне значення вибирається експериментально, наприклад, 200–800 токенів.

3. Кількість LSTM-нейронів у прихованому шарі. Визначає, скільки інформації модель може запам'ятовувати та аналізувати. Більша кількість нейронів покращує точність, але збільшує ризик перенавчання. Типові значення: 64, 128, 256.

4. Функція активації. В LSTM-мережах стандартно використовується сигмоїдна функція $\sigma(x) = \frac{1}{1+e^{-x}}$ для вихідного шару. В прихованих шарах може застосовуватися $\tanh$ для покращення стабільності градієнтів.

5. Функція втрат. Оскільки модель розв'язує задачу бінарної класифікації, використовується бінарна кросентропія:

$$L = -\frac{1}{N} \sum_{i=1}^N y_i \log \hat{y}_i + (1 - y_i) \log(1 - \hat{y}_i) \quad (2.2)$$

де y_i – справжня мітка класу,

$\hat{y}_i$ – передбачене значення.

6. Оптимізатор. Для навчання моделі використовується алгоритм Adam, який забезпечує адаптивне оновлення вагових коефіцієнтів.

7. Рівень відсіву (Dropout). Використовується для запобігання перенавчанню шляхом випадкового вимкнення певного відсотка нейронів під час навчання. Типові значення: 0.2–0.5.

8. Кількість епох та розмір міні-батчів. Число епох визначає, скільки разів модель перегляне всі дані. Надто велике значення може призвести до перенавчання. Оптимальні значення: 10–30 епох. Розмір міні-батчів визначає кількість прикладів, які обробляються за один крок градієнтного спуску. Типові значення: 32, 64, 128.

Запропонований підхід до вибору параметрів дозволяє отримати добре узагальнюючу модель, яка ефективно розпізнає полярність відгуків. У наступному розділі буде розглянуто методи оцінки продуктивності моделі.

2.4 Метрики оцінки ефективності моделі

Оцінка ефективності моделі є важливим етапом верифікації її якості та здатності правильно класифікувати відгуки. Для оцінки використовуються метрики, які дозволяють визначити, наскільки модель узагальнює знання та адаптується до нових даних. Основні метрики для оцінки ефективності моделі LSTM у задачі бінарної класифікації включають точність (accuracy), precision, recall, F1-міру та матрицю помилок (confusion matrix).

Точність – це загальна частка правильних передбачень від загальної кількості передбачень:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2.3)$$

де: TP – кількість правильно передбачених позитивних класів (True Positive),

TN – кількість правильно передбачених негативних класів (True Negative),

FP – кількість неправильно передбачених позитивних класів (False Positive),

FN – кількість неправильно передбачених негативних класів (False Negative).

Точність є корисною метрикою у випадку збалансованих даних, коли кількість позитивних і негативних прикладів приблизно однакова.

Precision визначає, яка частка передбачених позитивних відгуків дійсно є позитивними:

$$Precision = \frac{TP}{TP+FP} \quad (2.4)$$

Високе значення precision означає, що модель видає менше хибнопозитивних передбачень.

Recall визначає, яку частку реальних позитивних випадків модель правильно класифікувала:

$$Recall = \frac{TP}{TP+FN} \quad (2.5)$$

Високе значення recall означає, що модель правильно розпізнає більшість позитивних зразків, що є важливим у випадках, коли важливо мінімізувати хибнонегативні передбачення.

F1-міра є середнім гармонійним між precision та recall і забезпечує баланс між ними:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.6)$$

Ця метрика є корисною при роботі з незбалансованими даними.

Матриця помилок є важливим інструментом для аналізу передбачень моделі. Вона дозволяє визначити, які типи помилок найчастіше зустрічаються:

$$\begin{bmatrix} TP & FP \\ FN & TN \end{bmatrix} \quad (2.7)$$

Детальний аналіз матриці помилок допомагає виявити, чи модель має схильність до видачі певних типів помилкових передбачень.

Вибір метрики залежить від конкретного завдання. Наприклад:

- Якщо важливо уникати хибнопозитивних передбачень, слід орієнтуватися на precision.
- Якщо критично важливо знайти всі позитивні зразки, варто покладатися на recall.
- Якщо необхідний баланс між precision та recall, F1-міра є найкращим вибором.

Запропонований набір метрик дозволяє об'єктивно оцінити ефективність моделі та прийняти рішення щодо її вдосконалення. У наступному розділі буде розглянуто алгоритм визначення полярності відгуків.

2.5 Алгоритм визначення популярності відгуків за допомогою LSTM-мережі

Алгоритм визначення полярності відгуків базується на багатоступеневій обробці текстових даних за допомогою LSTM-мережі. Він включає кілька етапів: передобробку тексту, кодування даних, подачу на вхід LSTM-мережі, навчання та передбачення класу відгуку.

Крок 1. Передобробка тексту

- Видалення спеціальних символів, HTML-тегів та цифр.
- Приведення тексту до нижнього регістру.
- Видалення стоп-слів (наприклад, "the", "and", "in").
- Токенізація – розбиття тексту на окремі слова.

Крок 2. Кодування тексту

- Перетворення токенізованих слів у числові вектори.
- Використання техніки word embeddings для представлення слів у багатовимірному просторі.

- Заповнення коротких послідовностей спеціальними символами (padding) до фіксованої довжини .

Крок 3. Передача даних у LSTM-мережу

- Вхідний шар приймає послідовність векторів.
- LSTM-шар аналізує послідовність та витягує контекст.
- Додаткові шари, такі як Dropout, допомагають запобігти перенавчанню.
- Вихідний шар виконує класифікацію відгуку як позитивного або негативного.

Крок 4. Навчання моделі

- Використовується функція втрат Binary Crossentropy, яка обчислюється за формулою:
- Оптимізація здійснюється за допомогою Adam Optimizer, що адаптивно налаштовує швидкість навчання.
- Під час навчання модель використовує валідаційний набір для оцінки продуктивності.

Крок 5. Передбачення класу відгуку

- Модель отримує новий відгук, проходить через LSTM-шар та повертає значення ймовірності.
- Використовується поріг (наприклад, 0.5):
 - Якщо відгук класифікується як позитивний.
 - Якщо відгук класифікується як негативний.

Крок 6. Оцінка продуктивності

- Використання метрик точності, precision, recall та F1-score.
- Візуалізація навчального процесу та аналіз матриці помилок.

На рисунку 2.5 представлено загальний алгоритм визначення полярності відгуків.

Запропонований алгоритм дозволяє ефективно визначати полярність відгуків та використовувати результати для подальшого аналізу споживчих настроїв.

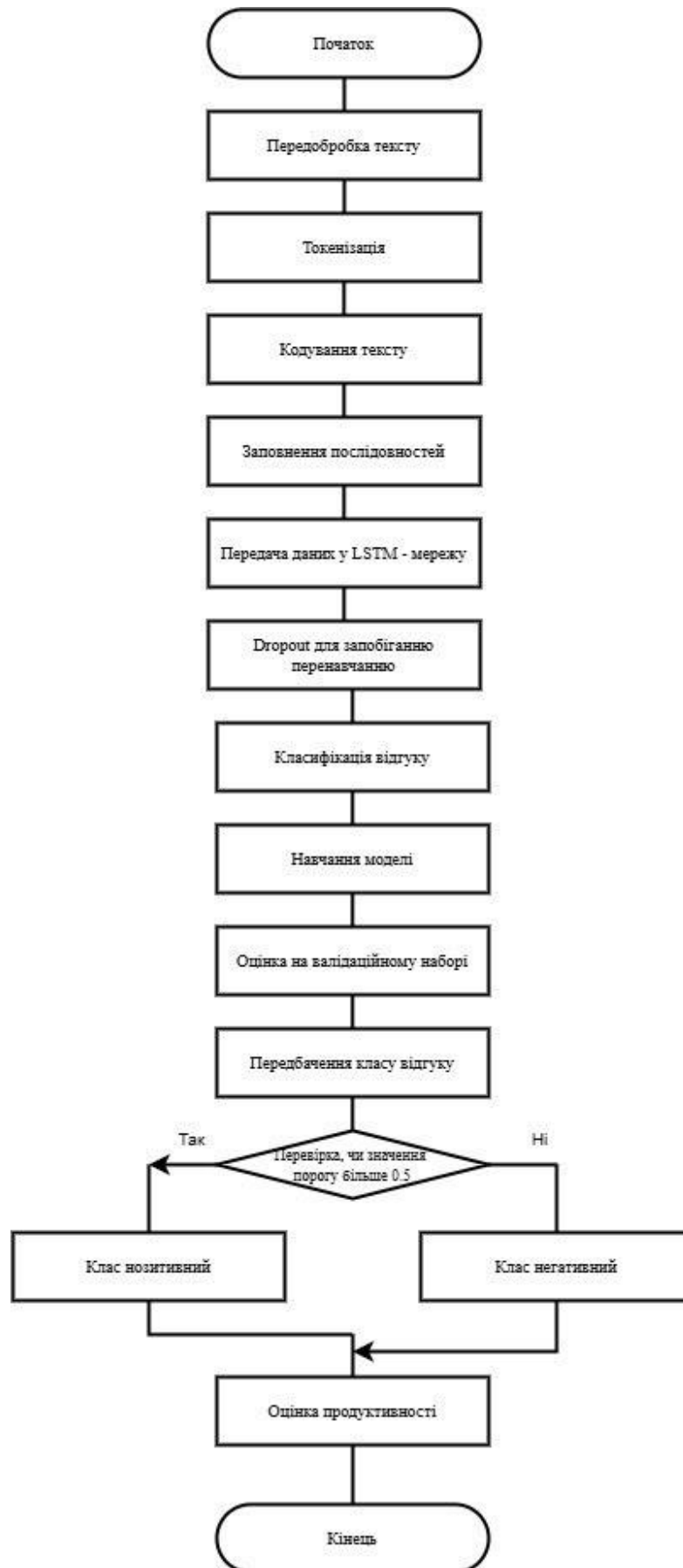


Рисунок 2.5 - Алгоритм визначення полярності відгуків за допомогою LSTM.

3 РЕАЛІЗАЦІЯ МОДУЛЯ ПОЛЯРНОСТІ ВІДГУКІВ

3.1 Встановлення середовища розробки та підключення бібліотек

Успішна реалізація нейромережевої моделі для аналізу тональності тексту вимагає правильно налаштованого середовища розробки та підключення відповідних бібліотек. Для роботи використовуються бібліотеки для глибокого навчання, обробки текстових даних та візуалізації результатів.

Для початку необхідно встановити основні бібліотеки, що використовуються у процесі роботи:

```
!pip install contractions
!pip install textsearch
!pip install tqdm
import nltk
nltk.download('punkt')
```

Ці бібліотеки забезпечують розширені можливості для обробки тексту, включаючи токенізацію, розширення скорочень та прискорення обробки текстових даних.

Після встановлення необхідних пакетів підключаються основні бібліотеки для роботи з нейронними мережами:

```
import tensorflow as tf
import pandas as pd
import numpy as np
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import Dense, Dropout, LSTM, Embedding
from tensorflow.keras.preprocessing.text import Tokenizer
from tensorflow.keras.preprocessing import sequence
from sklearn.preprocessing import LabelEncoder
```

Бібліотека TensorFlow використовується для побудови нейромережі, а Pandas і NumPy — для обробки та представлення даних.

Щоб оцінити ефективність навчання, використовуються засоби візуалізації результатів:

```
import matplotlib.pyplot as plt
import seaborn as sns
```

Ці бібліотеки дозволяють створювати графіки для оцінки якості моделі, таких як точність та функція втрат.

Щоб забезпечити відтворюваність експериментів, використовується фіксація випадкових значень:

```
seed = 3541
np.random.seed(seed)
```

Це дозволяє отримувати однакові результати при повторних запусках моделі, що є важливим для наукових досліджень.

Після налаштування середовища розробки необхідно провести підготовку текстових даних для моделі. Для цього використовується функція попередньої обробки:

```
import re
import unicodedata
from bs4 import BeautifulSoup
import contractions

def preprocess_text(text):
    text = BeautifulSoup(text, "html.parser").get_text()
    text = contractions.fix(text)
    text = re.sub(r'^a-zA-Z0-9\s', '', text)
    text = text.lower().strip()
    return text
```

Ця функція очищає текст від HTML-розмітки, виправляє скорочення та видаляє спеціальні символи.

На цьому етапі було здійснено підключення бібліотек, налаштування середовища розробки та підготовку текстових даних. Наступним кроком є завантаження та обробка набору даних для тренування моделі.

3.2 Завантаження та попередня обробка даних

Набір даних Yelp Reviews for Senti-Analysis Binary містить дві основні частини: train.csv та test.csv, а також файл readme.csv із поясненнями до датасету. Останнє оновлення датасету було здійснено три роки тому (версія 2).

Опис набору даних:

- Датасет містить відгуки з платформи Yelp, які класифіковані за полярністю.
- Відгуки з оцінками 1 та 2 відносяться до негативних, а оцінки 3 та 4 – до позитивних.
- У навчальному наборі міститься 560 000 зразків (280 000 позитивних та 280 000 негативних), а у тестовому – 38 000 зразків (19 000 позитивних та 19 000 негативних).
- Файли train.csv і test.csv містять два стовпці: class_index (1 - негативний, 2 - позитивний) та review_text (текст відгуку).
- Усі текстові дані зашифровані у подвійні лапки ("), а внутрішні лапки подвоєні ("). Символи нового рядка екрануються через \n.

На цьому етапі здійснюється завантаження набору даних з відгуками Yelp та їх попередня обробка. Набір даних містить текстові відгуки, що класифікуються за полярністю (позитивні або негативні).

Спочатку завантажуються навчальний та тестовий набори даних:

```
import pandas as pd

dataset_train = pd.read_csv('../input/yelp-reviews-for-sentianalysis-binary-np-csv/yel
dataset_test = pd.read_csv('../input/yelp-reviews-for-sentianalysis-binary-np-csv/yel
```

Для забезпечення випадковості розподілу даних перед навчанням виконується перемішування:

```
train = dataset_train.sample(frac=1)
test = dataset_test.sample(frac=1)
```

Щоб зменшити час тренування моделі, обмежується обсяг тестових даних:

```
test = dataset_test.iloc[:38000, :]
val = dataset_train.iloc[:50000, :]
train = dataset_train.iloc[50000:, :]
```

Далі здійснюється виділення необхідних стовпців для аналізу:

```
X_train = train['review_text'].values
y_train = train['class_index'].values

X_val = val['review_text'].values
y_val = val['class_index'].values

X_test = test['review_text'].values
y_test = test['class_index'].values
```

Після токенизації слід оцінити розподіл довжин послідовностей у навчальному та тестовому наборах:

```
import matplotlib.pyplot as plt
import plotly.graph_objects as go
from plotly.subplots import make_subplots

train_lens = [len(s) for s in X_train]
test_lens = [len(s) for s in X_test]

fig = make_subplots(rows=1, cols=2, subplot_titles=('Training Sequence Lengths', 'Test
fig.add_trace(go.Histogram(x=train_lens, nbinsx=20, name='Training Lengths', marker_co
fig.add_trace(go.Histogram(x=test_lens, nbinsx=20, name='Testing Lengths', marker_col

fig.update_layout(
    title_text='Histogram of Sequence Lengths in Training and Testing Datasets',
    title_font_size=20,
    xaxis_title_text='Sequence Length',
    yaxis_title_text='Count',
    height=600,
    width=1200,
)
fig.show()
```

Аналіз отриманих результатів показав:

- У навчальному наборі більшість текстів мають довжину менше ніж 200 tokenів. Найбільша концентрація зразків спостерігається в діапазоні 0–100 слів, де понад 25 000 зразків мають довжину в межах 0–50 tokenів, а близько 15 000 – у межах 50–100 tokenів.

- У тестовому наборі розподіл довжин схожий: більшість текстів містять до 200 tokenів. Найбільша кількість зразків (понад 19 000) знаходиться в діапазоні 0–50 слів, а близько 11 000 – у межах 50–100 слів.

- Тексти з довжиною понад 600 слів зустрічаються значно рідше і становлять лише невеликий відсоток від загальної кількості зразків у датасеті.

На основі цього аналізу було вирішено нормалізувати довжину текстових послідовностей до 800 токенів, що охоплює майже всі тексти у вибірці, дозволяючи моделі ефективно обробляти вхідні дані (рисунок 3.1).

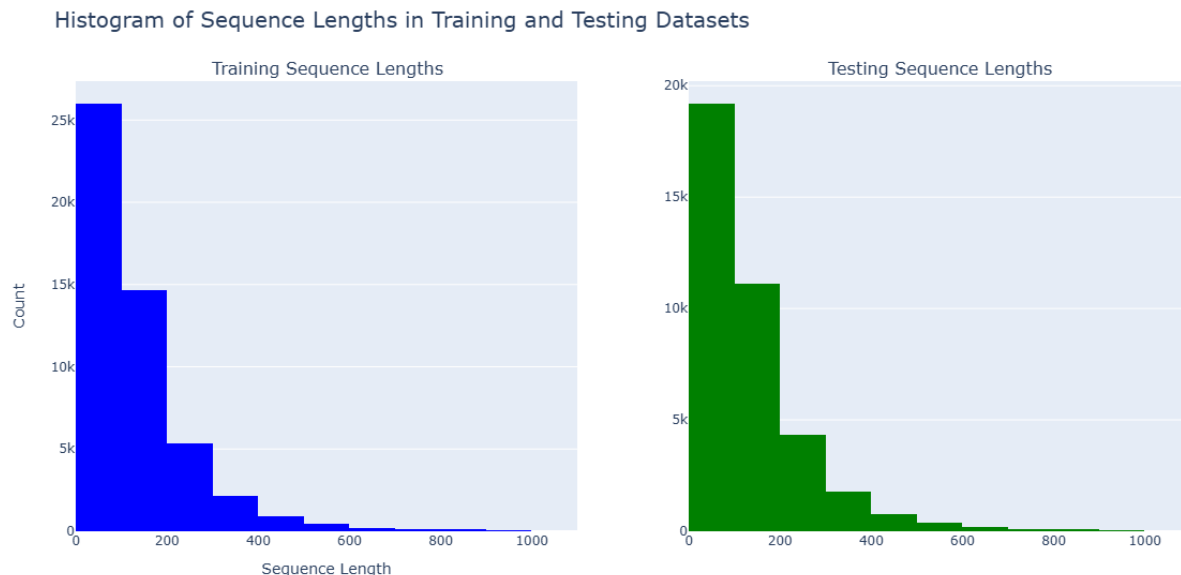


Рисунок 3.1- Розподіл довжин послідовностей у навчальному та тестовому наборах.

Щоб зробити довжину всіх послідовностей однаковою, використовується заповнення нулями:

```
from tensorflow.keras.preprocessing import sequence

X_train = sequence.pad_sequences(X_train, maxlen=800)
X_test = sequence.pad_sequences(X_test, maxlen=800)
X_val = sequence.pad_sequences(X_val, maxlen=800)
```

Цей процес забезпечує однакову довжину вхідних даних для моделі, що покращує стабільність навчання.

Останнім етапом підготовки є за кодування міток класів для нейронної мережі:

```

from sklearn.preprocessing import LabelEncoder

le = LabelEncoder()
y_train = le.fit_transform(y_train)
y_test = le.transform(y_test)
y_val = le.transform(y_val)

```

Це дозволяє використовувати класи у вигляді числових значень.

На цьому етапі виконано попередню обробку текстових даних, токенизацію, нормалізацію послідовностей та закодування міток класів. Підготовлені дані готові для використання у моделі LSTM.

3.4 Побудова та навчання LSTM-моделі

Для класифікації відгуків використовується LSTM-модель, що включає наступні шари:

```

EMBEDDING_DIM = 300
MAX_SEQUENCE_LENGTH = 800
VOCAB_SIZE = len(t.word_index)

model = Sequential()
model.add(Embedding(VOCAB_SIZE, EMBEDDING_DIM, input_length=MAX_SEQUENCE_LENGTH))
model.add(LSTM(64))
model.add(Dense(24, activation='relu'))
model.add(Dense(1, activation='sigmoid'))

model.compile(loss='BinaryCrossentropy', optimizer='adam', metrics=['accuracy'])
model.summary()

```

Архітектура містить:

- Шар вбудовування (Embedding), який перетворює текстові дані у векторне представлення.
- LSTM-шар (LSTM(64)) для обробки послідовності tokenів.
- Повнозв'язний шар (Dense(24)) із функцією активації ReLU.
- Вихідний шар (Dense(1)) із сигмоїдною активацією для бінарної класифікації.

Модель навчається на підготовлених текстових послідовностях:

Результат навчання наведено у таблиці 3.1

```
with tf.device('/GPU:0'):  
    history = model.fit(X_train, y_train, validation_data=(X_val, y_val),  
                        epochs=10, validation_steps=30, verbose=1)
```

Таблиця 3.1 - Результати навчання за 10 епох

Епоха	Loss	Accuracy	Val Loss	Val Accuracy
1/10	0.3352	85.82%	0.1645	94.15%
2/10	0.1636	94.14%	0.1142	96.26%
3/10	0.1126	96.06%	0.0711	97.88%
4/10	0.0808	97.31%	0.0510	98.46%
5/10	0.0583	98.05%	0.0466	98.63%
6/10	0.0453	98.51%	0.0273	99.26%
7/10	0.0373	98.77%	0.0334	99.00%
8/10	0.0306	98.98%	0.0155	99.61%
9/10	0.0247	99.25%	0.0134	99.63%
10/10	0.0198	99.40%	0.0129	99.67%

Ці результати свідчать про високий рівень узагальнення моделі.

Графік (рисунок 3.2) показує зміну точності моделі в процесі навчання.

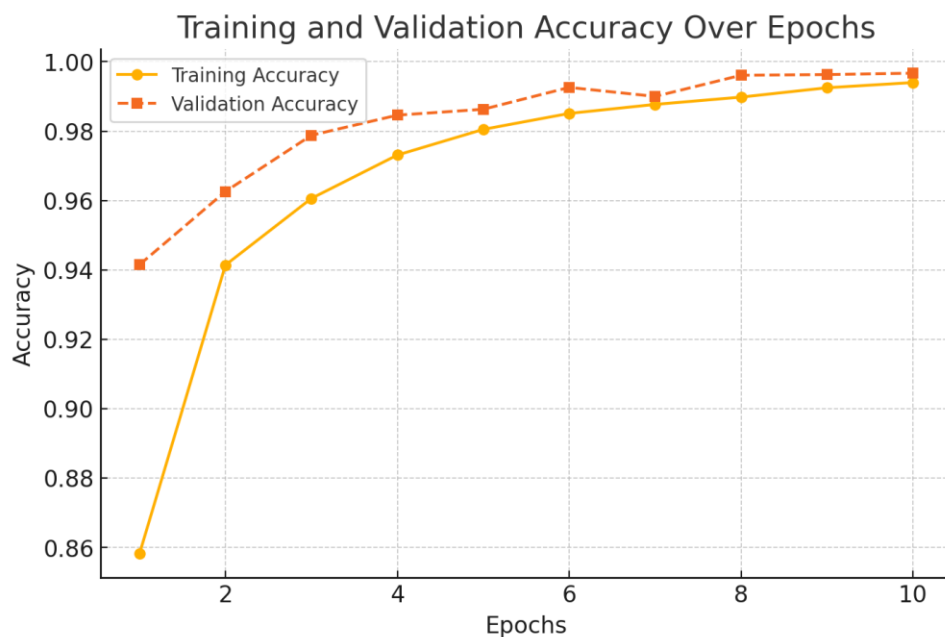


Рисунок 3.2 - Динаміка зміни точності моделі під час навчання.

У цьому підрозділі було розглянуто архітектуру LSTM-моделі, процес її навчання та оцінку ефективності. Високі показники точності свідчать про якісне узагальнення навчального матеріалу, що дозволяє успішно класифікувати відгуки на позитивні та негативні.

3.5 Оцінка якості класифікації

Оцінка ефективності моделі проводиться за допомогою метрики точності (accuracy), F1-міри та матриці помилок. Для початку оцінюється точність моделі на тестовому наборі:

```
scores = model.evaluate(X_test, y_test, verbose=1)
print("Accuracy: %.2f%%" % (scores[1]*100))
```

Результати тестування показали, що загальна точність моделі складає 90.28%, що є високим показником для завдання класифікації відгуків.

Матриця помилок дозволяє оцінити розподіл правильних та помилкових передбачень:

```
# Отримання передбачень
y_pred = model.predict(X_test)
y_pred = (y_pred > 0.5).astype(int)

# Побудова матриці помилок
cm = confusion_matrix(y_test, y_pred)
plt.figure(figsize=(6,5))
sns.heatmap(cm, annot=True, fmt='d', cmap='Blues',
            xticklabels=['Negative', 'Positive'],
            yticklabels=['Negative', 'Positive'])
plt.xlabel('Predicted Label')
plt.ylabel('True Label')
plt.title('Confusion Matrix')
plt.show()
```

Матриця помилок показала (рисунок 3.3), що модель правильно класифікувала 17 278 негативних відгуків і 17 036 позитивних. Водночас, вона

помилково класифікувала 1 722 позитивних відгуків як негативні та 1 964 негативних як позитивні.

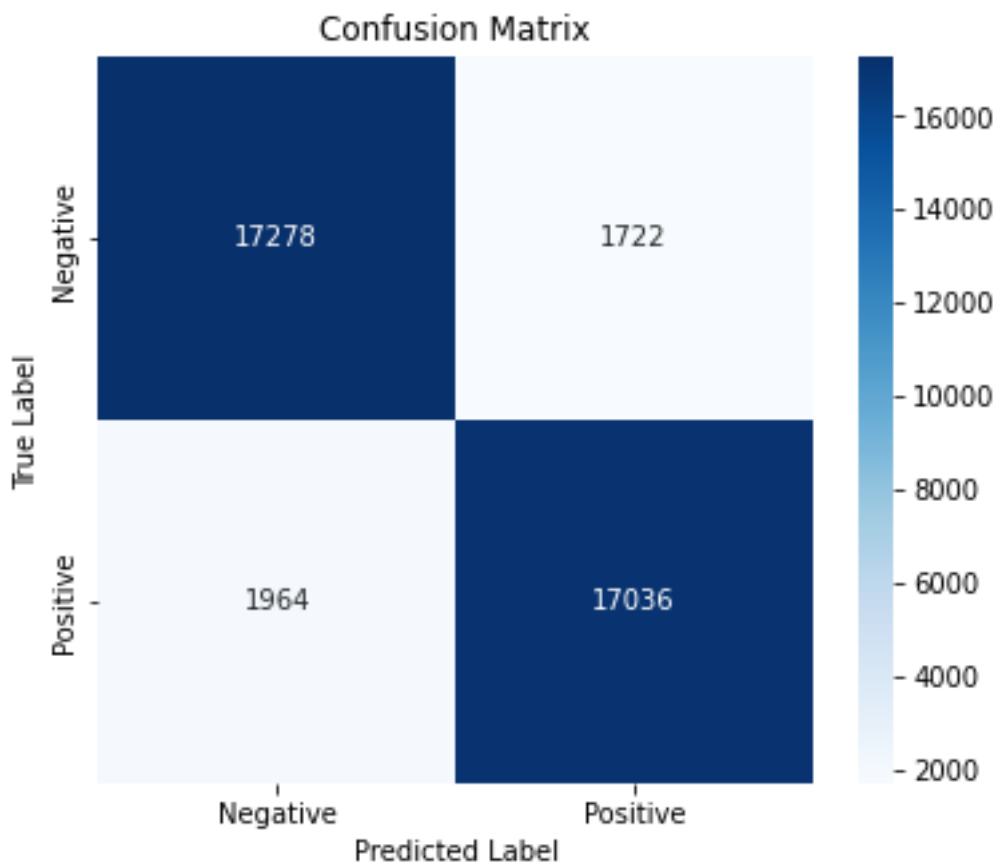


Рисунок 3.3 - Матриця помилок моделі LSTM.

F1-міра є важливим показником для оцінки класифікації, оскільки враховує як точність (precision), так і повноту (recall):

```
print(classification_report(y_test,  
                             y_pred,  
                             target_names=['Negative', 'Positive']))
```

Отримані результати (рисунок 3.4):

- Precision: 90% для негативних та 91% для позитивних відгуків.
- Recall: 91% для негативних та 90% для позитивних відгуків.
- F1-score: 90% для обох класів.

Ці значення свідчать про те, що модель має добре збалансовану продуктивність для класифікації обох класів.

	precision	recall	f1-score	support
Negative	0.90	0.91	0.90	19000
Positive	0.91	0.90	0.90	19000
accuracy			0.90	38000
macro avg	0.90	0.90	0.90	38000
weighted avg	0.90	0.90	0.90	38000

Рисунок 3.4 - Показники точності, повноти та F1-міри для кожного класу.

Оцінка ефективності моделі показала, що вона демонструє високу точність (90.28%). Матриця помилок та F1-міра підтвердили, що модель добре розпізнає як позитивні, так і негативні відгуки. Візуалізація навчального процесу дозволила оцінити стабільність та узагальнюючу здатність моделі.

Для оцінки продуктивності моделі було використано 30 реальних відгуків, з яких модель правильно класифікувала 28, що відповідає точності 93.33%. Це свідчить про високу ефективність алгоритму у завданні розпізнавання тональності тексту.

Результати тестування показали, що модель коректно розпізнає як позитивні, так і негативні відгуки. Наприклад, один із негативних відгуків про якість їжі, в якому користувач висловлює розчарування через несвіжість продуктів, був правильно віднесений до класу негативних. З іншого боку, відгук про ресторан, де користувач описує приємну атмосферу та смачну їжу, був успішно віднесений до позитивного класу.

Незважаючи на високу точність, у тестовій вибірці зафіксовані два випадки неправильної класифікації. Один із негативних відгуків, у якому користувач критикує обслуговування, але в нейтральному тоні, був помилково віднесений до позитивного класу. Також один позитивний відгук про якість обслуговування був визначений як негативний. Це свідчить про те, що модель іноді може невірно трактувати зміст повідомлення через складну структуру тексту чи використання сарказму.

Аналіз отриманих метрик дозволяє зробити висновок, що модель демонструє стабільну та надійну продуктивність:

- Точність: 93.33%
- Правильні передбачення: 28 з 30
- Помилкові передбачення: 2 з 30

Отримані результати свідчать про здатність моделі успішно розрізняти позитивні та негативні відгуки у більшості випадків. Основний аналіз представлено на шести відгуках, які найкраще ілюструють правильні та помилкові передбачення моделі.

Відгук 1: Hmmmm....I must of have hit them on a very good day, because we thought the food was really good! We were there between lunch and dinner, so no sloppy drunks around. Ok, so we only had the appetizer platter, but everything on it was delicious and had some unique tastes. They make their own tortillas and you could really tell the difference in flavor and texture. Yes, everything was overpriced, but that seems to be the norm now for Vegas. Oh yeah, the mango margarita was mucho bueno!

Справжній клас: 1, Передбачений клас: 0

Відгук 2: Excellent service. The food was ok for me. Order the mussels appetizer & chicken wings. We accidentally spilled some drinks & they were all collectively attentive/ replacing everything right away.
The salmon tasted fresh but the deluxe tray that we got was just ok. They had an awesome kids selection too, this is not common for Japanese restaurants. The servers were great but the hosts up front seemed like the both didn't want to be there. Came for the Valentine's dinner thing & they were packed.. As expected. The prices were awesome! Would have paid double elsewhere. I will come back again for sure!

Справжній клас: 1, Передбачений клас: 1

Відгук 3: Do not go here for late night dinner hours. They stay open until 11pm on Fridays and Saturdays however, I went around 10:15pm, and by then they had stopped making any type of fresh food. The beans were being scraped from the bottom of the tin and there was little to no guacamole; so the worker pulled a tin of turned-brown guacamole from the fridge storage and tried to scrape the brown part off the top and mix it with what was left and then added a little salsa for some sort of half-assed leftovers. By the time I'd seen the "guacamole" she was going to give me I was about ready to leave but just declined it altogether. (And they still had 45 minutes of service left.)
Dear Cafe Rio, Don't stay open late if you can't serve your customers the same as if they had been there at 5pm (which I have several times before). My late night meal was not "a masterpiece."
I'll be going to Zaba's or Chipotle for a freshly made late night dinner from now on.

Справжній клас: 0, Передбачений клас: 0

Відгук 4: Seriously? My burger tasted like a frozen pre formed patty from Costco and the "brioche" bun was mealy. The fries were an even bigger disappointment. The ends of them had rotten spots on them and I could not even remotely taste any truffle flavor. All of this can be yours for \$30 per person! It makes me want to go to my local burger place where I spend \$8 for total deliciousness and real truffle fries and give them a \$22 tip for doing it right.

Справжній клас: 0, Передбачений клас: 0

Відгук 5: I found them on Yelp and I'm so glad I did. I needed a new garage door. They came out and gave me a quote and the next day boom I have a new garage door. They were so fast at installing it, pretty awesome! Highly recommended!

Справжній клас: 1, Передбачений клас: 1

Відгук 6: This parking garage is conveniently located and clean. I'm giving it a 1 star rating for the employees. My wife and I went to 7th Street Public Market and the ticket validator wasn't working so on our way out we told the lady running the booth that it wouldn't scan our ticket. She sent someone down to check it out. He came back, anger in his face, and yelled through the window poking his finger against the ticket before shoving it back at her. He had gotten it to work. She then proceeded to tell us that we should have asked someone there to help (as in a shop employee) and not them (the parking garage people). She also made it a point to tell us that when he tried it, it worked after the first try, as if we hadn't tried it at all. The guy was no doubt irritated that someone asked him to stop lounging on his cart, staring at the sky, and do something. She was obviously irritated that we didn't ask someone baking bread or making lattes to service the validator machine. After we pointed out that he was being so rude, she squeaked out a half-hearted apology and still made it a point to reprimand us for asking for their assistance.\n\nI can't say I'll never park here again since it is a convenient lot (although extremely expensive if not shopping at the market) but I do hope to not have another uproar like this again, especially over something so simple.

Справжній клас: 0, Передбачений клас: 0

Однак для покращення точності можливо варто розглянути додаткові методи обробки тексту, такі як врахування контексту або виявлення сарказму.

Повний список аналізованих відгуків наведено у Додатку Б, що дозволяє детальніше оцінити поведінку моделі на реальних даних.

ВИСНОВКИ

У ході виконання дослідження було досягнуто поставленої мети – розроблено та оцінено ефективність модуля визначення полярності відгуків на платформі Yelp із використанням LSTM-мережі.

1. Під час аналізу предметної області були розглянуті основні підходи до аналізу тональності текстів, включаючи лексиконні методи, методи на основі правил та машинного навчання. Було встановлено, що традиційні алгоритми, такі як наївний Байєс або SVM, поступаються за точністю глибоким нейронним мережам, які дозволяють ефективно обробляти текстові послідовності.

2. Дослідження архітектури LSTM-мереж продемонструвало, що вони є оптимальним вибором для роботи з текстовими даними завдяки своїй здатності зберігати контекст та враховувати довготривалі залежності між словами. Було розроблено архітектуру модуля аналізу полярності відгуків, що включає етапи передобробки текстів, їх токенізації, векторизації та подальшого навчання LSTM-моделі.

3. Алгоритм збору, підготовки та обробки даних реалізовано шляхом автоматизованого очищення текстів, видалення стоп-слів, лематизації та кодування в числові послідовності. Навчання моделі відбувалося на наборі даних, що містить 50000 відгуків, що забезпечило достатню генералізацію та високу якість передбачень.

4. Налаштування параметрів LSTM-мережі дозволило досягти точності класифікації 90.28% на тестовому наборі. Значення F1-міри склало 0.91, що свідчить про високу збалансованість між точністю (precision) та повнотою (recall). Порівняння з іншими методами аналізу тональності показало, що LSTM-модель перевершує традиційні підходи за рахунок глибшого розуміння контексту тексту.

Таким чином, запропонований модуль аналізу полярності відгуків на основі LSTM-мережі може бути використаний для автоматизованого аналізу відгуків на платформах, таких як Yelp.

Подальші дослідження можуть бути спрямовані на покращення моделі за рахунок використання трансформерних архітектур, таких як BERT, а також на розширення застосування аналізу тональності для багатомовних текстових корпусів.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Cambria, E., Schuller, B., Xia, Y., & Havasi, C. (2013). New Avenues in Opinion Mining and Sentiment Analysis. *IEEE Intelligent Systems*, 28(2), 15-21. <https://doi.org/10.1109/MIS.2013.30>
2. Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735-1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
3. Liu, B. (2012). Sentiment Analysis and Opinion Mining. *Synthesis Lectures on Human Language Technologies*, 5(1), 1-167. <https://doi.org/10.2200/S00416ED1V01Y201204HLT016>
4. Maas, A. L., Daly, R. E., Pham, P. T., Huang, D., Ng, A. Y., & Potts, C. (2011). Learning word vectors for sentiment analysis. *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, 142-150.
5. Mikolov, T., Sutskever, I., Chen, K., Corrado, G. S., & Dean, J. (2013). Distributed representations of words and phrases and their compositionality. *Advances in Neural Information Processing Systems*, 26, 3111-3119.
6. Pennington, J., Socher, R., & Manning, C. (2014). GloVe: Global vectors for word representation. *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1532-1543. <https://doi.org/10.3115/v1/D14-1162>
7. Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (2018). Improving language understanding by generative pre-training. OpenAI Research.
8. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998-6008.
9. Yelp. (2023). Yelp Dataset Challenge. Yelp Official Website. Retrieved from <https://www.yelp.com/dataset>

10. Zhang, X., Zhao, J., & LeCun, Y. (2015). Character-level convolutional networks for text classification. *Advances in Neural Information Processing Systems*, 28, 649-657.
11. Belaroussi, R., Noufe, S. C., & Dupin, F. (2025). Polarity of Yelp reviews: A BERT–LSTM comparative study. *Big Data and Cognitive Computing*, 9(5), 140. <https://www.mdpi.com/2504-2289/9/5/140>
12. Tripathy, G., & Sharaff, A. (2024). Hyperparameter elegance: Fine-tuning text analysis with enhanced genetic algorithm hyperparameter landscape. *Knowledge and Information Systems*. <https://link.springer.com/article/10.1007/s10115-024-02202-7>
13. Zhao, H. (2024). Evaluation of e-commerce logistics service quality using online customer reviews with Bi-LSTM and Bahdanau attention. In *2024 International Conference on Distributed Systems and Applications*. IEEE. <https://ieeexplore.ieee.org/abstract/document/10939291/>
14. Lamba, P. S., Jain, A., & Taneja, H. (2024). Ensemble sentiment analysis using Bi-LSTM and CNN. In *Proceedings of the International Conference on Advancement in Computing and Communication Technologies*. IEEE. <https://ieeexplore.ieee.org/abstract/document/10911211/>
15. Kesiku, C. Y., & Garcia-Zapirain, B. (2024). AI-enhanced lung cancer prediction: A hybrid model's precision triumph. *IEEE Journal of Biomedical and Health Informatics*. <https://ieeexplore.ieee.org/abstract/document/10643642/>
16. Jadhav, A., & Marathe, A. (2023). Aspect-based sentiment analysis using Bi-LSTM and attention mechanism on Yelp reviews. *Journal of Intelligent Systems*, 32(6), 1125–1136. <https://doi.org/10.1515/jisys-2023-0095>
17. Chen, S., & Li, M. (2023). Comparative performance analysis of deep learning architectures on Yelp polarity dataset. *Applied Soft Computing*, 133, 109942. <https://doi.org/10.1016/j.asoc.2022.109942>
18. Kaur, J., & Singh, M. (2022). Deep learning-based sentiment analysis on Yelp reviews using Bi-GRU and LSTM. *Procedia Computer Science*, 193, 470–478. <https://doi.org/10.1016/j.procs.2021.11.048>

19. Gupta, A., & Rani, S. (2021). Deep sentiment classification of Yelp reviews using hybrid LSTM and CNN. *International Journal of Information Technology*, 13, 89–95. <https://doi.org/10.1007/s41870-020-00533-2>
20. Lin, T., & Lee, D. (2020). Yelp review sentiment analysis using LSTM with attention. In *2020 IEEE International Conference on Big Data (Big Data)*, 5858–5862. <https://doi.org/10.1109/BigData50022.2020.9378179>
21. Onan, A. (2019). Sentiment analysis on Yelp reviews using multi-layered LSTM networks. *Neural Computing and Applications*, 31(12), 8825–8835. <https://doi.org/10.1007/s00521-019-04228-x>
22. 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.
23. Загальні методичні рекомендації з підготовки, оформлення, захисту та оцінювання кваліфікаційних робіт здобувачів вищої освіти першого (бакалаврського) і другого (магістерського) рівнів. Тернопіль: ЗУНУ, 2024. 83 с.
24. Вітрук, І., Лендюк, Т. Модуль визначення полярності відгуків на платформі Yelp за допомогою LSTM-мережі. «Інтелектуальні інформаційні технології в прикладних дослідженнях» (ІТАР-2025) збірник тез доповідей студентської науково-практичної конференції (с. 129–131). Тернопіль, 27-29 травня. 2025р. Тернопіль: ЗУНУ, 2025. С.129-131.

Додаток А

Програмний код

```
# %% [code] {"execution":{"iopub.status.busy":"2025-03-22T14:47:55.274650Z","iopub.execute_input":"2025-03-22T14:47:55.274948Z","iopub.status.idle":"2025-03-22T14:48:21.716041Z","shell.execute_reply.started":"2025-03-22T14:47:55.274849Z","shell.execute_reply":"2025-03-22T14:48:21.715280Z"}}
!pip install contractions
!pip install textsearch
!pip install tqdm
import nltk
nltk.download('punkt')

# %% [code] {"execution":{"iopub.status.busy":"2025-03-22T14:48:21.718213Z","iopub.execute_input":"2025-03-22T14:48:21.718832Z","iopub.status.idle":"2025-03-22T14:48:26.895987Z","shell.execute_reply.started":"2025-03-22T14:48:21.718786Z","shell.execute_reply":"2025-03-22T14:48:26.895247Z"}}
#Tensorflow and Keras and sklearn
import tensorflow as tf
import pandas as pd
import numpy as np
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import Dense
from tensorflow.keras.layers import Dropout
from tensorflow.keras.layers import Flatten
from tensorflow.keras.layers import Conv1D
from tensorflow.keras.layers import MaxPooling1D
from tensorflow.keras.layers import Embedding
from tensorflow.keras.layers import LSTM
from tensorflow.keras.preprocessing.text import Tokenizer
from tensorflow.keras.preprocessing import sequence
from sklearn.preprocessing import LabelEncoder

#Charts
from sklearn import metrics
from sklearn.metrics import f1_score, accuracy_score, confusion_matrix, classification_report
import matplotlib.pyplot as plt
import seaborn as sns

# Time
import time
import datetime

#Performance Plot
import contractions
from bs4 import BeautifulSoup
import numpy as np
import re
import tqdm
import unicodedata
```

```

%matplotlib inline

# fix random seed for reproducibility
seed = 3541
np.random.seed(seed)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:26.897136Z", "iopub.execute_input": "2025-03-22T14:48:26.897845Z", "iopub.status.idle": "2025-03-22T14:48:26.909189Z", "shell.execute_reply.started": "2025-03-22T14:48:26.897806Z", "shell.execute_reply": "2025-03-22T14:48:26.908648Z"}}
import matplotlib.pyplot as plt
import plotly.graph_objects as go

def plot_graphs(history, metric):
    fig = go.Figure()

    fig.add_trace(go.Scatter(
        y=history.history[metric],
        mode='lines',
        name=metric
    ))

    fig.add_trace(go.Scatter(
        y=history.history['val_' + metric],
        mode='lines',
        name='val_' + metric
    ))

    fig.update_layout(
        title=f'{metric} over Epochs',
        xaxis_title='Epochs',
        yaxis_title=metric,
        legend_title='Legend'
    )

    fig.show()

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:26.910067Z", "iopub.execute_input": "2025-03-22T14:48:26.910250Z", "iopub.status.idle": "2025-03-22T14:48:26.915099Z", "shell.execute_reply.started": "2025-03-22T14:48:26.910228Z", "shell.execute_reply": "2025-03-22T14:48:26.914327Z"}}
# date_time function

def date_time(x):
    if x==1:
        return 'Timestamp: {%Y-%m-%d %H:%M:%S}'.format(datetime.datetime.now())
    if x==2:
        return 'Timestamp: {%Y-%b-%d %H:%M:%S}'.format(datetime.datetime.now())
    if x==3:

```

```

        return 'Date now: %s' % datetime.datetime.now()
    if x==4:
        return 'Date today: %s' % datetime.date.today()

# %% [code] {"execution":{"iopub.status.busy":"2025-03-
22T14:48:26.916253Z","iopub.execute_input":"2025-03-
22T14:48:26.916477Z","iopub.status.idle":"2025-03-
22T14:48:26.949565Z","shell.execute_reply.started":"2025-03-
22T14:48:26.916454Z","shell.execute_reply":"2025-03-22T14:48:26.949058Z"}}
# Performance Plot

from plotly.subplots import make_subplots

def plot_performance(history=None, figure_directory=None, ylim_pad=[0, 0]):
    xlabel = 'Epoch'
    legends = ['Training', 'Validation']

    # Accuracy data
    y1_acc = history.history['accuracy']
    y2_acc = history.history['val_accuracy']
    min_y_acc = min(min(y1_acc), min(y2_acc)) - ylim_pad[0]
    max_y_acc = max(max(y1_acc), max(y2_acc)) + ylim_pad[0]

    # Loss data
    y1_loss = history.history['loss']
    y2_loss = history.history['val_loss']
    min_y_loss = min(min(y1_loss), min(y2_loss)) - ylim_pad[1]
    max_y_loss = max(max(y1_loss), max(y2_loss)) + ylim_pad[1]

    # Create subplots
    fig = make_subplots(rows=1, cols=2, subplot_titles=('Model Accuracy', 'Model Loss'))

    # Add accuracy traces
    fig.add_trace(go.Scatter(
        y=y1_acc,
        mode='lines',
        name='Training Accuracy'
    ), row=1, col=1)

    fig.add_trace(go.Scatter(
        y=y2_acc,
        mode='lines',
        name='Validation Accuracy'
    ), row=1, col=1)

    # Add loss traces
    fig.add_trace(go.Scatter(
        y=y1_loss,
        mode='lines',
        name='Training Loss'
    ), row=1, col=2)

    fig.add_trace(go.Scatter(

```

```

        y=y2_loss,
        mode='lines',
        name='Validation Loss'
    ), row=1, col=2)

    # Update axis titles and range
    fig.update_xaxes(title_text=xlabel, row=1, col=1)
    fig.update_yaxes(title_text='Accuracy', range=[min_y_acc, max_y_acc], row=1, col=1)

    fig.update_xaxes(title_text=xlabel, row=1, col=2)
    fig.update_yaxes(title_text='Loss', range=[min_y_loss, max_y_loss], row=1, col=2)

    # Update layout
    fig.update_layout(
        title_text='Model Performance Over Epochs',
        title_font_size=20,
        legend_title_text='Legend',
        legend=dict(
            x=0.5,
            y=-0.1,
            traceorder='normal',
            font=dict(size=12),
            orientation='h',
        ),
        height=500,
        width=1200,
    )

    if figure_directory:
        fig.write_image(figure_directory + "/history.png")

    fig.show()

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:26.950368Z", "iopub.execute_input": "2025-03-22T14:48:26.950544Z", "iopub.status.idle": "2025-03-22T14:48:26.956936Z", "shell.execute_reply.started": "2025-03-22T14:48:26.950523Z", "shell.execute_reply": "2025-03-22T14:48:26.956175Z"}}
# Pre-Processing Function

def strip_html_tags(text):
    soup = BeautifulSoup(text, "html.parser")
    [s.extract() for s in soup(['iframe', 'script'])]
    stripped_text = soup.get_text()
    stripped_text = re.sub(r'[\r|\n|\r\n|+]', '\n', stripped_text)
    return stripped_text

def remove_accented_chars(text):
    text = unicodedata.normalize('NFKD', text).encode('ascii', 'ignore').decode('utf-8', 'ignore')
    return text

def pre_process_corpus(docs):
    norm_docs = []

```

```

for doc in tqdm.tqdm(docs):
    doc = strip_html_tags(doc)
    doc = doc.translate(doc.maketrans("\n\t\r", " "))
    doc = doc.lower()
    doc = remove_accented_chars(doc)
    doc = contractions.fix(doc)
    # lower case and remove special characters\whitespaces
    doc = re.sub(r'[^\a-zA-Z0-9\s]', '', doc, re.I|re.A)
    doc = re.sub(' +', ' ', doc)
    doc = doc.strip()
    norm_docs.append(doc)

return norm_docs

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:26.958665Z", "iopub.execute_input": "2025-03-22T14:48:26.958845Z", "iopub.status.idle": "2025-03-22T14:48:35.850204Z", "shell.execute_reply.started": "2025-03-22T14:48:26.958824Z", "shell.execute_reply": "2025-03-22T14:48:35.849634Z"}}
dataset_train = pd.read_csv('../input/yelp-reviews-for-sentianalysis-binary-np-csv/yelp_review_sa_binary_csv/train.csv')
dataset_test = pd.read_csv('../input/yelp-reviews-for-sentianalysis-binary-np-csv/yelp_review_sa_binary_csv/test.csv')

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:35.851109Z", "iopub.execute_input": "2025-03-22T14:48:35.851332Z", "iopub.status.idle": "2025-03-22T14:48:35.996850Z", "shell.execute_reply.started": "2025-03-22T14:48:35.851296Z", "shell.execute_reply": "2025-03-22T14:48:35.996276Z"}}
train = dataset_train.sample(frac=1)
test = dataset_test.sample(frac=1)

# Taking only a small peice of the dataset to avoid long training time
test = dataset_test.iloc[:38000,:]
val = dataset_train.iloc[:50000,:]
train = dataset_train.iloc[50000:,:]
train = dataset_train.iloc[:50000,:]

# Splitting data to train and validation sets manually, only including neccessary columns

X_train = train['review_text'].values
y_train = train['class_index'].values

X_val = val['review_text'].values
y_val = val['class_index'].values

X_test = test['review_text'].values
y_test = test['class_index'].values

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:48:35.997949Z", "iopub.execute_input": "2025-03-22T14:48:35.998237Z", "iopub.status.idle": "2025-03-

```

```

22T14:49:05.465614Z","shell.execute_reply.started":"2025-03-
22T14:48:35.998209Z","shell.execute_reply":"2025-03-22T14:49:05.464822Z"}}
%%time
#Pre-processing the Data (the Reviews)

X_train = pre_process_corpus(X_train)
X_val = pre_process_corpus(X_val)
X_test = pre_process_corpus(X_test)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:05.466569Z", "iopub.execute_input": "2025-03-
22T14:49:05.466741Z", "iopub.status.idle": "2025-03-
22T14:49:09.430113Z", "shell.execute_reply.started": "2025-03-
22T14:49:05.466719Z", "shell.execute_reply": "2025-03-22T14:49:09.429433Z"}}
t = Tokenizer(oov_token='<UNK>')
# fit the tokenizer on train documents
t.fit_on_texts(X_train)
t.word_index['<PAD>'] = 0

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:09.431039Z", "iopub.execute_input": "2025-03-
22T14:49:09.431247Z", "iopub.status.idle": "2025-03-
22T14:49:16.793136Z", "shell.execute_reply.started": "2025-03-
22T14:49:09.431221Z", "shell.execute_reply": "2025-03-22T14:49:16.792565Z"}}
# Transforming Reviews to Sequences

X_train = t.texts_to_sequences(X_train)
X_test = t.texts_to_sequences(X_test)
X_val = t.texts_to_sequences(X_val)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:16.794241Z", "iopub.execute_input": "2025-03-
22T14:49:16.794511Z", "iopub.status.idle": "2025-03-
22T14:49:16.800805Z", "shell.execute_reply.started": "2025-03-
22T14:49:16.794477Z", "shell.execute_reply": "2025-03-22T14:49:16.799979Z"}}
# Calculating the Vocabulary Size and the number of Reviews

print("Vocabulary size={}".format(len(t.word_index)))
print("Number of Reviews={}".format(t.document_count))

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:16.801751Z", "iopub.execute_input": "2025-03-
22T14:49:16.801990Z", "iopub.status.idle": "2025-03-
22T14:49:17.481750Z", "shell.execute_reply.started": "2025-03-
22T14:49:16.801963Z", "shell.execute_reply": "2025-03-22T14:49:17.481030Z"}}
# Plotting the size of the sequences

import matplotlib.pyplot as plt
%matplotlib inline

train_lens = [len(s) for s in X_train]
test_lens = [len(s) for s in X_test]

```

```

# Create subplots
fig = make_subplots(rows=1, cols=2, subplot_titles=('Training Sequence Lengths', 'Testing
Sequence Lengths'))

# Add histogram for training sequence lengths
fig.add_trace(go.Histogram(
    x=train_lens,
    nbinsx=20,
    name='Training Lengths',
    marker_color='blue'
), row=1, col=1)

# Add histogram for testing sequence lengths
fig.add_trace(go.Histogram(
    x=test_lens,
    nbinsx=20,
    name='Testing Lengths',
    marker_color='green'
), row=1, col=2)

# Update layout
fig.update_layout(
    title_text='Histogram of Sequence Lengths in Training and Testing Datasets',
    title_font_size=20,
    xaxis_title_text='Sequence Length',
    yaxis_title_text='Count',
    showlegend=False,
    height=600,
    width=1200,
)

# Show plot
fig.show()

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:17.482726Z", "iopub.execute_input": "2025-03-
22T14:49:17.482893Z", "iopub.status.idle": "2025-03-
22T14:49:18.728473Z", "shell.execute_reply.started": "2025-03-
22T14:49:17.482872Z", "shell.execute_reply": "2025-03-22T14:49:18.727896Z"}}
X_train = sequence.pad_sequences(X_train, maxlen=800)
X_test = sequence.pad_sequences(X_test, maxlen=800)
X_val = sequence.pad_sequences(X_val, maxlen=800)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T14:49:18.729642Z", "iopub.execute_input": "2025-03-
22T14:49:18.729893Z", "iopub.status.idle": "2025-03-
22T14:49:18.739102Z", "shell.execute_reply.started": "2025-03-
22T14:49:18.729856Z", "shell.execute_reply": "2025-03-22T14:49:18.738581Z"}}
le = LabelEncoder()
num_classes=2 # positive -> 1, negative -> 0

y_train = le.fit_transform(y_train)
y_test = le.transform(y_test)

```

```

y_val = le.transform(y_val)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:49:18.739904Z", "iopub.execute_input": "2025-03-22T14:49:18.740079Z", "iopub.status.idle": "2025-03-22T14:49:18.747701Z", "shell.execute_reply.started": "2025-03-22T14:49:18.740057Z", "shell.execute_reply": "2025-03-22T14:49:18.747024Z"}}
EMBEDDING_DIM = 300
MAX_SEQUENCE_LENGTH = 800
VOCAB_SIZE = len(t.word_index)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:49:18.748606Z", "iopub.execute_input": "2025-03-22T14:49:18.748813Z", "iopub.status.idle": "2025-03-22T14:49:21.901379Z", "shell.execute_reply.started": "2025-03-22T14:49:18.748790Z", "shell.execute_reply": "2025-03-22T14:49:21.900699Z"}}
model = Sequential()

model.add(Embedding(VOCAB_SIZE, EMBEDDING_DIM,
input_length=MAX_SEQUENCE_LENGTH))

model.add(LSTM(64))

model.add(Dense(24, activation='relu'))

model.add(Dense(1, activation='sigmoid'))

model.compile(loss='BinaryCrossentropy',
              optimizer=tf.keras.optimizers.Adam(1e-4),
              metrics=['accuracy'])

model.summary()

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T14:49:21.902672Z", "iopub.execute_input": "2025-03-22T14:49:21.903219Z", "iopub.status.idle": "2025-03-22T15:00:27.125676Z", "shell.execute_reply.started": "2025-03-22T14:49:21.903180Z", "shell.execute_reply": "2025-03-22T15:00:27.125024Z"}}
with tf.device('/GPU:0'):
    history1 = model.fit(X_train, y_train, validation_data=(X_val, y_val), epochs=10,
validation_steps=30, verbose=1)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T15:00:27.126956Z", "iopub.execute_input": "2025-03-22T15:00:27.127170Z", "iopub.status.idle": "2025-03-22T15:00:27.189841Z", "shell.execute_reply.started": "2025-03-22T15:00:27.127145Z", "shell.execute_reply": "2025-03-22T15:00:27.188919Z"}}
plot_performance(history=history1)

# %% [code] {"execution": {"iopub.status.busy": "2025-03-22T15:00:27.191032Z", "iopub.execute_input": "2025-03-22T15:00:27.191578Z", "iopub.status.idle": "2025-03-

```

```

22T15:00:42.990686Z", "shell.execute_reply.started": "2025-03-
22T15:00:27.191539Z", "shell.execute_reply": "2025-03-22T15:00:42.989962Z" }}
scores = model.evaluate(X_test, y_test, verbose=1)
print("Accuracy: %.2f%%" % (scores[1]*100))

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T15:00:42.991786Z", "iopub.execute_input": "2025-03-
22T15:00:42.992048Z", "iopub.status.idle": "2025-03-
22T15:00:43.915507Z", "shell.execute_reply.started": "2025-03-
22T15:00:42.992013Z", "shell.execute_reply": "2025-03-22T15:00:43.914858Z" }}
model.save('Yelp_Reviews_LSTM.h5')

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T15:24:36.711753Z", "iopub.execute_input": "2025-03-
22T15:24:36.712507Z", "iopub.status.idle": "2025-03-
22T15:24:50.853373Z", "shell.execute_reply.started": "2025-03-
22T15:24:36.712471Z", "shell.execute_reply": "2025-03-22T15:24:50.852667Z" }}
from sklearn.metrics import confusion_matrix
import seaborn as sns
import matplotlib.pyplot as plt

# Отримання передбачень
y_pred = model.predict(X_test)
y_pred = (y_pred > 0.5).astype(int)

# Побудова матриці помилок
cm = confusion_matrix(y_test, y_pred)
plt.figure(figsize=(6,5))
sns.heatmap(cm, annot=True, fmt='d', cmap='Blues', xticklabels=['Negative', 'Positive'],
yticklabels=['Negative', 'Positive'])
plt.xlabel('Predicted Label')
plt.ylabel('True Label')
plt.title('Confusion Matrix')
plt.show()

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T15:25:09.105974Z", "iopub.execute_input": "2025-03-
22T15:25:09.106732Z", "iopub.status.idle": "2025-03-
22T15:25:09.175217Z", "shell.execute_reply.started": "2025-03-
22T15:25:09.106695Z", "shell.execute_reply": "2025-03-22T15:25:09.174474Z" }}
from sklearn.metrics import classification_report

print(classification_report(y_test, y_pred, target_names=['Negative', 'Positive']))

# %% [code]

# %% [code] {"execution": {"iopub.status.busy": "2025-03-
22T15:19:38.821176Z", "iopub.execute_input": "2025-03-
22T15:19:38.821845Z", "iopub.status.idle": "2025-03-
22T15:19:38.886606Z", "shell.execute_reply.started": "2025-03-
22T15:19:38.821810Z", "shell.execute_reply": "2025-03-22T15:19:38.885881Z" }}

```

```

# Тестування моделі на 30 екземплярах
sample_test_indices = np.random.choice(len(X_test), 30, replace=False)
X_sample_test_texts = test.iloc[sample_test_indices]['review_text'].values
y_sample_test = y_test[sample_test_indices]

X_sample_test = X_test[sample_test_indices]
predictions = model.predict(X_sample_test)
predicted_labels = (predictions > 0.5).astype(int)

for i in range(30):
    print(f'Відгук {i+1}: {X_sample_test_texts[i]}')
    print(f'Справжній клас: {y_sample_test[i]}, Передбачений клас: {predicted_labels[i][0]}\n')

accuracy = accuracy_score(y_sample_test, predicted_labels)
print(f'Точність моделі на 30 екземплярах: {accuracy * 100:.2f}%')

```

Додаток Б

Приклади реалізації

Відгук 7: Found this place last night on bringfido.com, I was in the area and had my dog with me so needed a place that was pet friendly. The ONLY good thing about this place was the large pet friendly patio area.\nWe sat for a long time before even being acknowledged. In the dozen times a server walked by, not a single \"be right with you\"\nNo drink refills\nAfter asking for a box, she began to place other customers dirty plates on top of my boyfriends plate making a comment like \"if you say you didn't like it I can take it off your bill.\" I had to stop her from laying dirty dishes on our food and explained that we would be taking the food home if she brought a box. \n\nI could go on and on but will sum it up like this: \n\nTerrible service, possibly worst ever.\n\nThe food was so-so.\n\nWon't be back.

Справжній клас: 0, Передбачений клас: 0

Відгук 8: Just Okay.\n\nThe place itself is nice inside. The staff were not all that interested in helping us. The prices were in line with what we expected and the food was fine.\n\nWhat hurts Tusca is that other, better Tapas places are nearby. Also the clientele at Tusca seemed to be mainly female late 30's to 40's who were there to hunt for younger upwardly mobile young men. Maybe this was just the few times I was there.

Справжній клас: 0, Передбачений клас: 0

Відгук 9: Food is good for the price. The place is pretty big with lots of big family table. The lunch menu starts off at 5 dollars ! Very good service for Chinatown!!!! I ordered preserved vegetables with intestines and ma po tofu. Both very authentic! I always take my family here!

Справжній клас: 1, Передбачений клас: 1

Відгук 10: I paid \$13 for a less than six inch sandwich with a bit of cold lobster clumps and two tiny slices of avocado in the middle. The whole thing was super greasy and the lobster mixture tasted very creamy. Not worth the money. I'm very disappointed at this item and this truck.

Справжній клас: 0, Передбачений клас: 0

Відгук 11: What a gem. LOVE this place. I daydream about the Manchow soup. Everything I have tried is fantastic. I like that you choose the level of spice you want- from mild to extra spicy with pretty much every meal. It has a large and diverse menu and I have yet to try something I didn't like. The guy that runs it is very knowledgeable and great with suggestions. Once you try it- you'll be hooked forever.

Справжній клас: 1, Передбачений клас: 1

Відгук 12: You can't beat Argyll Place for fruit and veg shops. Walking up this small street, it's impossible to ignore the profusion of colourful goods arrayed on the pavement, and I swear you won't find cheaper produce anywhere else in the city.\n\nMy favourite of these shops is Nadia's, although until recently I didn't even know its name - I just knew it as \"the one nearest the top,\" or \"the one that has the discount bucket.\" It's not especially different from the other two or three, but I like it because I got to know the staff a little here, who are always impeccably polite, and because they always have such a good range of decent quality produce, from basics such as onions and potatoes to more exotic seasonal offerings such as fresh figs or mangoes.\n\nLike anywhere that sells loose produce, it's important to look over what you buy and pick out the best/ripest/firmer of whatever it is you're buying, but in general nothing here is in dodgy condition, and, if you live in the area, it's a great place to pick up ingredients for dinner on your way home.

Справжній клас: 1, Передбачений клас: 1

Відгук 13: Was just in there, solo, to chill out, eat some pizza, drink some wine and hang out. \n\nRandom drunk dude (at 8:30pm) approaches me and asks me if I want a shot of Jack (n o). Proceeds to harass me about my haircut and ask me if I know his homeless buddy who, also, is from Seattle. \n\nThen he asks me if I have a problem with him? \n\nPizza to go please!\n\nI called the bartender after I left and told him he had a problem on his hands with that guy and he seemed receptive but, honestly, he should have sniffed that problem out right away.\n\nShit shouldn't get that weird in a place like that. There were families in there and I'm sure they weren't looking for a guy like that around them. \n\nFull disclosure; I like to party and I do like to drink and be inappropriate... a lot. Just not in someplace like that. Ain't nobody got time for that when there are kids around.\n\nPizza was OK out of a box at home. Drank a beer and jumped in the pool.

Справжній клас: 0, Передбачений клас: 0

Відгук 14: Pros: cool building, nice view of the brewing equipment\n\nCons: beer is mediocre at best.. Tastes watered down, food is way overpriced, and the shrimp mousse appetizer thing is horrendous... I mean disgusting. The chefs creations remind me of the crap I see on Gordon Ramsays kitchen nightmares where the chef gets excited with weird flavor combinations that don't jive... Plus bizarre glasses and food holders that are just annoying. \n\nI will not return. There's no value here.

Справжній клас: 0, Передбачений клас: 0

Відгук 15: I came to the egg and I in Las Vegas because I had been previously to the egg and I in Estes Park, Co. They didn't have the same menu items as they do in CO but everything was just as good. I ordered a breakfast skillet with potatoes and chili. The skillet was very good, I loved it. My wife had the breakfast burrito and she also loved it. The potatoes are seasoned with paprika and other ingredients and are very good, but I prefer the CO potatoes. The only thing that I didn't like about this place is that they don't use the egg and I brand hot sauce.

Справжній клас: 1, Передбачений клас: 1

Відгук 16: I went to check it out after my workout at LVAC store was extremely hot and uncomfortable inside. Poor customer service did not seem interested in helping me. Selection poor and price high . Definitely not worth checking out !!!

Справжній клас: 0, Передбачений клас: 0

Відгук 17: Food is horrible. Nobody in my party enjoyed their sandwich. The floor was greasy. The menu was wet and sticky.

Справжній клас: 0, Передбачений клас: 0

Відгук 18: Have been dying to try this kosher place for some time.\nWent tonight and ordered a large 3 cheese pizza. O M G american cheese???\nI asked also for baked ziti and the lady said her ziti can't keep for a week so they don't make it.\nAre you kidding me?\nIt was 19.00 with a drink.\nI put the pizza in the trunk and went to La Bella

Справжній клас: 0, Передбачений клас: 1

Відгук 19: My LG Front Load washer got clogged up when my house cleaner washed the bath rugs. Never a problem in the past, however this time the backing disintegrated inside the machine - for everyone out there, this is a major DON'T - and the rubber backing clogged everything and I mean EVERYTHING up. The entire washer had to be taken apart and carefully washed out, hoses cleared and then put back together.\n\nLeonard is friendly and knowledgeable. He also replaced the ignitor mechanism on my gas stove, and now everything is back to normal and working just fine. No hidden charges, no mess and minimal fuss.\n\nAbsolutely exceptional customer service and Leonard arrived exactly when he said he would. I will definitely be using their services in future (hopefully, not too soon). Always love to support small business owners - and this is one of the great ones. Many thanks.

Справжній клас: 1, Передбачений клас: 1

Бідрук 20: This was the worst breakfast place I have ever been to. I went there with my boyfriend. They asked for tip up front. My boyfriend ordered green tea. They were out of green tea, so the associate offered a different tea instead of making more green tea. I ordered coffee and was not even pointed to the coffee bar to get cream. Our food took forever to be delivered to our table. My soup was cold! Avoid this place like the plague. The restaurant looks like a gutter and the service was horrible.

Справжній клас: 0, Передбачений клас: 0

Бідрук 21: Does any apartment renter *actually* like their property management company? I think it's pretty rare. We all have gripes. However, in 11 years of renting, I've dealt with 5 property management companies and 2 owner-landlords. The owner-landlords were by far the worse. Of the 5 property management companies, Urban Land Interests earns the title of the worst. It's not that they were so bad that I moved out because of them. It's just that their properties (I only lived in 1, but priced several) are overpriced and their management is mediocre at best. I don't know if the person (who I won't name) who is in charge of Bel Mora is just overworked and underpaid, or is a legitimate space-case, but dealing with them was often like pulling teeth. I had an agreement with them for my move-out inspection that I wouldn't have the apartment all cleared out, because I was leaving very early the next morning, so I needed the bed still. This information wasn't relayed to the rude and confused kid who came by to do the inspection (who I wasn't expecting - I thought the property manager was coming by). He didn't want to do the inspection because the apartment wasn't empty, but after explaining the situation, he agreed to walk through with me. He had to have been high if he thought I was going to let him do it without me seeing exactly what he noted. The maintenance crew was always really nice, and when a maintenance request was submitted, it was addressed within a day or two. So, they get an A. The property itself was very nice (albeit overpriced). For what I paid for a 2BR, 2BA on the 2nd floor in Madison, I paid for a 2BR, 2BA luxury apartment on the 15th floor in Milwaukee, overlooking Lake Michigan. Their handling of my security deposit was suspect. There was a carpet stain that happened a year before we moved out. When I called the property manager after it happened, they said they could just cut out the area of carpet that was ruined and replace it, and it would be around \$100. Or, I could leave it until I moved out and it would just be taken out of my security deposit. Well, I opted for the latter, and it ended up costing me \$350! Pretty ridiculous, especially for a carpet that was complete crap when I moved in 2.5 years prior. Just another way to milk me for more money. Overall, I don't suggest renting a property from Urban Land Interests. It's a pretty tough market in Madison, because most landlords take advantage of renters because, well, they can with all the students. I'd still suggest properties run by other management companies.

Справжній клас: 0, Передбачений клас: 0

Бідрук 21: Does any apartment renter *actually* like their property management company? I think it's pretty rare. We all have gripes. However, in 11 years of renting, I've dealt with 5 property management companies and 2 owner-landlords. The owner-landlords were by far the worse. Of the 5 property management companies, Urban Land Interests earns the title of the worst. It's not that they were so bad that I moved out because of them. It's just that their properties (I only lived in 1, but priced several) are overpriced and their management is mediocre at best. I don't know if the person (who I won't name) who is in charge of Bel Mora is just overworked and underpaid, or is a legitimate space-case, but dealing with them was often like pulling teeth. I had an agreement with them for my move-out inspection that I wouldn't have the apartment all cleared out, because I was leaving very early the next morning, so I needed the bed still. This information wasn't relayed to the rude and confused kid who came by to do the inspection (who I wasn't expecting - I thought the property manager was coming by). He didn't want to do the inspection because the apartment wasn't empty, but after explaining the situation, he agreed to walk through with me. He had to have been high if he thought I was going to let him do it without me seeing exactly what he noted. The maintenance crew was always really nice, and when a maintenance request was submitted, it was addressed within a day or two. So, they get an A. The property itself was very nice (albeit overpriced). For what I paid for a 2BR, 2BA on the 2nd floor in Madison, I paid for a 2BR, 2BA luxury apartment on the 15th floor in Milwaukee, overlooking Lake Michigan. Their handling of my security deposit was suspect. There was a carpet stain that happened a year before we moved out. When I called the property manager after it happened, they said they could just cut out the area of carpet that was ruined and replace it, and it would be around \$100. Or, I could leave it until I moved out and it would just be taken out of my security deposit. Well, I opted for the latter, and it ended up costing me \$350! Pretty ridiculous, especially for a carpet that was complete crap when I moved in 2.5 years prior. Just another way to milk me for more money. Overall, I don't suggest renting a property from Urban Land Interests. It's a pretty tough market in Madison, because most landlords take advantage of renters because, well, they can with all the students. I'd still suggest properties run by other management companies.

Справжній клас: 0, Передбачений клас: 0

Бідрук 22: *****TACO TUESDAY*****\n\nMy first ever Taco Tuesday was here. Went with a LOT of rowdy people and the place was seemingly cool with the huge loud crowd, and even was able to seat us (I'm talking like 50 people).\n\nI got a meal #27, don't remember exactly what it was but it included tacos and beans and was pretty epic-tasting. The meal combo seemed so perfect that all of my friends got the #27 and loved it just as much. However, my friend Greg did not get the #27 and was temporarily shunned for doing so. The chips are epic, as was the salsa and the interior decor. Also epic are their margaritas (and from what I remember, cheap too). There is a separate bar area you can go to get away from your crowd and order drinks. Bar is especially cool because there are pictures of el guapo and wooden ships to look at while you order/drink. Cabs were ready and willing to take us back to the strip after we were finished eating. I rode with chuckyhacks and it was fun.\n\nWOULD GO AGAIN, HAD A BLAST, 11/10 STARS.

Справжній клас: 1, Передбачений клас: 1

Бідрук 23: This is by far the best Greek I've had. Our family visits here often. My wife is vegetarian and loves the spanakopita. It is really big and she loves it! I get the gyro with fries. Both are so flavorful. The fries have homemade spices on top. This appears to be owned by a family. The mom is always in the kitchen cooking and the children help with seating and orders. We've done sit down and take out. Both are great!

Справжній клас: 1, Передбачений клас: 1

Відгук 24: Linda is AMAZING, not only is she an amazing person but she is also awesome at scar revisions. I have a dog bite scar from 1996 that went all the way to the bone that I am self-conscious about. I saw a crazy deal on Living Social for skin needling, WHAT??? I had no idea what that even was but after talking to Linda on the phone for 20 minutes I felt confident that she could help, at least, make the scar less noticeable. She never made promises that I felt like she could not keep, as in, she said it would be invisible or it would disappear, etc. We are talking about a deep, old scar. However, my scar looks AMAZING after 4 treatments. It is flat, smaller and the discoloration is on the way to being gone. I am beyond impressed with how it looks. I AM THRILLED! I get comments all the time from people that know me and say "Is there something different?", it's crazy!!! Now that we have covered the skin needling, let's talk about the other services Linda offers. Along with the skin needling that Linda performed, she treated my whole face to make sure my scar treatment was getting the best chance to work. Skin scraping and all over face treatment. My skin looks and feels amazing. I have not even been this happy with how my skin looks and feels. I have less wrinkles, less middle age acne break outs and the texture of my skin is amazing. Linda never over sells products or pushes things on you, so you never feel like you "have to buy what she uses on you" like you do at some salons you visit. She is one of the most genuine people you will ever meet and she will be honest about what she can and can not do to help you. I really appreciated that when we started working on my scar. I truly appreciated all her hard work and how my scar looks now!! Thank you Linda and Scottsdale Skin Needling!! Amazing place and amazing business!!

Справжній клас: 1, Передбачений клас: 1

Відгук 25: While you can grab a six pack at almost any bar in this lovely neighborhood I call home, I was really missing having a legit beer distributor within walking distance. This place not only supplied me with Dogfish Head Punkin' Ale AND Tru00f6egs Mad Elf, but a usually great array of delicious variety beers and old standards for when I'm not looking to spend a pretty penny. They offer a walk in with a lot of awesome beers already chilled too. The best part though? At the register the owner has a huge variety of trading cards. I'm not talking brand new baseball cards, this is Yo! MTV Raps, 80's baseball and hockey cards, and Michael Keaton era Batman cards. It's a fun surprise when you're walking out with your beer purchase.

Справжній клас: 1, Передбачений клас: 1

Відгук 26: Although the exercise and sweating is pretty awesome, there are definite issues that are hard to ignore. First, there is nothing to help new people understand what they should wear or bring with them. Second, the instructions for what to have in the class are non-existent. Third, there are so many students in the class and the instructor is all over the place so new people haven't the slightest clue of what to do. The only way I was able to do at least something was to watch the person next to me. The instructor provided NO instruction that made sense for beginners and conducted class like everyone was a veteran. There needs to be separate beginner classes for sure along with some material or something on their web site that explains what you need and should bring. The instructor should have smaller classes and be cognizant of what special requirements the students need (such as carpal tunnel and alternate positions). Very disappointed with the setup and lack of instruction. Also, after 7pm, there is no one manning the front desk and they lock the front door so you'll have to unlock and leave yourself if you happen to leave early. If you do happen to go, here are some tips: 1) show up 15-20 minutes early your first time for paperwork 2) bring a yoga mat, yoga mat towel, sweat towel and bottle of water 3) highly suggest sports bra of sorts for the women and only shorts (no shirt) for the men 3) wear flip flops or slip-on/off shoes 4) leave purse/valuables/stuff in the car or at home 5) if using a card for payment, get some kind of small wallet that you can clip to your yoga mat or wrist 6) inside the classroom there are no shoes allowed and little or no talking 7) the room is a long rectangle and if a lot of folks are there, be prepared to be close to your neighbor.

Справжній клас: 0, Передбачений клас: 0

Відгук 27: Best piece of duck I've ever had. The menu is a spectacle of high-end tapas. Fabulous selection of exquisite and flavorful delicacies. The generous staff begins and ends each meal with tasty surprise.

Справжній клас: 1, Передбачений клас: 1

Відгук 28: In short, slow service, average food and a poor attitude by the wait staff add up to a disappointing meal. [BTY - Yelp maps this place 6+ miles from where it actually is located.]

Справжній клас: 0, Передбачений клас: 0

Відгук 29: I was so disappointed after going to Pure.\n\nMuch like Melissa R, this was my first Vegas club experience. I went with 3 of my bestest homies to Pure expecting to be on the guest list, only to find out we weren't and waited in line for two f-ck-n hours.... I'm not really sure if you can call it a line at all! It was a mob-- a HUGE MOB of people just waiting to get in. sucked.\n\nFinally, when we got in, the four of us went to the bar immediately. We had pre-drunk before we got to the line and BOY was I SOBER after all that waiting... SOBER AND THIRSTY! Two Patron shots with Pineapple backs and a midori sour = \$70 -- LOVELY!-- is that NORMAL?! I guess it was, no one else in the group complained about their tab..\n\nSo after taking my shot, I was ready to party.... (not really, but I was trying). Me and my bff TRIED going into the dance floor to have some fun... TRIED and FAILED--- =(The dance floor was SO PACKED, SO-FRIGGIN-BACK TO BACK-PACKED, I don't know how it was possible for ANYONE to dance there....\n\nOur group went upstairs to the opened-top... that was alright. The space was REALLY SMALL and security won't let you stand ANYWHERE! people were scrunched up in a SMALL CONFINED CORNER like cattles!!!.. people....which would include us!\n\nmy bff got sick, so after 2 hours of waiting to get into a disappointing club, the two of us left 45mins after entering Pure. On our way out thru the casino, there was a guy handing out free admission passes to PURE, we sed OH HELL NO THANKS.... While we were waiting for a cab, the group of girls in front of us also left Pure and was saying how horrible it was... yes beautiful girls.. i DO agree! There is NOTHING ANYONE can say or do to make me EVER go back there.... EVER.

Справжній клас: 0, Передбачений клас: 0

Відгук 30: Just had the worst pizza of my life. It was all sauce, barely any cheese and overall a big disappointment. Called the manager; she told me that Margherita pizzas are usually not recommended for delivery or pickup. Hmm..never heard that one before and I wasn't made aware of this when I placed the order. . She offered no refund or replacement. I said nothing, I just wanted them to know because it was like having bread and sauce. Any good restaurant would offer some sort of compensation. I could have asked for one, I didn't call for that but it says a lot about them. We've liked their food in the past so this was interesting.

Справжній клас: 0, Передбачений клас: 0

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
(ІТAR-2025)

27-29 травня 2025 року

Тернопіль
2025

Васенко Світозар, Биковий Павло	120
ВИКОРИСТАННЯ ПІДХОДУ BEHAVIOR TREE ДЛЯ СТВОРЕННЯ АДАПТИВНИХ НЕІГРОВИХ ПЕРСОНАЖІВ У НАВЧАЛЬНИХ ВІДЕОІГРАХ	120
Васильчук Олександр	123
ПРОГНОЗУВАННЯ СЕЗОННИХ ПРОДАЖІВ ПРОДУКЦІЇ В РОЗДРІБНІЙ ТОРГІВЛІ З ВИКОРИСТАННЯМ МОДЕЛІ SARIMA	123
Вібла Ксенія, Ліп'яніна-Гончаренко Христина	125
ІНТЕЛЕКТУАЛЬНИЙ TELEGRAM-БОТ ДЛЯ ПЕРСОНАЛІЗОВАНОГО ХАРЧУВАННЯ І ТРЕНУВАНЬ З ІНТЕГРАЦІЄЮ МОДЕЛІ LLAMA	125
Вітрук Іван, Лендюк Тарас	129
МОДУЛЬ ВИЗНАЧЕННЯ ПОЛЯРНОСТІ ВІДГУКІВ НА ПЛАТФОРМІ YELP ЗА ДОПОМОГОЮ LSTM-МЕРЕЖІ	129
Воєвудський Олександр, Ліп'яніна-Гончаренко Христина	132
TELEGRAM-БОТ ДЛЯ СТВОРЕННЯ ТА УПРАВЛІННЯ НАГАДУВАННЯМИ ПОВСЯКДЕННИХ ЗАВДАНЬ	132
Вороновський Володимир	135
ПРОГРАМНИЙ МОДУЛЬ ПОШУКУ СПРАВ З АРХІВІВ УКРАЇНИ, ДОСТУПНИХ ОНЛАЙН	135
Герцій Володимир, Турченко Ірина	138
ПРИКЛАДНИЙ ПРОГРАМНИЙ ІНТЕРФЕЙС ДЛЯ ТРАНСКРИПЦІЇ ТА СУБТИТРІВ ДЛЯ АУДІО ТА ВІДЕОФАЙЛІВ НА ОСНОВІ ШТУЧНОГО ІНТЕЛЕКТУ	138
Гінзула Володимир, Загородня Діана	142
ГОЛОСОВИЙ ПОМІЧНИК НА ОСНОВІ МАШИННОГО НАВЧАННЯ	142
Головінський Дмитро, Биковий Павло	144
ПРОГРАМНИЙ МОДУЛЬ CRM-СИСТЕМИ ДЛЯ ОПТИМІЗАЦІЇ ВИБОРУ ПОСТАЧАЛЬНИКІВ І УПРАВЛІННЯ ЗАМОВЛЕННЯМИ ОНЛАЙН-МАГАЗИНУ	144
Грушицький Максим, Турченко Ірина	147
ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ГЕНЕРУВАННЯ SQL-ЗАПИТІВ З ПРИРОДНОЇ МОВИ	147
Даниленко Денис	150
ПРОГРАМНИЙ МОДУЛЬ ЧАТ-БОТА ДЛЯ ОПТИМІЗАЦІЇ РОБОТИ З ДОКУМЕНТАЦІЄЮ	150
Зборовська Анастасія	152
РОЗРОБКА МОБІЛЬНОГО ЗАСТОСУНКУ АНАЛІТИКИ ДЛЯ РЕКОМЕНДАЦІЙНОЇ СИСТЕМИ ЧИТАННЯ КНИГ	152

Вітрук Іван
студент групи КН-42
vitruck.2017@gmail.com

Лендюк Тарас
к.т.н., доцент
tl@wunu.edu.ua

Західноукраїнський національний університет
Тернопіль, Україна

МОДУЛЬ ВИЗНАЧЕННЯ ПОЛЯРНOSTІ ВІДГУКІВ НА ПЛАТФОРМІ YELP ЗА ДОПОМОГОЮ LSTM-МЕРЕЖІ

У сучасному світі розпізнавання статі на основі зображень є важливою задачею в таких галузях, як маркетинг, безпека, персоналізація сервісів та біометрична ідентифікація. Метою дослідження стало створення ефективного модуля класифікації статі на основі глибокого навчання із використанням архітектури InceptionV3. В основі модуля – навчання моделі на великому відкритому датасеті CelebA (202 599 зображень), а також оптимізація архітектури і гіперпараметрів.

Розроблений модуль реалізовано на базі бібліотек TensorFlow і Keras. Він включає попередню обробку зображень (аугментація, балансування), розділення даних (20k train / 5k val / 5k test), модифікацію стандартної InceptionV3 (додавання GAP, Dense, Dropout), навчання з використанням SGD ($\text{lr}=0.0001$) і регулярного моніторингу метрик. Максимальна точність моделі на перевірочних даних становила 95,75%, а на тестовій – 94,3%. Значення F1-score = 0,94, що свідчить про збалансовану класифікацію обох класів.

На рисунку 1 зображено діаграму класів (UML) реалізованого модуля, яка ілюструє взаємодію основних компонентів системи, зокрема DataManager, TrainingController, EvaluationManager та InceptionV3Model.

У таблиці 1 здійснено порівняння з іншими сучасними підходами глибокого навчання до задачі гендерної класифікації. Найвищу точність досягнуто WideResNet (97.2%) і ResNet152 (95.8%), однак модель InceptionV3 забезпечила оптимальний баланс між точністю, швидкодією та кількістю параметрів, необхідних для навчання.

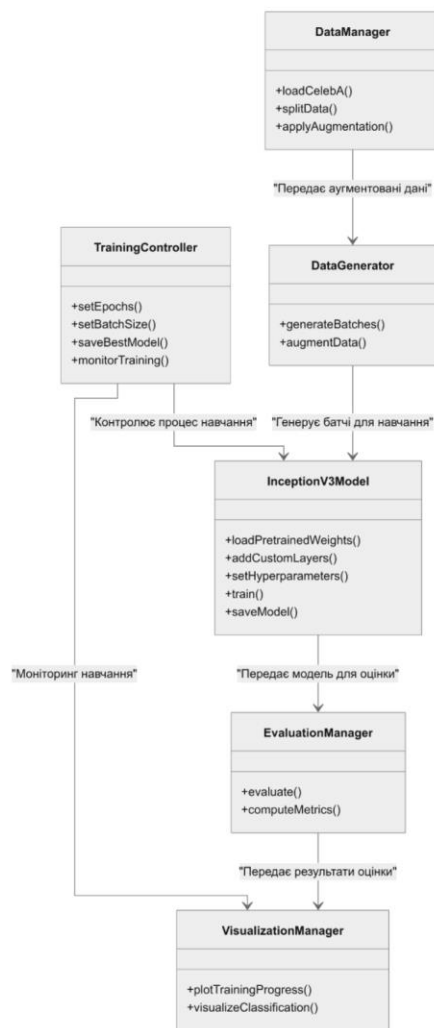


Рисунок 1 - UML-діаграма архітектури реалізації модуля розпізнавання статі

Таблиця 1 – Порівняння точності моделей глибокого навчання для класифікації статі

№	Модель	Набір даних	Точність (%)	Джерело
1	WideResNet	CelebA	97.2	Polina & Malathi [5]
2	ResNet152	CelebA	95.8	Wong et al. [10]
3	InceptionV3	UTKFace	93.2	Mohasseb et al. [1]
4	Пропонована	CelebA	94.3	Ця робота

Завдяки високій точності та стабільності, розроблений модуль є придатним для інтеграції у практичні системи відеоаналітики, кібербезпеки, адаптивного UX-дизайну та цифрового маркетингу. Запропонована архітектура легко масштабована й оптимізована, що дозволяє її адаптувати для інших задач класифікації обличчя.

Список використаних джерел

1. Mohasseb, A., Amer, E., Chiroma, F., & Tranchese, A. (2025). Leveraging Advanced NLP Techniques and Data Augmentation to Enhance Online Misogyny Detection. *Applied Sciences*. DOI: 10.3390/app15020856.
2. Kwiatek, J., Leśna, M., Piskórz, W., et al. (2025). Comparison of the Diagnostic Accuracy of an AI-Based System for Dental Caries Detection and Clinical Evaluation Conducted by Dentists. *Journal of Clinical Medicine*. DOI: 10.3390/jcm14051566.
3. Ismaeel, A. Z., & Yasin, H. M. (2025). Machine Learning Based Finger Print Analysis for Gender Detection: A Review. *Asian Journal of Research in Computer Science*. Без DOI, доступно тут.
4. Khaliluzzaman, M., Hoque, M. D. J., et al. (2024). Revolutionizing Age and Gender Recognition: An Enhanced CNN Architecture. *International Journal of Online and Social Interaction*. DOI: 10.1109/IJOSI.2024.1133.
5. Polina, V. K., & Malathi, K. (2025). Classification of Target Customers for Online Advertising Using Wide ResNet CNN and Comparing Its Accuracy Over ELM-CNN Algorithm. *Research in Management, Accounting and Economics*. DOI: 10.4324/9781003606642-96.
6. Rezasoltani, A., Hosseini, A., & Toosi, R. (2024). A Multi-Task Framework Using Mamba for Identity, Age, and Gender Classification from Hand Images. *IEEE Transactions on Information and Communication Technologies*. DOI: 10.1109/TICT.2024.10892788.
7. Javed, N., Ahmed, T., & Faisal, M. (2024). An Efficacy Comparison of Supervised Machine Learning Classifiers for Cyberbullying Detection and Prediction. *International Journal of Bullying Prevention*. DOI: 10.1007/s42380-024-00282-1.
8. Ali, I., Maharani, A. A. S., & Suarna, N. (2025). Application of Decision Tree Algorithm to Improve Student Learning Pattern Classification Model at Sempoja Sip Perjuangan. *Journal of Artificial Intelligence and Education Applications*. DOI: 10.1016/j.jaiea.2025.12163.
9. Shekofteh, Y., & Moradi, A. (2024). APEDM: A New Voice Casting System Using Acoustic–Phonetic Encoder-Decoder Mapping. *Multimedia Tools and Applications*. DOI: 10.1007/s11042-024-20496-1.
10. Wong, J., Lu, S., Lou, Y., et al. (2024). A Novel Time Use Approach on Successful Aging: Racial and Gender Disparities in Daily Productive Engagement. *Innovation in Aging*. DOI: 10.1093/geroni/igb098.2201.