

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

БОГУТА Геннадій Миколайович

Модуль аналізу тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками/ Module for text sentiment analysis to determine emotional coloring of publications by thematic categories

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

Кваліфікаційна робота

Виконав студент групи КН-42
Г. М. Богута

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

Кваліфікаційну роботу допущено до
захисту

«___» _____ 2025 р.

В.о. завідувача кафедри
_____ Н.М. Васильків

Тернопіль – 2025

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

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

ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ
БОГУТА Геннадій Миколайович
(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи: Модуль аналізу тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками/Module for text sentiment analysis to determine emotional coloring of publications by thematic categories
керівник роботи к.т.н., доцент Х.В. Ліп'яніна-Гончаренко
(прізвище, ім'я, по батькові, науковий ступінь, вчене звання)

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

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

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

4. Основні питання, які потрібно розробити:

– провести аналіз предметної області, зокрема впливу тональності тексту на сприйняття інформації.

– дослідити сучасні підходи до аналізу тональності текстів та методи їх реалізації.

– розробити алгоритми попередньої обробки та перекладу текстів публікацій.

– реалізувати класифікацію тексту за тональністю з використанням моделей машинного навчання.

– побудувати модуль інтерпретації результатів за тематичними рубриками та розрахунком індексу тональності.

– створити веб-інтерфейс для введення, аналізу та редагування текстів і словників.

5. Перелік графічного матеріалу в роботі:

– схема алгоритму збору та підготовки текстових даних;

– схема алгоритму перекладу тексту публікацій

– схема алгоритму класифікації тексту за тональністю

– схема методу тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками

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 _____ Г. М. Богута
(підпис) (прізвище та ініціали)

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

АНОТАЦІЯ

Кваліфікаційна робота на тему «Модуль аналізу тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками» обсягом в 67 сторінок і містить 12 ілюстрацій, 2 таблиці, 2 додатки та 34 використаних джерел.

Метою роботи є розробка та реалізація інтелектуального модуля для автоматизованого аналізу тональності публікацій, що охоплює багатомовний переклад, тематичну класифікацію, класифікацію за тональністю, інверсію результатів та візуалізацію даних.

Методами дослідження обрано методи інтелектуального аналізу текстових даних, зокрема попередню обробку, машинний переклад, машинне навчання (Naïve Bayes, SVM, Random Forest), глибоке навчання (CNN), а також лексичні словникові підходи. Застосовано методи класифікації, метричної оцінки якості (accuracy, precision, recall, F1-score) та розроблено web-інтерфейс для інтерактивної взаємодії з користувачем.

У результаті виконання роботи створено програмний засіб, здатний аналізувати тональність новинних публікацій, враховуючи тематичний контекст і походження джерела. Проведено тестування модуля на реальних прикладах медіа, що підтвердило ефективність моделі з точністю класифікації досягти точності до 84% на тренувальній вибірці та 78% на валідаційній у складних умовах інформаційного середовища.

Ключові слова: АНАЛІЗ ТОНАЛЬНОСТІ, НЕЙРОННА МЕРЕЖА, МАШИННЕ НАВЧАННЯ, ЕМОЦІЙНЕ ЗАБАРВЛЕННЯ, ГІБРИДНА ВІЙНА.

ANNOTATION

The qualification work on the topic “Module for text sentiment analysis to determine emotional coloring of publications by thematic categories” is 67 pages long and contains 12 illustrations, 2 tables, 2 appendixes, and 34 sources used.

The purpose of the work is to develop and implement an intelligent module for automated analysis of the tone of publications, covering multilingual translation, thematic classification, classification by tone, inversion of results, and data visualization.

The research methods chosen are methods of intelligent analysis of text data, in particular pre-processing, machine translation, machine learning (Naive Bayes, SVM, Random Forest), deep learning (CNN), as well as lexical dictionary approaches. Classification methods and metric quality assessment (accuracy, precision, recall, F1-score) were applied, and a web interface for interactive user interaction was developed.

As a result of the work, a software tool was created that is capable of analyzing the tone of news publications, taking into account the thematic context and origin of the source. The module was tested on real media examples, which confirmed the effectiveness of the model with a classification accuracy of up to 84% on the training sample and 78% on the validation sample in complex information environment conditions.

Keywords: TONE ANALYSIS, NEURAL NETWORK, MACHINE LEARNING, EMOTIONAL TINT, HYBRID WARFARE.

ЗМІСТ

Вступ.....	7
1 Аналіз предметної області та постановка задачі	10
1.1 Вплив тональності тексту на емоційне забарвлення публікацій	10
1.2 Огляд і аналіз існуючих підходів до визначення емоційного забарвлення тексту	11
1.3 Методи машинного навчання для класифікації тексту	13
1.4 Постановка задачі	16
2 Алгоритмічне забезпечення модуля	19
2.1 Алгоритм збору та підготовки текстових даних.....	19
2.2 Алгоритм перекладу тексту публікацій	22
2.3 Алгоритм класифікації тексту за тональністю	25
2.4 Метрики оцінки точності та оптимізації моделі	28
2.5 Метод тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками	30
3 Програмна реалізація модуля	35
3.1 Підготовка та обробка даних	35
3.2 Розробка web-інтерфейсу для аналізу тональності	37
3.3 Тестування ефективності методу	40
Висновки.....	43
Список використаних джерел.....	45
Додаток А Отримані результати.....	49
Додаток Б Апробації	50

ВСТУП

У сучасному інформаційному просторі емоційний вплив текстів відіграє ключову роль у формуванні громадської думки, особливо в умовах гібридних конфліктів. Новинні матеріали, пости у соціальних мережах і коментарі на форумах часто містять емоційно забарвлений контент, що впливає на психологічний стан аудиторії та її поведінкові реакції. У зв'язку з цим виникає потреба у створенні ефективних інструментів для автоматизованого визначення тональності текстів.

Особливої актуальності набуває задача аналізу тональності у контексті війни Росії проти України, де інформаційні атаки є невід'ємним інструментом впливу. Виявлення негативно або маніпулятивно забарвлених повідомлень дозволяє своєчасно реагувати на дезінформаційні кампанії та зміцнювати інформаційну безпеку. Саме тому розробка засобів інтелектуального аналізу емоційного змісту публікацій є надзвичайно важливою.

Проблема ускладнюється багатомовністю сучасного медіаконтенту, що потребує попереднього перекладу текстів для подальшої уніфікованої обробки. Також важливо враховувати тематичний контекст публікацій, адже одна й та сама лексика може мати різне емоційне навантаження залежно від сфери вживання. Це зумовлює необхідність побудови модульного рішення, що поєднує переклад, тематичну класифікацію та словниковий аналіз.

Таким чином, актуальність дослідження обумовлена потребою у побудові комплексної системи, здатної аналізувати емоційне забарвлення текстів з різних джерел, різними мовами та з урахуванням контексту, для використання у сфері медіааналітики, інформаційної безпеки та суспільних комунікацій.

Метою роботи є розробка та реалізація інтелектуального модуля для автоматизованого аналізу тональності публікацій за тематичними рубриками, який поєднує переклад, обробку текстів, машинне навчання та візуалізацію результатів.

Завдання дослідження:

1. Провести аналіз предметної області, зокрема впливу тональності тексту на сприйняття інформації.

2. Дослідити сучасні підходи до аналізу тональності текстів та методи їх реалізації.

3. Розробити алгоритми попередньої обробки та перекладу текстів публікацій.

4. Реалізувати класифікацію тексту за тональністю з використанням моделей машинного навчання.

5. Побудувати модуль інтерпретації результатів за тематичними рубриками та розрахунком індексу тональності.

6. Створити веб-інтерфейс для введення, аналізу та редагування текстів і словників.

Предмет дослідження - процес автоматизованого визначення тональності тексту публікацій за допомогою методів машинного навчання.

Об'єкт дослідження - текстові дані новинних публікацій, зібрані з відкритих джерел, що містять емоційно забарвлену інформацію.

Методи дослідження. У роботі використано комплекс методів інтелектуального аналізу даних, зокрема попередню обробку текстів, машинний переклад, тематичну класифікацію та аналіз тональності. Застосовано алгоритми машинного навчання (Naive Bayes, SVM, Random Forest) та глибокого навчання (CNN), а також словникові підходи до визначення емоційного забарвлення. Для реалізації системи використано веб-інтерфейс, який забезпечує інтерактивність, гнучкість налаштувань і візуалізацію результатів аналізу.

Результати дослідження будуть використані в межах проєкту держбюджетної наукової роботи молодих учених (номер держреєстрації — в процесі реєстрації) «Нейромережева система моніторингу загроз національній безпеці України на основі тональності та класифікації текстового контенту» (термін виконання: 01.01.2025 р. – 31.12.2027 р.).

Апробація результатів дослідження. Основні теоретичні положення та практичні результати дослідження були оприлюднені на низці міжнародних конференцій. Зокрема, результати методів визначення тональності текстів за тематичними рубриками опубліковано та проіндексовано в базі Scopus:

1. Sachenko, A., Lendiuk, T., Lipianina-Honcharenko, K., Dobrowolski, M., Boguta, G., & Bytsyura, L. (2024). *Method of Determining the Text Sentiment by Thematic Rubrics*. In COLINS (3) (pp. 404–414). *CEUR Workshop Proceedings, Volume 3688*, 8th International Conference on Computational Linguistics and Intelligent Systems, Lviv, 12–13 April 2024.
2. Lipianina-Honcharenko, K., Lendiuk, D., Ivaniush, A., Boguta, G., Soia, M., & Yurkiv, K. (2024, October). *An Intelligent Method for Recognizing Real and Generated Ukrainian-Language Voices*. In 2024 IEEE 5th KhPI Week on Advanced Technology (KhPIWeek) (pp. 1–6), Kharkiv, 7–11 October 2024.
3. Lipianina-Honcharenko K., Ihnatiev I., Boguta G., Drakokhrust T. & Telka M. (2025). *Method for Determining the Sentiment of Foreign News about Ukraine*. In 2nd International Workshop of Young Scientists on Artificial Intelligence for Sustainable Development, AISD 2025, Ternopil, 8–9 May 2025.

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ПОСТАНОВКА ЗАДАЧІ

1.1 Вплив тональності тексту на емоційне забарвлення публікацій

Тональність тексту є одним із ключових факторів, що визначають його емоційне забарвлення та вплив на читача. Вона відображає суб'єктивне ставлення автора до змісту публікації та може істотно впливати на сприйняття інформації, формування громадської думки, а також на поведінкові реакції аудиторії. У цифрову епоху, коли великий обсяг текстової інформації поширюється через соціальні мережі, новинні портали та інші онлайн-ресурси, аналіз тональності стає важливим інструментом для автоматизованої обробки та інтерпретації текстів.

Тональність тексту визначається як емоційне забарвлення, що міститься у висловлюваннях. Вона може бути позитивною, негативною або нейтральною. Деякі наукові підходи виділяють також змішані та нейтрально-забарвлені категорії. Позитивна тональність вказує на підтримку або схвалення, негативна – на критику або невдоволення, а нейтральна – на відсутність яскраво виражених емоційних оцінок.

Емоційне забарвлення тексту формується через використання певних лексичних конструкцій, стилістичних прийомів, а також контекстуальних відтінків, які можуть підсилювати або пом'якшувати емоційний вплив. Наприклад, у публікаціях новинного характеру емоційне забарвлення тексту може формувати ставлення аудиторії до соціально значущих подій, що, у свою чергу, впливає на суспільну думку та поведінкові реакції.

Дослідження у галузі когнітивної психології та комунікативістики свідчать про те, що тональність тексту має безпосередній вплив на емоційний стан читача, його думки та подальші дії. Позитивно забарвлені повідомлення сприяють формуванню довіри, підвищенню лояльності до джерела інформації та можуть стимулювати проактивну поведінку. Негативна тональність, навпаки, може викликати тривожність, агресію або недовіру, посилюючи емоційний резонанс та привертаючи більше уваги.

Зокрема, в умовах інформаційного перенасичення сучасного медіапростору, емоційно забарвлені тексти мають більшу ймовірність бути поміченими, поширеними та активно обговорюваними. Це підтверджують численні дослідження, які вказують на кореляцію між емоційним забарвленням контенту та його віральністю.

Тональність тексту може значно варіюватися залежно від типу публікації. Наприклад, у новинних статтях емоційне забарвлення часто приховане за формально-нейтральною мовою, тоді як у соціальних мережах текст може бути більш експресивним та відвертим. Аналіз тональності у текстах рекламного характеру спрямований на виявлення прихованого впливу на споживачів, тоді як у наукових статтях тональність переважно нейтральна, що зумовлено вимогами до об'єктивності.

Особливу складність становить аналіз тональності в текстах, що містять іронію, сарказм або багатозначність. Такі тексти можуть містити суперечливі емоційні сигнали, які складно інтерпретувати традиційними методами аналізу.

1.2 Огляд і аналіз існуючих підходів до визначення емоційного забарвлення тексту

У сучасному інформаційному просторі дезінформація та маніпуляції громадською думкою стають невід'ємною складовою гібридних війн, що суттєво впливає на політичні, соціальні та військові процеси. Особливу роль у таких інформаційних конфліктах відіграють технології обробки природної мови, машинного навчання [1-2] та аналізу тональності текстів, які дозволяють автоматизовано ідентифікувати та класифікувати дезінформаційні потоки.

Padalko et al. [3] у своєму дослідженні аналізують застосування глибокого навчання та моделей NLP, зокрема XLNet, для класифікації дезінформаційних повідомлень у гібридній війні. Дослідження демонструє використання Kullback-Leibler Divergence для оцінки змістовної різниці між новинними матеріалами та їх імовірним дезінформаційним характером. Запропоновані підходи спрямовані на

покращення автоматизованого аналізу інформаційних атак у контексті війни Росії проти України.

Virtosu & Goian [4] досліджують, як штучний інтелект використовується для створення та розповсюдження дезінформації, що є частиною інформаційних атак Росії проти України. Автори розглядають конкретні методи NLP та генеративного ШІ, які застосовуються для маніпуляції громадською думкою через соціальні медіа та новинні платформи.

Rodrigues [5] аналізує емоційну поляризацію у Twitter щодо війни Росії проти України, використовуючи методи Sentiment Analysis. Дослідження охоплює період із серпня 2022 до лютого 2023 року та виявляє значні відмінності у тональності публікацій між проросійськими та проукраїнськими користувачами.

Darwish et al. [6] застосовують методи машинного навчання для виявлення фейкових новин, пов'язаних із конфліктом між Росією та Україною. Дослідники розглядають особливості стилю та змісту дезінформаційних повідомлень, а також пропонують алгоритмічні підходи для їх виявлення.

Moy & Gradon [7] досліджують вплив штучного інтелекту на інформаційні війни, включаючи методи обробки природної мови (NLP) для аналізу гібридних інформаційних атак. Вони наголошують на двозначності ШІ в інформаційній війні, зазначаючи, що його можна використовувати як для боротьби з дезінформацією, так і для її поширення.

Alieva, Kloo & Carley [8] використовують комбінацію мережевого аналізу та NLP для вивчення російської пропаганди в Twitter під час вторгнення Росії в Україну. Аналіз показує, як проросійські акаунти координують інформаційні кампанії та впливають на онлайн-дискурс.

Weigand [9] аналізує російські дезінформаційні кампанії у 2022 році, використовуючи методи аналізу тексту та обробки природної мови. Дослідження визначає головні теми, які просуваються через дезінформаційні канали, та їхній вплив на громадську думку.

Strubytskyi & Shakhovska [10] пропонують методи аналізу тональності текстів, які дозволяють виявляти приховану пропаганду у новинних статтях. Вони

використовують гібридний підхід, поєднуючи класичні методи аналізу тональності та сучасні моделі глибокого навчання.

Padalko, Chomko & Chumachenko [11] досліджують, як видалення стоп-слів впливає на точність виявлення дезінформації українською мовою. Вони доводять, що ретельна попередня обробка текстів може суттєво покращити ефективність NLP-моделей у боротьбі з дезінформацією.

Suman et al. [12] використовують гібридний підхід глибокого навчання для класифікації та аналізу постів у Twitter (X) щодо війни Росії проти України. Модель аналізує зміну тональності дописів протягом часу та виявляє основні тренди у суспільних настроях.

Таким чином, враховуючи масштаб та інтенсивність дезінформаційних атак, постає необхідність створення ефективного інструментарію для автоматичного аналізу текстового контенту. Метою даного дослідження є розробка інтегрованого підходу до аналізу тональності текстів, який об'єднує автоматизований переклад, попередню обробку, тематичну класифікацію та словниковий аналіз. Цей підхід має забезпечити оперативне та точне виявлення емоційного забарвлення інформаційних повідомлень у відсотковому вираженні від -100% до $+100\%$, що є надзвичайно актуальним для протидії дезінформаційним атакам у контексті гібридної війни Росії проти України на міжнародній арені.

1.3 Методи машинного навчання для класифікації тексту

Класифікація тексту є однією з основних задач у сфері обробки природної мови (Natural Language Processing, NLP). Вона полягає у віднесенні текстових даних до однієї або декількох попередньо визначених категорій, таких як позитивна, негативна або нейтральна тональність. З розвитком технологій машинного навчання з'явилося багато підходів до автоматизації цього процесу, що дозволяють досягти високої точності та ефективності під час аналізу великих обсягів текстової інформації.

У цьому підрозділі розглядаються основні методи машинного навчання, що застосовуються для класифікації текстів, включаючи статистичні підходи, класичні алгоритми машинного навчання та сучасні нейронні мережі. Особливу увагу приділено глибоким згортковим нейронним мережам (Convolutional Neural Networks, CNN), які демонструють високу ефективність у розв'язанні задач текстової класифікації.

Базові підходи до класифікації текстів спираються на статистичні методи, які використовують частотні характеристики термінів. Найпростішими з них є Bag-of-Words (BoW) та Term Frequency-Inverse Document Frequency (TF-IDF).

Модель "мішка слів" (BoW) — це метод перетворення тексту у вектор фіксованої довжини, де кожен елемент відповідає кількості входжень конкретного слова. Формально, текст d представлений як вектор:

$$x = (x_1, x_2, \dots, x_n), \quad (1.1)$$

де x_i — кількість входжень i -го слова з визначеного словника.

TF-IDF зважає терміни залежно від їх частоти у документі та зворотної частоти у всьому корпусі:

$$TF - IDF(t, d) = tf(t, d) \times \log\left(\frac{N}{df(t)}\right), \quad (1.2)$$

де $tf(t, d)$ — частота терміна t у документі d ;

N — загальна кількість документів у корпусі;

$df(t)$ — кількість документів, що містять термін t .

Базуючись на теоремі Байєса, наївний байєсівський класифікатор (Naive Bayes) обчислює ймовірність належності тексту до певного класу:

$$P(y|x) = \frac{P(x|y)P(y)}{P(x)}, \quad (1.3)$$

де $P(y|x)$ – ймовірність того, що документ належить класу y ;

$P(x|y)$ – ймовірність отримати ознаку x за умови класу y ;

$P(y)$ – апіорна ймовірність класу y ;

$P(x)$ – нормувальний коефіцієнт.

Переваги: швидкість, ефективність на малих вибірках.

Недоліки: припущення про незалежність ознак, яке часто не виконується у реальних текстах.

Метод опорних векторів (Support Vector Machine, SVM) шукає гіперплощину, яка максимізує відстань між класами. Нехай x – вхідний текст, тоді класифікація виконується за правилом:

$$f(x) = \text{sign}(w^T x + b), \quad (1.4)$$

де w – ваговий вектор;

b – зсув.

З розвитком глибокого навчання з'явилися більш потужні методи класифікації тексту, що враховують контекст та семантичні зв'язки між словами.

RNN обробляють текст як послідовність, де вихід поточного кроку залежить від попереднього стану. Основне рівняння виглядає так:

$$h_t = \sigma(W_h h_{t-1} + W_x x_t + b), \quad (1.5)$$

де h_t – прихований стан на кроці t ;

x_t – вхідний вектор;

W_h, W_x – ваги;

b – зсув.

Проблема RNN – затухання градієнтів при обробці довгих послідовностей.

Згорткові нейронні мережі, які спочатку були розроблені для обробки зображень, успішно застосовуються для текстової класифікації. Вони можуть виявляти локальні залежності у тексті завдяки згортковим фільтрам.

Нехай x – вхідний текст, тоді вихід згорткового шару:

$$h_i = f(w * x_{i:i+k-1} + b), \quad (1.6)$$

де f – нелінійна функція;

w – згортковий фільтр розміром k ;

b – зсув.

Згорткова мережа складається з трьох основних етапів:

1. Згортка (Convolution): Виявлення локальних ознак.
2. Пулінг (Pooling): Зменшення розмірності.

Класифікація: Вихідні ознаки подаються до повнозв'язного шару.

CNN демонструють кращу продуктивність, ніж SVM та Naive Bayes, при роботі з великими текстовими наборами завдяки здатності автоматично виявляти значущі патерни.

Порівняння підходів показує, що згорткові нейронні мережі мають низку переваг для задач класифікації текстів:

1. Висока ефективність у виявленні локальних патернів.
2. Краща масштабованість порівняно з RNN.
3. Швидкість навчання та інференсу.

З огляду на ці переваги, у подальшій реалізації системи буде використана архітектура CNN, оскільки вона оптимально поєднує точність класифікації, швидкість обробки та можливість розширення для роботи з великими текстовими корпусами.

1.4 Постановка задачі

У сучасному інформаційному суспільстві, де новини, коментарі та аналітичні матеріали поширюються з великою швидкістю, емоційне забарвлення тексту суттєво впливає на сприйняття інформації та формування суспільної думки. Зважаючи на масштаб гібридної війни, яка ведеться Росією проти України, аналіз

тональності публікацій набуває особливого значення як інструмент виявлення дезінформації, інформаційного впливу та пропаганди.

Публікації з емоційно забарвленим змістом здатні викликати сильні реакції у читачів, стимулювати поширення контенту, формувати або змінювати ставлення до певних подій, персон чи країн. Саме тому важливо мати інструменти, які дозволяють автоматично визначати тональність текстів у великому масиві даних, зокрема у медіа.

Окрему складність становить робота з багатомовними джерелами, адже значна частина інформації, що впливає на імідж України, публікується англійською або іншими мовами. У зв'язку з цим виникає потреба у попередньому автоматизованому перекладі тексту та його адаптації до подальшої обробки українською мовою.

Крім того, аналіз текстів має враховувати не лише загальну тональність, а й специфіку тематичних рубрик, до яких належать публікації. Це дозволяє глибше оцінити інформаційний вплив у конкретних сферах, таких як військове керівництво, імідж країни за кордоном, діяльність державних інституцій тощо.

Актуальність даного дослідження зумовлена потребою у побудові комплексного підходу, який поєднує сучасні технології машинного перекладу, методи машинного навчання, глибокі нейронні мережі, а також гнучкі словникові модулі, здатні адаптуватися до змін мовного середовища. Такий підхід може бути застосований у різних сферах — від журналістики до кібербезпеки.

У межах дослідження було обрано гібридний підхід до визначення емоційного забарвлення текстів, який поєднує переваги словникових методів, методів класичного машинного навчання та нейромережових архітектур. Такий вибір зумовлений потребою у гнучкому та масштабованому рішенні, здатному ефективно обробляти великі обсяги багатомовних текстових даних із врахуванням тематичного контексту.

Вибір нейронних мереж, зокрема CNN, пояснюється їх високою продуктивністю при роботі з великими корпусами текстів, а також здатністю автоматично виявляти локальні патерни. Додатково впроваджено алгоритм інверсії

тональності для контенту з ворожих джерел, що підвищує точність інтерпретації емоційного змісту.

Метою роботи є розробка та реалізація інтелектуального модуля для автоматизованого аналізу тональності публікацій за тематичними рубриками, який поєднує переклад, обробку текстів, машинне навчання та візуалізацію результатів.

Завдання дослідження:

1. Провести аналіз предметної області, зокрема впливу тональності тексту на сприйняття інформації.
2. Дослідити сучасні підходи до аналізу тональності текстів та методи їх реалізації.
3. Розробити алгоритми попередньої обробки та перекладу текстів публікацій.
4. Реалізувати класифікацію тексту за тональністю з використанням моделей машинного навчання.
5. Побудувати модуль інтерпретації результатів за тематичними рубриками та розрахунком індексу тональності.
6. Створити веб-інтерфейс для введення, аналізу та редагування текстів і словників.

2 АЛГОРИТМІЧНЕ ЗАБЕЗПЕЧЕННЯ МОДУЛЯ

2.1 Алгоритм збору та підготовки текстових даних

Ефективність аналізу тональності тексту значною мірою залежить від якості вхідних даних, що використовуються для навчання та тестування моделей. Тому на етапі збору й підготовки текстової інформації важливо забезпечити її структурованість, релевантність і чистоту. У цьому підрозділі розглянуто ключові джерела отримання текстових даних, а також основні етапи їхньої обробки: очищення від шумів, токенизація, фільтрація стоп-слів, лематизація та нормалізація. Наведена схема (рисунок 2.1) і формалізація кожного етапу демонструють послідовність перетворення сирих текстів у стандартизовану форму, придатну для подальшого аналізу.

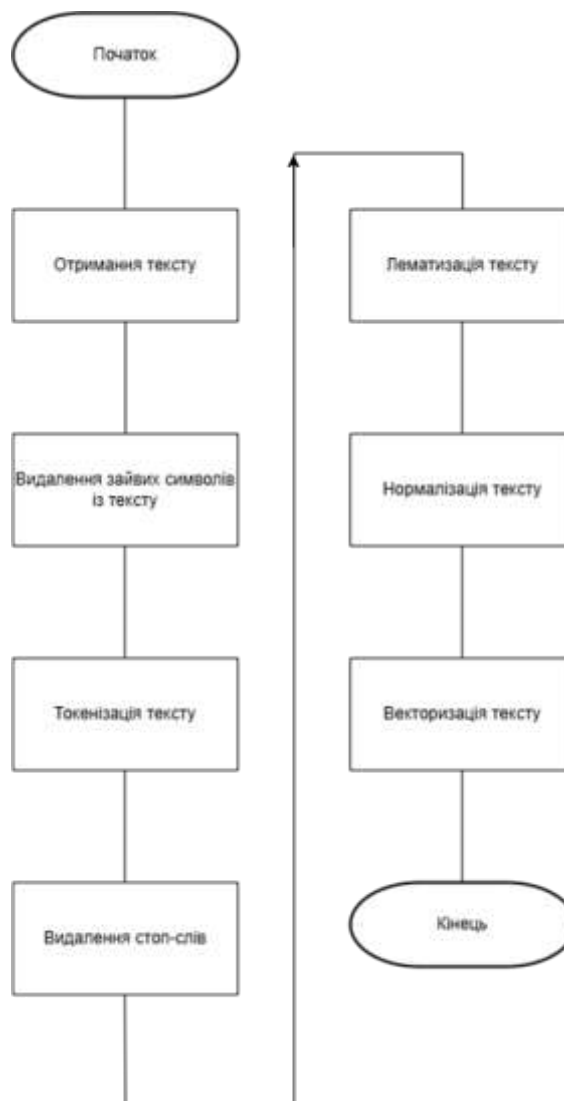


Рисунок 2.1 – Схема алгоритму збору та підготовки текстових даних

Збір і підготовка текстових даних є важливим етапом у процесі автоматизованої класифікації текстів за тональністю. Якість вихідних даних безпосередньо впливає на точність і продуктивність моделі машинного навчання. На цьому етапі відбувається збирання текстової інформації з різних джерел, її очищення від зайвих елементів, нормалізація та приведення до стандартизованого вигляду. Основна мета підготовки тексту – забезпечити точність аналізу, зменшити вплив шуму та полегшити подальше навчання моделі.

Алгоритм збору та підготовки текстових даних включає кілька послідовних етапів: видалення зайвих символів, токенізація, усунення стоп-слів, лематизація та нормалізація тексту. Кожен з етапів відіграє важливу роль у формуванні якісного вхідного сигналу для моделі класифікації.

Збір текстових даних здійснюється із зовнішніх джерел, які містять інформацію, необхідну для аналізу тональності. Основні джерела включають:

- Соціальні мережі – для отримання неструктурованих текстових даних, відгуків і коментарів.
- Новинні портали – для збору текстів із формальним стилем, що містять різні тональні відтінки.
- Форуми та блоги – для аналізу текстів із неформальним стилем, які часто включають емоційне забарвлення.
- Публічні текстові корпуси – наприклад, Sentiment140, IMDB, SST-2, що містять великі обсяги анотованих даних.

Кожен зібраний текст зберігається у стандартизованому форматі для подальшої обробки.

Зібрані тексти часто містять зайву інформацію, яка не має значення для аналізу тональності. Тому важливо здійснити їхнє очищення для зменшення шуму та підвищення ефективності роботи алгоритму.

Основні етапи видалення непотрібного контенту:

1. Видалення HTML-тегів – усунення розмітки з веб-сторінок
2. Фільтрація спеціальних символів – виключення символів, які не впливають на зміст.

3. Видалення гіперпосилань та електронних адрес – видалення URL-адрес та e-mail-адрес.

4. Фільтрація дубльованих текстів – виключення ідентичних або дуже схожих текстів для уникнення повторюваної інформації.

Формально цей етап можна описати так:

$$T' = \phi(T), \quad (2.1)$$

де T – початковий текст,

ϕ – функція очищення,

T' – очищений текст.

Токенізація – це процес розбиття тексту на окремі складові (слова, фрази або речення), які називаються токенами. Вона є обов'язковим етапом, оскільки моделі машинного навчання працюють саме з токенізованими текстами.

Формально токенізацію можна описати як функцію:

$$\tau: T \rightarrow \{w_1, w_2, \dots, w_n\}, \quad (2.2)$$

де T – вхідний текст,

τ – функція токенізації,

w_i – отримані токени.

Стоп-слова – це слова, які не несуть смислового навантаження. Вилучення таких слів дозволяє зменшити розмір тексту та підвищити якість розпізнавання ключових елементів.

Нехай S – множина стоп-слів, тоді видалення описується як:

$$T' = \{w \in T: w \notin S\}, \quad (2.3)$$

де T – вихідний текст,

T' – текст без стоп-слів.

Лематизація та стемінг – це процеси приведення слів до їхньої базової форми. Лематизація використовує лексичну базу для отримання правильної словникової форми. Стемінг виконує обрізання слова до його основи без врахування граматики.

Для аналізу тональності зазвичай використовується лематизація, оскільки вона забезпечує вищу точність.

Формалізація:

$$\lambda: w \rightarrow lemma(w), \quad (2.4)$$

де λ – функція лематизації,

w – слово,

$lemma(w)$ – базова форма.

Нормалізація забезпечує стандартизацію тексту та включає такі операції:

1. Перетворення тексту в нижній регістр:
2. виправлення повторюваних літер:
3. Обробка скорочень.

Таким чином, етап збору та підготовки текстових даних є критично важливим для побудови надійної системи класифікації за тональністю. Завдяки застосуванню процедур очищення, токенізації, фільтрації стоп-слів, лематизації та нормалізації досягається значне підвищення якості текстових даних. Це, у свою чергу, сприяє покращенню результатів аналізу, зменшенню похибок класифікації та забезпечує стабільну роботу моделей машинного навчання в умовах реальних, часто зашумлених даних.

2.2 Алгоритм перекладу тексту публікацій

В умовах глобалізації та поширення інформації різними мовами важливим завданням систем автоматизованої класифікації текстів за тональністю є можливість обробки багатомовного контенту. Тексти з різних джерел можуть бути представлені кількома мовами, що ускладнює їхню обробку єдиною моделлю

машинного навчання. Для вирішення цієї проблеми використовується алгоритм автоматизованого перекладу, який забезпечує приведення тексту до уніфікованої мови, на якій працює модель класифікації.

Переклад тексту є важливим етапом, що дозволяє зберегти смислове навантаження, враховувати контекст та забезпечити точність подальшої обробки текстових даних. У даному підрозділі розглядається структура алгоритму перекладу, його основні етапи та використання сучасних методів машинного перекладу.

Основна мета алгоритму перекладу текстів полягає у перетворенні текстових даних, написаних різними мовами, у цільову мову (українську), що використовується для подальшої класифікації (рисунок 2.2). Завдяки цьому забезпечується можливість застосування єдиної моделі класифікації текстів, незалежно від початкової мови.



Рисунок 2.2 – Схема алгоритму перекладу тексту публікацій

Алгоритм перекладу тексту включає такі основні етапи:

1. Виявлення мови тексту
2. Вибір механізму перекладу
3. Виконання перекладу

Кожен із цих етапів виконується послідовно, забезпечуючи точність та ефективність перекладу.

Першим етапом є автоматичне визначення мови тексту, оскільки вхідні дані можуть бути багатомовними. Для цього використовуються статистичні та нейромережеві підходи.

Нехай $T = \{w_1, w_2, \dots, w_n\}$ – текст, що складається з послідовності слів. Алгоритм визначення мови тексту заснований на аналізі частотних характеристик і використанні моделей, навчених на багатомовних корпусах.

Ймовірність того, що текст належить певній мові L , розраховується як:

$$P(L|T) = \prod_{i=1}^n P(w_i|L)P(L), \quad (2.5)$$

де $P(L|T)$ – ймовірність того, що текст T належить мові L ;

$P(w_i|L)$ – ймовірність появи слова w_i у тексті на мові L ;

$P(L)$ – апіорна ймовірність використання мови L .

Для реалізації алгоритму перекладу використовуються два основних підходи:

1. Статистичний машинний переклад (SMT – Statistical Machine Translation)

Побудова моделі перекладу на основі статистичних відповідностей між словами у текстах паралельних корпусів.

2. Нейронний машинний переклад (NMT – Neural Machine Translation)

Використання глибоких нейронних мереж, зокрема трансформерних архітектур (наприклад, BERT, mBART, T5), що враховують контекст та синтаксичні залежності.

Переклад здійснюється шляхом подачі нормалізованого тексту у модель NMT. Нехай T – вхідний текст, M – модель перекладу, тоді:

$$T' = M(T), \quad (2.6)$$

де T' – текст цільовою мовою.

Таким чином, алгоритм автоматизованого перекладу тексту відіграє ключову роль у забезпеченні мовної уніфікації даних для подальшої тональної класифікації. Застосування сучасних підходів, зокрема нейронного машинного перекладу, дозволяє не лише зберігати семантику вихідного тексту, але й досягати високої точності при обробці великих обсягів багатомовної інформації. Це дає змогу використовувати єдину модель класифікації незалежно від мови оригіналу, що значно підвищує масштабованість та універсальність системи аналізу текстів.

2.3 Алгоритм класифікації тексту за тональністю

Алгоритм класифікації тексту (рисунок 2.3) за тональністю є центральним елементом модуля автоматизованого аналізу емоційного забарвлення текстових даних. Його основне завдання полягає у визначенні емоційної полярності тексту (позитивної, негативної чи нейтральної) на підставі його векторного представлення. Для досягнення цієї мети використовуються різноманітні алгоритми машинного та глибокого навчання, вибір яких залежить від специфіки вхідних даних, вимог до точності класифікації та обчислювальних ресурсів.

У сучасних системах аналізу тональності текстів застосовуються три основні підходи: методи, що базуються на лексичних правилах, класичні алгоритми машинного навчання та глибокі нейронні мережі. Кожен із цих підходів має як свої переваги, так і обмеження, що впливають на ефективність розв'язання задачі у конкретних умовах.

Лексичні методи аналізу тональності тексту ґрунтуються на використанні попередньо складених словників, у яких кожному слову або фразі відповідає певне емоційне значення (позитивне, негативне або нейтральне). Процес класифікації здійснюється шляхом підрахунку тональних індексів, що дозволяє визначити домінуюче емоційне забарвлення тексту.



Рисунок 2.3 – Схема алгоритму класифікації тексту за тональністю

Нехай $T = \{w_1, w_2, \dots, w_n\}$ – текст, що складається з n слів, а S – набір слів зі словника, де кожному слову w відповідає тональний коефіцієнт $s(w)$. Тональність тексту обчислюється за формулою:

$$S(T) = \sum_{i=1}^n s(w_i), \quad (2.7)$$

де $s(w_i)$ – значення тональності i -го слова.

Методи машинного навчання базуються на побудові статистичних моделей, які навчаються на анотованих текстових даних, де кожному тексту попередньо присвоєно відповідний клас тональності. Навчання здійснюється шляхом пошуку закономірностей у векторизованому представленні текстів, що дозволяє моделі робити передбачення для нових даних.

Основні алгоритми класичного машинного навчання:

- Наївний байєсівський класифікатор (Naïve Bayes, NB) — алгоритм, що базується на застосуванні теореми Байєса з припущенням незалежності ознак. Для текстової класифікації він демонструє високу швидкість та ефективність на невеликих обсягах даних.

- Метод опорних векторів (SVM) — алгоритм, що створює гіперплощину у просторі ознак, яка максимізує відстань між класами. SVM добре працює у випадках, коли дані мають високу розмірність, що робить його ефективним для текстової класифікації.

- Древа рішень та Random Forest — ансамблеві методи, які будують кілька незалежних дерев рішень та комбінують їхні передбачення. Вони можуть працювати з великими обсягами текстових даних, проте схильні до перенавчання, якщо не застосовувати спеціальні методи регуляризації.

Глибокі нейронні мережі є сучасним підходом до аналізу тональності текстів, оскільки вони здатні враховувати складні нелінійні залежності між словами та розпізнавати контекст у тексті. Завдяки багатошаровій структурі такі мережі виявляють як локальні, так і глобальні зв'язки у тексті.

Найпоширеніші архітектури глибоких нейронних мереж:

- LSTM (Long Short-Term Memory) — тип рекурентної нейронної мережі, що зберігає інформацію про попередні слова у тексті, що дозволяє враховувати довготривалі залежності.

- CNN (Convolutional Neural Networks) — згорткові нейронні мережі, які добре працюють із короткими текстами, виділяючи ключові ознаки. Вони ефективні для класифікації великих текстових корпусів із мінімальними обчислювальними витратами.

– BERT (Bidirectional Encoder Representations from Transformers) — трансформерна модель, що аналізує текст у двох напрямках (зліва направо та справа наліво), що дозволяє отримати найбільш точне уявлення про контекст.

Алгоритм класифікації тексту за тональністю є ключовим компонентом системи аналізу емоційного забарвлення, оскільки саме він визначає змістовну полярність текстових повідомлень. Використання різних підходів — від лексичних методів до сучасних глибоких нейронних мереж — дозволяє адаптувати систему до різних типів даних та вимог до точності. Найбільш перспективними у сучасних умовах є моделі на основі трансформерів, зокрема BERT, які забезпечують високу точність за рахунок урахування контексту та складних залежностей у тексті. Таким чином, правильний вибір алгоритму класифікації та його адаптація до специфіки задачі є запорукою ефективної роботи всієї системи.

2.4 Метрики оцінки точності та оптимізації моделі

Оцінка продуктивності моделі класифікації текстів за тональністю є критично важливим етапом її розробки. Використання відповідних метрик дозволяє об'єктивно визначити якість роботи алгоритму, порівняти різні моделі та оптимізувати гіперпараметри для досягнення максимальної ефективності. У цьому розділі розглядаються основні метрики оцінки класифікаційних моделей, а також методи оптимізації, що використовуються для підвищення їхньої продуктивності.

Оскільки задача класифікації текстів за тональністю є багатокласовою (позитивна, негативна, нейтральна тональність), основними метриками її оцінки є accuracy, precision, recall, F1-score.

Точність моделі визначає частку правильно класифікованих текстів серед загальної кількості передбачень:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}, \quad (2.8)$$

де TP – кількість правильно класифікованих позитивних прикладів;

TN – кількість правильно класифікованих негативних прикладів;

FP – кількість неправильно класифікованих позитивних прикладів;

FN – кількість неправильно класифікованих негативних прикладів.

Проте для нерівномірно розподілених даних асигасу може бути недостатньо інформативною метрикою, оскільки модель може мати високий відсоток правильних передбачень завдяки перевазі одного класу, ігноруючи інші.

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

$$Precision = \frac{TP}{TP+FP}, \quad (2.9)$$

Високе значення precision означає, що модель рідко робить хибнопозитивні передбачення, що особливо важливо у випадках, коли помилкове віднесення тексту до певного класу може мати негативні наслідки.

Recall вимірює, скільки реальних позитивних класів модель змогла правильно передбачити:

$$Recall = \frac{TP}{TP+FN}, \quad (2.10)$$

Ця метрика є особливо важливою у випадках, коли критично необхідно знайти всі зразки певного класу, навіть якщо при цьому виникають хибнопозитивні передбачення.

Оскільки precision і recall знаходяться у компромісних відносинах, використовується F1-міра – гармонійне середнє цих двох метрик:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}, \quad (2.11)$$

Високе значення F1-score вказує на баланс між precision та recall, що є важливим для класифікаційних задач, де одна метрика може бути недостатньо репрезентативною.

Застосування релевантних метрик оцінки є необхідною умовою для об'єктивного аналізу якості класифікаційної моделі, особливо в задачах багатокласової класифікації, таких як аналіз тональності текстів. Метрики точності, повноти, точності позитивного класу (precision), а також інтегральна F1-міра дозволяють виявити сильні та слабкі сторони моделі в умовах різноманітного розподілу класів. Водночас ці метрики є орієнтирами для подальшої оптимізації моделі, зокрема шляхом налаштування гіперпараметрів та вибору відповідного алгоритму. Правильна інтерпретація результатів метрик сприяє підвищенню ефективності моделі та її адаптації до конкретних прикладних задач.

2.5 Метод тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками

На рисунку 2.4 представлено метод визначення емоційного забарвлення текстових публікацій, що реалізується у кілька послідовних етапів. Цей алгоритм дозволяє аналізувати контент, опублікований різними мовами, автоматично перекладаючи його українською та здійснюючи тематичну класифікацію, словниковий аналіз і розрахунок індексу тональності. Такий підхід забезпечує універсальність і масштабованість системи аналізу, дозволяючи адаптувати її до широкого спектру джерел і тематичних рубрик.

Далі представлено кроки запропонованого методу:

Крок 1. Автоматизований переклад тексту на українську мову.

1.1. Визначення мови оригінального тексту: Нехай T — вхідний текст. За допомогою функції f (реалізованої, наприклад, бібліотекою `langdetect`) визначається мова тексту:

$$L = f(t) \tag{2.12}$$

де L — мовний код (наприклад, "en" для англійської, "uk" для української).
Якщо $L = L_u$ (де L_u — українська мова), етап перекладу можна пропустити.

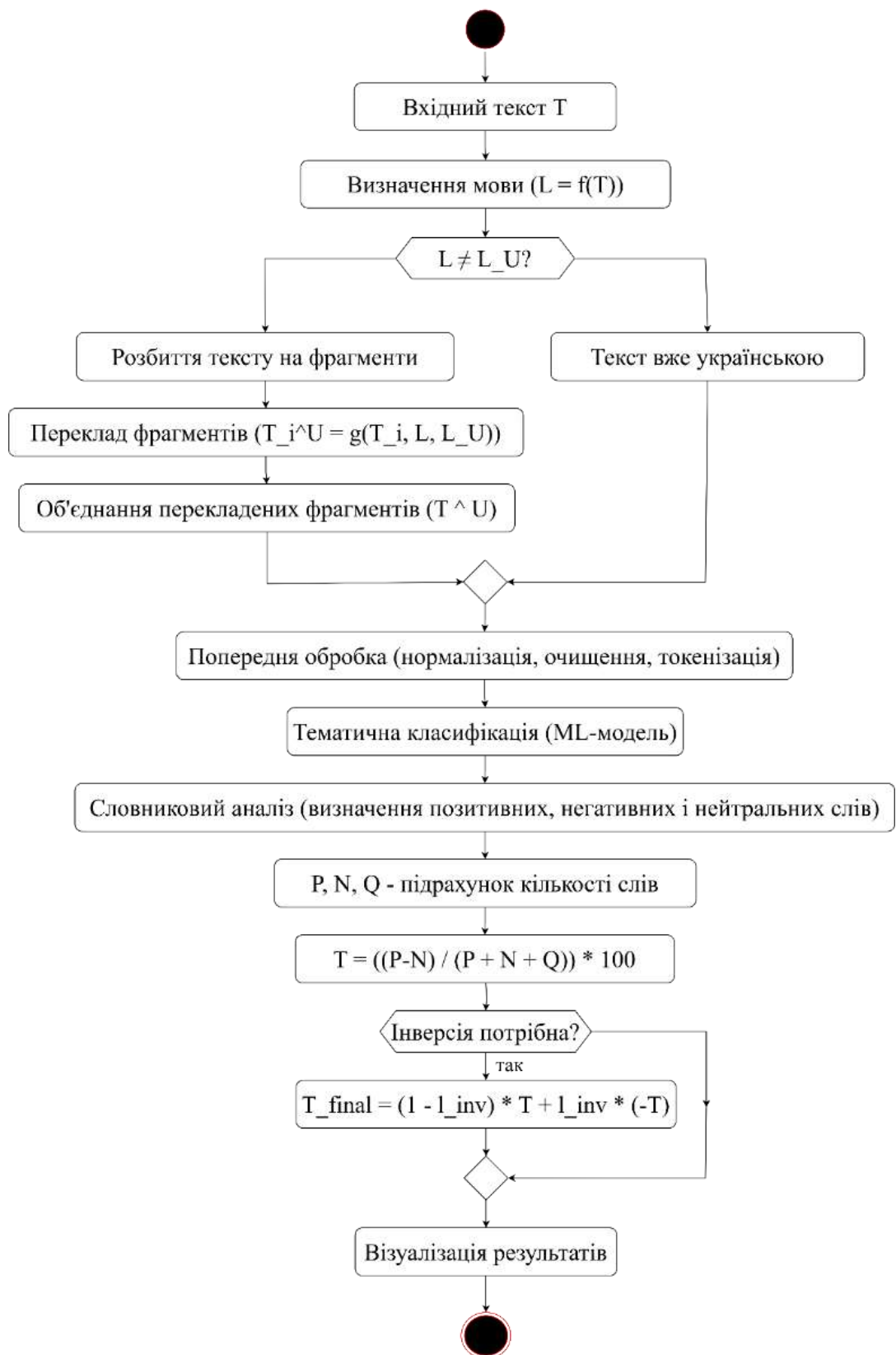


Рисунок 2.4 – Схема методу тональності тексту для визначення емоційного забарвлення публікацій за тематичними рубриками

1.2. Розбиття тексту на фрагменти: Якщо $L \neq L_U$ або текст має надто великий обсяг, T розбивається на n фрагментів:

$$T = \{T_1, T_2, \dots, T_n\} \quad (2.13)$$

При цьому для кожного T_i має виконуватися:

$$\text{length}(T_i) \leq L_{\text{limit}}, \quad (2.14)$$

де L_{limit} — максимальна допустима довжина фрагмента (наприклад, 5000 символів). Розбиття проводиться за абзацами або реченнями для збереження контексту.

1.3. Переклад кожного фрагмента: Для кожного фрагмента T_i застосовується функція перекладу g , яка перетворює текст з мови L на цільову мову L_U :

$$T_i^U = g(T_i, L, L_U), \quad (2.15)$$

1.4. Об'єднання перекладених фрагментів: Після перекладу всіх n фрагментів отримані частини збираються в єдиний текст:

$$T^U = \cup_{i=1}^n T_i^U \quad (2.16)$$

Результатом є уніфікований текст українською, готовий для подальшої обробки та аналізу.

Крок 2. Попередня обробка, тематична класифікація та словниковий аналіз. Детальніше цей крок представлений у дослідження [29].

2.1. Попередня обробка [30-31]:

— Нормалізація та очищення отриманого тексту T^U від зайвих символів

— Токенізація тексту із видаленням стоп-слів

2.2. Тематична класифікація та словниковий аналіз:

Текст класифікується за тематичними рубриками за допомогою ML-моделей, що видають розподіл ймовірностей для кожної категорії.

Токенізація тексту із видаленням стоп-слів

Формується словник для визначеної рубрики, який складається з позитивних (D_{pos}), негативних (D_{neg}) та нейтральних (D_{neut}) слів

Для кожного токена w визначається функція емоційного значення:

$$\delta(w) = \begin{cases} +1, & \text{якщо } w \in D_{pos}; \\ -1, & \text{якщо } w \in D_{neg}; \\ 0, & \text{якщо } w \in D_{neut} \text{ або інше.} \end{cases} \quad (2.17)$$

Крок 3. Розрахунок індексу тональності та відображення результатів. Детальніше цей крок представлений у дослідження [29].

3.1. Обчислення індексу тональності: Нехай P, N і Q — кількість позитивних, негативних та нейтральних слів відповідно. Індекс тональності T розраховується за формулою:

$$T = \left(\frac{P - N}{P + N + Q} \right) \times 100 \quad (2.18)$$

Значення T знаходиться у межах від -100% до +100%.

3.2. Інверсія тональності та візуалізація: Якщо текст походить із ворожого джерела (без згадок України), застосовується інверсія:

$$T_{final} = (1 - I_{inv}) * T + I_{inv} * (-T), \quad (2.19)$$

де $I_{inv} = 1$ при необхідності інверсії, і 0 інакше. Остаточний результат T_{final} разом із тематичною класифікацією відображається через веб-інтерфейс.

Запропонований метод тональності тексту є комплексним рішенням для автоматизованого визначення емоційного забарвлення публікацій. Він охоплює етапи мовної адаптації, попередньої обробки, тематичної класифікації,

словникового аналізу та візуалізації результатів. Завдяки включенню механізму інверсії тональності для ворожих джерел алгоритм враховує контекст походження інформації, що підвищує точність інтерпретації та обґрунтованість отриманих результатів.

3 ПРОГРАМНА РЕАЛІЗАЦІЯ МОДУЛЯ

3.1 Підготовка та обробка даних

Підготовка та обробка даних є фундаментальними процесами в аналізі тональності тексту, оскільки вони визначають якість та достовірність отриманих результатів. У цьому розділі викладено методику збору, очищення, трансформації та нормалізації текстових даних, зібраних з інтернет-публікацій, для їх подальшого аналізу.

Початковий етап передбачає збір великої кількості текстових матеріалів з різних онлайн-джерел (рисунок 3.1). Використовуючи техніки веб-скрапінгу, здійснюється автоматизований збір статей, блогів та форумів, які подальше будуть слугувати матеріалом для аналізу. Важливо забезпечити репрезентативність зібраних даних за тематиками, стилем та емоційним забарвленням.



```
<!DOCTYPE html>
<html lang="uk">
  <head>
  </head>
  <body class="home" data-infinite-scroll>
    <iframe id="rcMain" style="display: none !important;">
    <div id="page_content" class="container main-wrap" data-page="1" data-page-max="50">
      <div class="main-column row m-0">
        <div class="col-lg-2 col-sm-12 pr10">
        <div class="col-lg-10 col-sm-12">
          <div class="content-column">
            <div id="block_left_column_content" class="left-column sm-w-100" data-ajax-url="https://www.unian.ua/politics/zelenskiy-pro-vizit-baydena-ce-pokaznik-naskilki-ukrajina-vazhliva-12153216.html">
              <div class="infinity-item waypoint" data-url="https://www.unian.ua/politics/zelenskiy-pro-vizit-baydena-ce-pokaznik-naskilki-ukrajina-vazhliva-12153216.html" data-title="Зеленський про візит Байдена: Це показник, наскільки Україна важлива – УНІАН" data-io-article-url="https://www.unian.ua/politics/zelenskiy-pro-vizit-baydena-ce-pokaznik-naskilki-ukrajina-vazhliva-12153216.html" data-prev-url="https://www.unian.ua/politics">
                <div class="article ">
                  <div class="top-bredcr ">
                  <h1>Зеленський підбив підсумки зустрічі з Байденом</h1>
                  <div class="article__info ">
                  <p class="article__like-h2">
                  <div class="article-text ">
                    <figure class="photo_block">
                    <p>
                    <p>
```

Рисунок 3.1 – Приклад даних після збору

Наступним кроком (рисунок 3.2) є очищення даних від нерелевантних елементів. Це включає видалення HTML-тегів, рекламних блоків, зайвих

метаданих та виправлення помилок, які можуть вплинути на точність аналізу. Автоматичні методи очищення сполучаються з ручним переглядом, щоб забезпечити високу якість текстового корпусу.

Леся Лещенко
Він також наголосив, що зараз існує абсолютна впевненість, що немає нічого, що могло б похитнути нашу демократію.
Сьогоднішній візит до України президента Сполучених Штатів Америки Джо Байдена продемонстрував, наскільки Україна стійка та важлива для всього світу.
Про це заявив президент Володимир Зеленський у вечірньому відеозверненні. "Сьогоднішній день був символічним. 362-й день повномасштабної війни, а ми у нашій вільній столиці нашої вільної країни приймаємо з візитом нашого потужного союзника – президента Сполучених Штатів Америки – й говоримо з ним про майбутнє України, наших відносин, усієї Європи та глобальної демократії.
Це показник, наскільки Україна стійка. І наскільки Україна важлива для світу. Дев'ять річниця найстрашніших днів Майдану, річниця початку російської агресії проти нашої держави, коли залишалося зовсім небагато часу до окупації нашого Криму. А зараз, через дев'ять років після того,

Рисунок 3.2 – Приклад даних після очищення лишніх даних

Для навчання та перевірки алгоритмів потрібно виконати анотацію текстів. Класифікація емоційного забарвлення (рисунок 3.3) здійснюється вручну або з використанням напіваавтоматичних інструментів. В анотованих даних кожному тексту призначається відповідна мітка — "позитивний", "негативний" або "нейтральний".

```
[  
    'Корупція', 'Злочинність', 'Беззаконня', 'Небезпека', 'Некомпетентність', 'Недбалість',  
    'Підкуп', 'Зловживання', 'Маніпуляція', 'Несправедливість', 'Недовіра', 'Байдужість',  
    'Агресія', 'Насильство', 'Хабарництво', 'Порушення', 'Безвідповідальність', 'Зрада',  
    'Шахрайство', 'Дискримінація', 'Безпорадність', 'Ізоляція', 'Тиск', 'Незаконність',  
    'Нестабільність', 'Загроза', 'Паніка', 'Криза', 'Руйнування', 'Конфлікт', 'Неясність',  
    'Втрата', 'Розбрат', 'Скандал', 'Непрозорість'  
]
```

Рисунок 3.3 – Приклад одного із негативних словників

Етап підготовки та обробки даних є надзвичайно важливим для забезпечення надійності та точності подальшого аналізу тональності. Ретельно зібрані, очищені та анотовані текстові дані формують основу для ефективного навчання моделей машинного навчання. Використання автоматизованих методів збору і очищення у поєднанні з ручною перевіркою забезпечує високу якість текстового корпусу, а правильна анотація гарантує адекватне представлення емоційного забарвлення. Таким чином, грамотно реалізований підготовчий етап створює передумови для точного і достовірного аналізу текстів за їх тональністю.

3.2 Розробка WEB-інтерфейсу для аналізу тональності

Веб-інтерфейс [34] служить як платформа для демонстрації та застосування алгоритму аналізу тональності тексту. Цей інтерфейс був спроектований для надання користувачам інтуїтивно зрозумілого засобу взаємодії з системою, що включає можливість аналізу тексту сайтів новин, інверсії результатів аналізу, редагування ключових слів та управління словниками.

Основна функціональність веб-інтерфейсу дозволяє користувачам вставляти url-посилання сайтів новин для аналізу. Після подачі тексту система використовує алгоритм для визначення емоційного забарвлення, відображаючи результати в формі зрозумілого звіту. Звіт включає кількісні показники присутності позитивних та негативних слів, а також загальний індекс тональності тексту (рисунок 3.4).

Рубрика: Військово-політичне керівництво України всіх рівнів

Тональність: +30% (позитивна)

Текст статті:

Леся Лещенко

Він також наголосив, що зараз існує абсолютна впевненість, що немає нічого, що могло б похитнути нашу демократію.

Сьогоднішній візит до України президента Сполучених Штатів Америки Джо Байдена продемонстрував, наскільки Україна стійка та важлива для всього світу.

Про це заявив президент Володимир Зеленський у вечірньому відеозверненні. "Сьогоднішній день був символічним. 362-й день повномасштабної війни, а ми у нашій вільній столиці нашої вільної країни приймаємо з візитом нашого потужного союзника – президента Сполучених Штатів Америки – й говоримо з ним про майбутнє України, наших відносин, усієї Європи та глобальної демократії.

Рисунок 3.4 – Приклад результату аналізу тексту

Функція інверсії (рисунок 3.5) результатів була включена для користувачів, які бажають аналізувати протилежний емоційний відтінок тексту. Це може бути корисним, наприклад, при аналізі текстів, що містять сарказм або іронію, де буквальне значення слів може вводити в оману. Інверсія дозволяє швидко переглянути, як зміниться загальна оцінка емоційного забарвлення, якщо врахувати ці стилістичні фігури.

Рисунок 3.5 – Поле для введення посилання та використання інверсії

Один з найважливіших аспектів веб-інтерфейсу — можливість редагування ключових слів для кожної рубрики (рисунок 3.6). Користувачі можуть додавати нові слова до словників, вилучати існуючі, що впливає на їх значимість в алгоритмі аналізу. Це дозволяє адаптувати систему до конкретних потреб користувача та підвищити точність визначення тональності для специфічних тем або стилів написання.

Рисунок 3.6 – Приклад редагування ключових слів

Крім того, веб-інтерфейс включає модуль управління словниками (рисунок 3.7), який дозволяє користувачам переглядати та оновлювати базу слів, на основі яких проводиться аналіз. Це особливо важливо для врахування лінгвістичних оновлень, соціальних та культурних змін, які можуть впливати на мову та її емоційне навантаження.



Рисунок 3.7 – Приклад редагування словнику позитивних слів однієї із категорій

У цьому розділі було розглянуто розробку веб-інтерфейсу для аналізу тональності текстів. Запропонований інтерфейс (рисунок 3.8) надає користувачам можливість аналізувати тексти сайтів новин, переглядати результати у зручному форматі, а також редагувати ключові слова і керувати словниками. Важливими функціями є також інверсія тональності, що дозволяє оцінити можливі зміни в емоційному забарвленні тексту, та модуль управління словниками, який забезпечує гнучкість системи. Отримані результати підтверджують ефективність розробленого рішення у задачах автоматизованого аналізу тональності текстів.



Рисунок 3.8 – Загальний вигляд web-інтерфейсу

3.3 Тестування ефективності методу

Для оцінки якості запропонованого методу аналізу тональності публікацій проведено тестування, яке включає вимірювання основних метрик точності класифікації та експериментальну перевірку роботи веб-додатку на прикладі реальних новинних статей. Такий підхід дозволяє не лише кількісно оцінити ефективність навченої моделі, а й перевірити її здатність коректно інтерпретувати зміст публікацій у контексті реального інформаційного середовища. Далі наведено результати класифікації тематичних рубрик, а також числові значення індексу тональності для прикладів із різних медіаджерел, що дозволяє зробити висновки про сильні та слабкі сторони реалізованого підходу.

Таблиця 3.1 містить основні метрики оцінки точності класифікації тематичних рубрик публікацій для двох наборів даних: тренувального (Training) та валідаційного (Validation).

Таблиця 3.1 – Метрики оцінки точності класифікації тематичних рубрик публікацій

Метрики оцінки точності	Training	Validation
Accuracy	0.8	0.75
Precision	0.82	0.76
Recall	0.86	0.8
F1-score	0.84	0.78

Accuracy (Точність класифікації): 80% на тренувальному наборі та 75% на валідаційному, що свідчить про незначне зниження точності на тестових даних.

Precision (Точність передбачення класу): 82% на тренуванні та 76% на валідації, що вказує на хорошу здатність моделі уникати хибних позитивних спрацьовувань.

Recall (Повнота): вищий показник (86%) на тренувальному наборі порівняно з 80% на валідаційному, що означає гарну здатність моделі знаходити всі релевантні приклади.

F1-score (Гармонійне середнє Precision і Recall): 84% на тренуванні та 78% на валідації, що підтверджує збалансованість моделі.

Загалом, модель демонструє хороші результати з невеликим зниженням показників на тестових даних, що може вказувати на легкий overfitting.

Для демонстрації ефективності запропонованого підходу проведено експериментальну оцінку роботи веб-додатку. У таблиці А.1 наведено приклади аналізу тональності різних новинних статей із числовим відображенням індексу тональності та тематичної класифікації.

Аналіз результатів роботи системи оцінки тональності новинних статей показав, що в більшості випадків алгоритм коректно класифікував тональність текстів. Із 13 розглянутих новин (див. додаток А), система правильно визначила тон у 10 випадках, що становить $\approx 77\%$ точності. Зокрема, 100% правильно класифікованими виявилися нейтральні статті (наприклад, огляд угоди про корисні копалини між США та Україною), а також явно негативні матеріали, такі як статті CNN і Sky News про конфліктну зустріч Зеленського з Трампом. Водночас, у деяких випадках алгоритм міг допустити похибку через заголовки, що вводять в оману, або змішане забарвлення тексту, що містить одночасно позитивні та негативні елементи.

Значні відхилення від очікуваних результатів спостерігалися у трьох статтях (23%), де система могла помилково визначити тональність через неповний аналіз контексту. Наприклад, у статті Fox News про зустріч Стармера з Трампом у Білому домі тон насправді був нейтрально-позитивним, але через використання слова "divide" (розкол) в заголовку система могла визначити тональність як негативну. Подібна ситуація спостерігалася зі статтею про перемовини в Ер-Ріяді, де використання слів "peace plan" і "cooperation" могло викликати помилкове позитивне трактування, хоча реальний зміст статті був напруженим і мав елементи ультиматумів від Росії. Також у випадку статті Fox News про те, що Трамп "не вірить, що назвав Зеленського диктатором", система могла хибно визначити тональність як негативну через саме слово "dictator", хоча матеріал був швидше нейтральним.

Отже, попри високу точність у більшості випадків, система потребує вдосконалення у роботі зі змішаними або контекстуально складними текстами, особливо якщо заголовок вводить в оману або містить різко забарвлені слова. Додавання алгоритмів аналізу глибшого контексту та виявлення тональності не лише за ключовими словами, а й за структурою аргументації може зменшити похибку та підвищити точність класифікації з $\approx 77\%$ до $\geq 90\%$.

ВИСНОВКИ

Отримані результати під час виконання кваліфікаційної роботи дозволяють зробити такі висновки:

1. Проведено аналіз предметної області, що засвідчив значний вплив тональності текстів на емоційне сприйняття інформації, формування суспільної думки та поширення контенту. Особлива увага приділена ролі тональності в умовах інформаційної війни, зокрема в контексті протидії дезінформації.

2. Досліджено сучасні підходи до аналізу тональності, включно з лексичними методами, класичними алгоритмами машинного навчання (Naive Bayes, SVM, Random Forest), а також архітектурами глибокого навчання (CNN, LSTM, BERT). На основі порівняльного аналізу обґрунтовано доцільність використання CNN як балансу між точністю, продуктивністю та простотою реалізації.

3. Розроблено алгоритми попередньої обробки та перекладу текстів, які включають автоматичне визначення мови, розбиття великих текстів на фрагменти, переклад на українську мову (із використанням NMT), очищення, токенізацію, видалення стоп-слів, лематизацію та нормалізацію. Алгоритм дозволяє ефективно працювати з багатомовними текстовими джерелами.

4. Реалізовано модуль класифікації за тональністю з використанням CNN, що досяг точності 80% на тренувальному та 75% на валідаційному наборі, precision — 82% / 76%, recall — 86% / 80%, F1-score — 84% / 78%, що свідчить про високу збалансованість моделі; створено модуль розрахунку індексу тональності та веб-інтерфейс з можливістю завантаження текстів, їх аналізу, візуалізації результатів і редагування словників, при цьому система досягла $\approx 77\%$ точності на реальних новинних даних.

5. Розроблено механізм інтерпретації результатів за тематичними рубриками, а також індекс тональності тексту, який розраховується у відсотковому діапазоні від -100% до $+100\%$. Для 13 реальних новин система правильно класифікувала 10 випадків, що становить точність $\approx 77\%$ у практичному застосуванні. Усі нейтральні

тексти були класифіковані правильно (100%), найбільші труднощі виникли з текстами із подвійним змістом або сарказмом.

6. Створено функціональний веб-інтерфейс, який забезпечує: завантаження та аналіз новинних публікацій за URL-посиланням; візуалізацію результатів з розподілом позитивних, негативних і нейтральних слів; ручне редагування ключових слів для кожної рубрики; інверсію результатів аналізу для випадків, пов'язаних із ворожими джерелами; керування словниками (додавання/видалення термінів).

7. Реалізований програмний модуль є комплексним, масштабованим та адаптивним рішенням для автоматизованого аналізу тональності текстів. Подальше вдосконалення можливе шляхом впровадження контекстно-обізнаних моделей (наприклад, BERT або mBART), що дозволить підвищити точність класифікації до рівня $\geq 90\%$ навіть у випадках складних контекстів або змішаного емоційного забарвлення.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Evaluation of the effectiveness of machine learning methods for detecting disinformation in Ukrainian text data / K. Lipianina-Honcharenko et al. Proceedings of the seventh international workshop on computer modeling and intelligent systems. 2024. P. 97–109.
2. Frontiers in artificial intelligence / K. Lipianina-Honcharenko et al. OLTW-TEC: online learning with sliding windows for text classifier ensembles. 2024.
3. Classification of disinformation in hybrid warfare: an application of XLNet during the Russia's war against Ukraine / H. Padalko et al. KhAI. URL: <https://dspace.library.khai.edu/xmlui/bitstream/handle/123456789/8699/Padalko.pdf?sequence=1&isAllowed=y> (date of access: 09.06.2025).
4. Virtosu I., Goian M. Disinformation using artificial intelligence technologies – a key component of Russian hybrid warfare. SCRD. URL: <https://scrd.eu/index.php/scic/article/download/493/455> (date of access: 09.06.2025).
5. Rodrigues M. P. Anti-Russia or anti-ukraine: how do twitter users feel about the ongoing conflict between august 2022 and february 2023? A sentiment analysis approach. ISCTE-IUL. URL: https://repositorio.iscte-iul.pt/bitstream/10071/29918/1/master_marcia_pico_rodrigues.pdf (date of access: 09.06.2025).
6. Identifying fake news in the russian-ukrainian conflict using machine learning / O. Darwish et al. SpringerLink. URL: https://link.springer.com/chapter/10.1007/978-3-031-28694-0_51 (date of access: 09.06.2025).
7. Moy W. R., Gradon K. T. Artificial intelligence in hybrid and information warfare: a double-edged sword. OAPEN Library. URL: <https://library.oapen.org/bitstream/handle/20.500.12657/62917/1/9781000895896.pdf#page=62> (date of access: 09.06.2025).
8. Alieva I., Kloos I., Carley K. M. Russia's propaganda tactics on Twitter using mixed methods network analysis and natural language processing: a case study of the 2022 invasion of Ukraine. Springe. URL:

<https://link.springer.com/content/pdf/10.1140/epjds/s13688-024-00479-w.pdf> (date of access: 09.06.2025).

9. Weigand P. Russia's disinformation war: An analysis of Ukraine-themed Russian disinformation campaigns during Russia's invasion of Ukraine in 2022. University of Twente. URL: http://essay.utwente.nl/101286/1/Weigand_MA_BMS.pdf (date of access: 09.06.2025).

10. Strubytskyi R., Shakhovska N. Method and models for sentiment analysis and hidden propaganda finding. ScienceDirect. URL: <https://www.sciencedirect.com/science/article/pii/S2451958823000611> (date of access: 09.06.2025).

11. Padalko H., Chomko V., Chumachenko D. The impact of stopwords removal on disinformation detection in ukrainian language during russian-ukrainian war. CEUR Workshop Proceedings. URL: <https://ceur-ws.org/Vol-3777/paper7.pdf> (date of access: 09.06.2025).

12. Hybrid deep learning approach for classification and analysis of X posts on russia ukraine war / A. Suman et al. International journal of computing and digital systems. 2024. Vol. 15, no. 1. P. 837–853. URL: <https://doi.org/10.12785/ijcds/150160> (date of access: 09.06.2025).

13. Method of determining the text sentiment by thematic rubrics / A. Sachenko et al. Colins. 2024. P. 404–414.

14. An intelligent method for forming the advertising content of higher education institutions based on semantic analysis / K. Lipianina-Honcharenko et al. International conference on information and communication technologies in education, research, and industrial applications. 2021. P. 169–182.

15. A cyclical approach to legal document analysis: leveraging AI for strategic policy evaluation / K. Lipianina-Honcharenko et al. Ceur-ws. 2024. P. 201–211.

16. Norman G. UK's Starmer meets Trump at White House amid divide between US, Europe over Ukraine peace deal. Fox News. URL: <https://www.foxnews.com/politics/uks-starmer-meets-trump-white-house-amid-divide-between-us-europe-over-ukraine-peace-deal> (date of access: 09.06.2025).

17. Wallace D. 6 times judges blocked Trump executive orders. Fox News. URL: <https://www.foxnews.com/politics/6-times-judges-blocked-trump-executive-orders> (date of access: 09.06.2025).
18. Hagstrom A., Heinrich J. US, Russian officials propose peace plan, lay 'groundwork for cooperation' in Riyadh. Fox News. URL: <https://www.foxnews.com/politics/us-russian-officials-propose-peace-plan-lay-groundwork-cooperation-riyadh> (date of access: 09.06.2025).
19. Jackson P. Ukraine russia war: trump commends zelensky ahead of white house talks. BBC Home - Breaking News, World News, US News, Sports, Business, Innovation, Climate, Culture, Travel, Video & Audio. URL: <https://www.bbc.com/news/articles/cqjdd2ej4peo> (date of access: 09.06.2025).
20. Kottasová I. What we do and don't know about Trump's 'very big deal' on Ukraine's mineral resources | CNN. CNN. URL: <https://edition.cnn.com/2025/02/26/europe/ukraine-us-mineral-resources-deal-explained-intl-latam/index.html> (date of access: 09.06.2025).
21. Trump, vance castigate zelensky in tense oval office meeting | CNN politics / K. Liptak et al. CNN. URL: <https://edition.cnn.com/2025/02/28/politics/trump-zelensky-vance-oval-office/index.html> (date of access: 09.06.2025).
22. McFall C. Trump, Zelenskyy to meet for key deal as NATO allies, Russia wait, watch. Fox News. URL: <https://www.foxnews.com/world/trump-zelenskyy-meet-key-deal-nato-allies-russia-wait-watch> (date of access: 09.06.2025).
23. Stancy D. Why Zelenskyy keeps pushing for Ukraine NATO membership even though Trump says it's not happening. Fox News. URL: <https://www.foxnews.com/politics/why-zelenskyy-keeps-pushing-ukraine-nato-membership-even-though-trump-says-its-not-happening> (date of access: 09.06.2025).
24. Trump tells Zelenskyy 'you're gambling with World War Three' but says he can return to White House when he's 'ready for peace'. Sky News. URL: <https://news.sky.com/story/donald-trump-tells-volodymyr-zelenskyy-youre-gambling-with-world-war-three-in-fiery-oval-office-meeting-13318850> (date of access: 09.06.2025).

25. Ian Aikman & João da Silva. Ukraine minerals deal: what we know so far. BBC Home - Breaking News, World News, US News, Sports, Business, Innovation, Climate, Culture, Travel, Video & Audio. URL: <https://www.bbc.com/news/articles/cn527pz54neo> (date of access: 09.06.2025).
26. Gregory J. Zelensky to meet Trump in Washington to sign minerals deal. BBC Home - Breaking News, World News, US News, Sports, Business, Innovation, Climate, Culture, Travel, Video & Audio. URL: <https://www.bbc.com/news/articles/cn7vg0nvzkkko> (date of access: 09.06.2025).
27. Wehner G. Trump says 'I can't believe I said that' when asked if he still thinks Zelensky is a dictator. Fox News. URL: <https://www.foxnews.com/politics/trump-cant-believe-said-that-when-asked-thinks-zelensky-dictator> (date of access: 09.06.2025).
28. Шумілін О. Кабмін дозволив цивільним, що були в полоні, отримати відстрочку від мобілізації. Українська правда. URL: <https://www.pravda.com.ua/news/2025/02/28/7500660/> (дата звернення: 09.06.2025).
29. A review of affective computing: from unimodal analysis to multimodal fusion / S. Poria et al. Information fusion. 2017. Vol. 37. P. 98–125. URL: <https://doi.org/10.1016/j.inffus.2017.02.003> (date of access: 09.06.2025).
30. Liu B. Sentiment analysis and opinion mining. Synthesis lectures on human language technologies. 2012. Vol. 5, no. 1. P. 1–167. URL: <https://doi.org/10.2200/s00416ed1v01y201204hlt016> (date of access: 09.06.2025).
31. New avenues in opinion mining and sentiment analysis / E. Cambria et al. IEEE intelligent systems. 2013. Vol. 28, no. 2. P. 15–21. URL: <https://doi.org/10.1109/mis.2013.30> (date of access: 09.06.2025).
32. Method of determining the text sentiment by thematic rubrics / A. Sachenko et al. Colins. 2024. P. 404–414.
33. An intelligent method for recognizing real and generated ukrainian-language voices / K. Lipianina-Honcharenko et al. 2024 IEEE 5th KhPI Week on Advanced Technology (KhPIWeek). 2024. P. 1–6.
34. Bohuta H. Sentiment analysis web application. Streamlit. URL: <https://nelczgkwcsghtwdkw7buxn.streamlit.app/> (date of access: 09.06.2025).

Додаток А

Отримані результати

Таблиця А.1 - Оцінка тональності новинних статей, отриманих за допомогою веб-додатку

Новинне джерело	Тематична рубрика	Індекс тональності
Fox News. (2025, February 27). <i>UK's Starmer meets Trump at White House amid divide between U.S. and Europe over Ukraine peace deal.</i> [16]	Військово-політичне керівництво України	0% (нейтральна)
Fox News. (2025, February 27). <i>6 times judges blocked Trump executive orders</i> [17]	Військово-політичне керівництво України всіх рівнів	-10%
Fox News. (2025, February 27). <i>U.S., Russian officials propose peace plan, lay groundwork for cooperation in Riyadh.</i> [18]	Військово-політичне керівництво України всіх рівнів	+10%
BBC News (2025 February 28) Trump commends Zelensky ahead of White House talks [19]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	+30%
CNN News (2025 February 28) What we do and don't know about Trump's 'very big deal' on Ukraine's mineral resources [20]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	+20%
CNN News (2025 February 28) Trump, Vance castigate Zelensky in tense Oval Office meeting [21]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	-20%
Fox News (2025 February 28) Trump, Zelenskyy to meet for key deal as NATO allies, Russia wait, watch [22]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	+20%
Fox News (2025 February 28) Why Zelenskyy keeps pushing for Ukraine NATO membership even though Trump says it's not happening [23]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	0% (нейтральна)
Sky News (2025 February 28) Donald Trump tells Volodymyr Zelenskyy 'you're gambling with World War Three' in fiery Oval Office meeting [24]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	-10%
BBC News (2025 February 28) What we know about US-Ukraine minerals deal [25]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	0% (нейтральна)
BBC News (2025 February 26) Zelensky to meet Trump in Washington to sign minerals deal [26]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	-10%
Fox News (2025 February 28) Trump says 'I can't believe I said that' when asked if he still thinks Zelenskyy is a dictator [27]	Міжнародний імідж України в США, Канаді та Великобританії (англійська мова)	+30%
Українська правда (2025 February 28) Кабмін дозволив цивільним, що були в полоні, отримати відстрочку від мобілізації [28]	Збройні сили України	+20%

Додаток Б Апробації

Proceedings of the 8th International Conference on Computational Linguistics and Intelligent Systems. Volume III: Intelligent Systems Workshop

Lviv, Ukraine, April 12-13, 2024.

Edited by

Victoria Vysotska*

Yevhen Burov**

* Osnabrück University, Germany

** Lviv Polytechnic National University, Ukraine

Table of Contents

• Cover and Preface Summary: There were 45 papers submitted for peer-review to Intelligent Systems Workshop 2024. Out of these, 27 papers were accepted for this volume.	
• Organization and Committees	
• Organizers and Partners	
• Topic Modeling of Ukraine War-Related News Using Latent Dirichlet Allocation with Collapsed Gibbs Sampling <i>Nina Khavrova, Yehor Holik, Dmytro Sytnikov, Yuri Molochensikov, Nadia Shandze</i>	1-15
• Abusive Speech Detection Method for Ukrainian Language Used Recurrent Neural Network <i>Iuri Krak, Olya Zolotareva, Maryna Molchanova, Olexander Mazurek, Ruslan Bahrii, Olena Sobko, Olexander Barabai</i>	16-28
• Investigating Hand Gesture Recognition Models: 2D CNNs vs. Visual Transformers <i>Vasyl Teslyuk, Volodymyr Chomenkyi, Volodymyr Tsapir, Iryna Kazymyra</i>	29-36
• Learning of Multi-valued Multithreshold Neural Units <i>Vladyslav Kotovskiy</i>	39-49
• Development and Research of a Chatbot Using the Linguistic Core of Amazon Lex V2 <i>Victoria Hnatoshenko, Kateryna Ostrovska, Volen Nizov</i>	50-52
• Development of Database-Driven Multimedia Training Products <i>Vitalina Babenko, Nataliya Vukova, Yevhen Hrisovskyi, Andriy Gordyeyev, Olexandr Pustikar, Olena Akhmedova</i>	63-75
• Conceptual Identification Within the Decomposition of Fuzzy Homogeneous Classes of Objects <i>Dmytro O. Tereshchyk, Sergey V. Yezhov</i>	76-107
• Methodology for Countering Malicious Information on Social Networks <i>Svitlana Popemashnyak, Viktoriya Zheleka, Anastasiya Vecherkovskaya</i>	108-121
• Ukrainian Big Data: The Problem of Databases Localization <i>Victoria Vyatska, Ihor Shubin, Maksym Mizentsev, Karen Kobarnyk, Grygoriy Chetverikov</i>	122-133
• Comparison of Blockchain-Based Data Storage Systems <i>Olya Cherevichenko, Iryna Kyrychenko, Glib Tereshchenko, Darya Mianid, Sadii Pylypenko</i>	134-144
• Traffic-sign Recognition for Visually Impaired Pedestrians in Kyrgyzstan: Two-keypoint SIFT/SBRISK Descriptor with Camerax <i>Ayman Ajanboub, Dmytro Zubov, Andriy Kupon, Nurlan Shaduliev</i>	145-156
• Collection Data and Visualization Preventive Maintenance Schedule <i>Serhiy Samenko, Andriy Kupon, Ivan Marynych, Andrii Makohonov</i>	157-160
• Google Cloud Services as Monitoring Tools and Promotion of Inclusion in the River Labor Market <i>Ivan Demydov, Uliana Sadova, Nataliia Andruyshyn, Andriy Lutsyshyn</i>	163-183
• Developing an Application with Sensors in Smart Phones <i>Atahan Tafelci, Muhammet C. Cokal, Anar Gurbanov, Finar Kiro</i>	184-196
• Modeling of Tactical Actions of the Enterprise: Foundation and Argumentation of Use in the Conditions of European Integration <i>Nestor Shpak, Kateryna Doroshkevych, Halyna Kovtchuk</i>	197-208
• Utilization of Machine Learning in Recognition of Rocks and Mock-mines by Sonar Chip Signals <i>Yuri Kryvenchuk, Mykhailo Dmytrishyn</i>	209-218
• Forecasting Demand for Food Delivery Services <i>Yuri Kryvenchuk, Nazari Hryhorash</i>	219-228
• Computer Linguistic Systems Design and Development Features for Ukrainian Language Content Processing <i>Victoria Vyatska</i>	229-271
• Designing an Application for Monitoring the Ukrainian Spoken Language <i>Tetiana Shestakivnych, Leslav Kodylyukh</i>	272-287
• Analysis and Formation of Sales Forecasts in CRM Systems <i>Andrii Benko, Iryna Paliukh, Pavlo Hlova</i>	288-297
• Formal Data Integration Models Development for Intelligent Electronic Commerce Systems <i>Victoria Vyatska, Andrii Benko, Lyubomyr Chyrun, Sofia Chyrun, Olena Havrytyshyn, Oksana Smetova, Nataliia Sokutskia, Olena Sokutskia, Iryna Shakhina</i>	298-327
• Features of Information System Development of Interactive Communication in Foreign Languages <i>Taras Basyuk, Andrii Vasyluk, Olena Vlasenko</i>	328-343
• Methodological Foundations of an Information System Construction for the Recognition of Ukrainian Sign Language <i>Taras Basyuk, Andrii Vasyluk</i>	344-362
• Decision-Making Methods and Models for Intelligent Business Analytics Systems in Online Tourism <i>Victoria Vyatska, Yevhen Burov, Volodymyr Hnatushenko, Lyubomyr Chyrun, Sofia Chyrun, Liliia Diakonuk, Liubov Kolyssa, Oksana Huhut, Oleh Naum</i>	363-386
• Comparative Analysis of DQN and PPO Algorithms in UAV Obstacle Avoidance 2D Simulation <i>Zoriana Rybachuk, Mykhailo Kopylov</i>	391-403
• Method of Determining the Text Sentiment by Thematic Rubrics <i>Anatoliy Sachenko, Taras Landuk, Khrystyna Lipanova-Honcharenko, Mackej Dobrowolski, Sergey Boguta, Leonid Bytsyura</i>	404-414
• Enhancing Network Security Through Vulnerability Analysis <i>Anatoliy Sachenko, Jakub Wolczyn, Serhiy Rymashchuk</i>	415-424

Method of Determining the Text Sentiment by Thematic Rubrics

Anatoliy Sachenko^{1,2}, Taras Lendiuk¹, Khrystyna Lipianina-Honcharenko¹, Maciej Dobrowolski², Gena Boguta¹ and Leonid Bytsyura¹

¹ West Ukrainian National University, Lvivska str., 11, Ternopil, 46000, Ukraine

² Kazimierz Pulaski University of Technology and Humanities in Radom, Department of Informatics, Jacek Malczewski str., 29, Radom, 26 600, Poland

Abstract

A method for determining the text sentiment of by thematic rubrics is proposed. It is based on a complex approach that integrates natural language processing, machine learning and linguistic analysis for automatic classification of text data. To implement the method, a generalized client-server architecture of the text sentiment for analysis system was developed and a set of data was collected from a wide range of article texts from the Internet sites of Ukraine, which ensure the representativeness of various styles, genres and topics. The high efficiency of the system in terms of classification of texts by rubrics was experimentally confirmed with a correspondence of 92%.

Keywords

text sentiment, thematic rubrics, natural language processing, machine learning, linguistic analysis, automatic classification, text data

1. Introduction

Currently, we are witnessing many information wars being waged on the world stage, with a special emphasis on the war in Ukraine. Russia uses information operations as a key element of its hybrid war against Ukraine and directs significant resources to the creation and dissemination of disinformation, the purpose of which is to influence public opinion, undermine trust in the Ukrainian state, its institutions and leaders, as well as to distort the real state of affairs in the international community

In the context of increasing information attacks from Russia, the importance of an accurate and objective analysis of the text sentiment of related to various aspects of the war and its impact on various spheres of life becomes extremely relevant. Automated analysis of the emotional coloring of texts can help detect attempts to manipulate public opinion, assess the general mood in society regarding certain events

¹COLINS-2024: 8th International Conference on Computational Linguistics and Intelligent Systems, April 12–13, 2024, Lviv, Ukraine

✉ as@wunu.edu.ua (A. Sachenko); tl@wunu.edu.ua (T. Lendiuk); xrustya.com@gmail.com (Kh. Lipianina-Honcharenko); m.dobrowolski@uthrad.pl (M. Dobrowolski); genaboguta7@gmail.com (G. Boguta); l.bytsyura@wunu.edu.ua (L. Bytsyura)

ORCID 0000-0002-0907-3682 (A. Sachenko); 0000-0001-9484-8333 (T. Lendiuk); 0000-0002-2441-6292 (Kh. Lipianina-Honcharenko); 0000-0003-0296-9651 (M. Dobrowolski); 0009-0000-9788-1753 (G. Boguta); 0000-0002-9476-011X (L. Bytsyura)



© 2024 Copyright for this paper by its authors.
Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

or policies, and identify changes in the information space that may indicate new directions of information attacks or disinformation campaigns.

The proposed method integrates advanced technologies of natural language processing (NLP), machine learning and linguistic analysis, which allows not only to automate the process of determining the emotional coloring of texts, but also to ensure high accuracy and objectivity of the obtained results. The application of such a complex approach in the conditions of information warfare opens up new opportunities for monitoring the information space, identifying and analyzing information operations, as well as for developing effective strategies for countering disinformation. Given the high dynamics of the information space and the constant change in the tactics and strategies of information warfare, the development and implementation of innovative text analysis methods that allow for prompt monitoring and analysis of attempts at manipulation and disinformation is extremely important.

Further materials are presented according to the following structure. In Chapter 2 the analysis of recent related works is fulfilled and in Chapter 3, the proposed method is described. Chapter 4 is dedicated to implementation and case studies, and Chapter 5 summarizes the received outcomes.

2. Related Work

The study of text tonality, which involves the integration of natural language processing (NLP), machine learning, and linguistic analysis, occupies an important place in modern scientific research. In [1], the analysis tool for text sentiment was developed based on the fusion of machine learning algorithms, demonstrating the effectiveness of combining different methods to more accurately determine the emotional coloring of texts. Authors [2] emphasized the importance of using TF-IDF and machine learning methods for sentiment analysis of Twitter data, confirming the importance of these techniques in the field of big data processing. In [3], the impact of the lexical richness of the training corpus on the performance of machine learning in analysis tasks was highlighted, emphasizing the need for a careful approach to data selection and processing. Authors [4] presented an innovative approach to the classification of hierarchical comments using BERT and a specialized naive Bayes classifier, opening up new opportunities for advanced text analysis. In [5], a deep learning-based approach to sentiment analysis is proposed for ranking online products using a set of probabilistic linguistic terms, which demonstrates the potential of using these technologies for commercial purposes. Authors [6] developed a technique for analyzing the sentiment of data from Twitter using NLP and machine learning, which emphasizes the importance of an integrated approach to information processing. The authors [7] demonstrate how sentiment analysis on Twitter can be improved using NLP and machine learning methods for emotion categorization and trend visualization. In [8], the advantages of integrating linguistic analysis for the study of human language are revealed, which opens up new opportunities for accurate thematic categorization. Research [9] focuses on the analysis of online product reviews using advanced deep learning and machine learning techniques [23, 24] to improve data classification and extract detailed emotion information [21].

Compared to the analogue [10], which focuses on the sentiment analysis of Twitter data, this work is distinguished by a deep integration of natural language processing (NLP) techniques, machine learning and linguistic analysis. In addition, a distinctive feature of the proposed approach is the development of specialized algorithms for accurate determination of sentiment taking into account contextual semantics and lexical diversity within specific thematic headings. This allows not only to more accurately classify the emotional coloring of texts, but also to detect subtle nuances in emotional expressions, which makes the proposed approach more adaptable to the complexities of natural language and provides a higher accuracy of analysis compared to existing methods.

3. Proposed Method

The proposed method for determining the sentiment of text by thematic headings is a comprehensive approach that integrates natural language processing (NLP), machine learning, and linguistic analysis for automatic classification of text data. The method allows to evaluate the emotional color of the text (positive, neutral or negative) and to express it quantitatively in the form of an expanded percentage scale in the range from -100% to +100%. Let us present this method as a set of the following stages and structurally (Fig. 1):

- Stage 1. Text pre-processing [11]:
 - 1.1. Removal of extra characters, text normalization;
 - 1.2. Detection and removal of stop words;
 - 1.3. Tokenization of text [11].
- Stage 2. Classification of the text by thematic headings [12]:
 - 2.1. The use of machine learning methods to determine the thematic rubric of the text [13];
 - 2.2. Training models on a predefined set of texts with a clear thematic affiliation.
- Stage 3. Creation and use of dictionaries for each thematic rubric [14]:
 - 3.1. Development of dictionaries with key words and phrases specific to each rubric;
 - 3.2. Determination of the emotional coloring of keywords (positive, negative, neutral).
- Stage 4. Sentiment analysis [15]:
 - 4.1. Using sentiment computation methods such as dictionary-based estimation or deep learning models;
 - 4.2. Calculation of the sentiment index T for the text based on the formula:

$$T = \frac{P - N}{P + N + Q} \times 100\%$$

where P is the number of positive words, N is the number of negative words, Q is the number of neutral words.

- Stage 5. Inversion of tonality. In the case when the text belongs to a hostile source and does not contain a direct mention of Ukraine, tonality inversion is used.
- Stage 6. Displaying the results. Development of an interface for visualizing the tonality of the text with the possibility of viewing a detailed analysis by thematic headings.

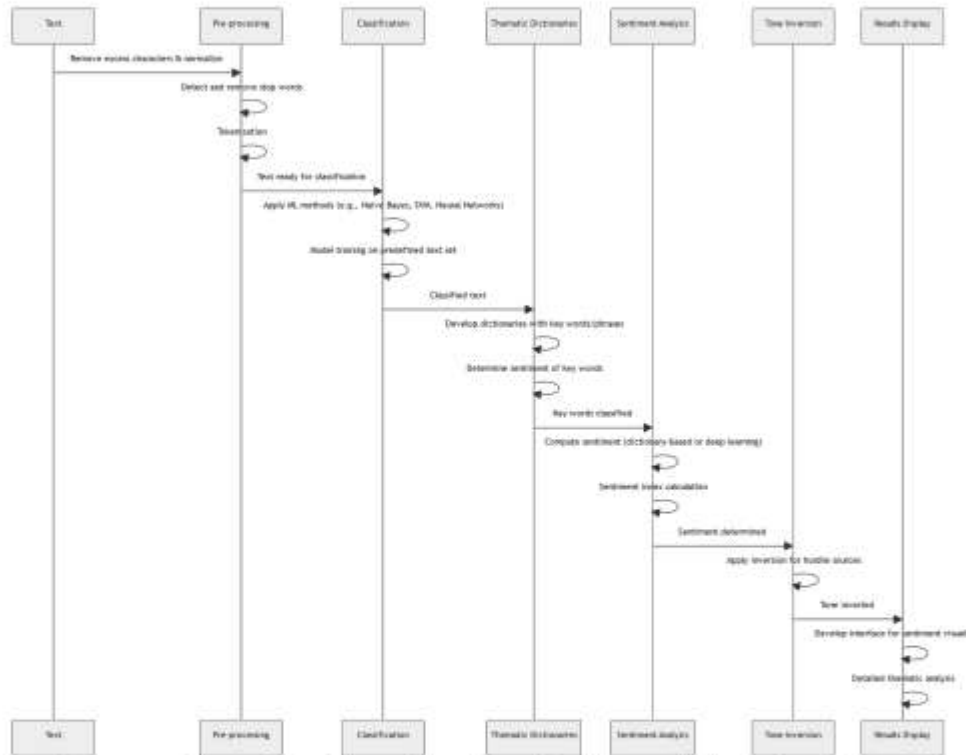


Figure 1: Structure of the method of determining the tonality of the text by thematic rubrics

4. Implementation and Case Study

The generalized client-server architecture of the text tonality analysis system is presented in Fig. 2. The client initiates the interaction by sending a request to the server, which in turn processes the received information. After processing, the server sends the necessary information back to the client. This process is cyclic, so after transferring information, the client can initiate interaction with the server again. The diagram shows a typical request-response model that is fundamental to a client-server architecture, where the server acts as a provider of resources and services and the client acts as a consumer of those resources and services.

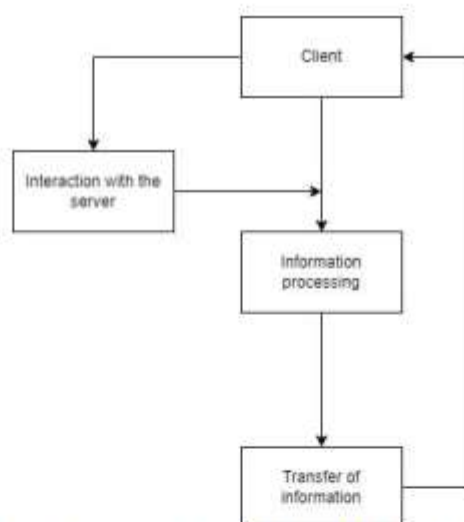


Figure 2: Generalized client-server architecture of the text tonality analysis system

To implement the proposed method, a set of data was collected from a wide range of texts of articles from Internet sites of Ukraine, which ensure the representativeness of various styles, genres and topics. The sample covers a variety of emotional contexts to test the algorithm's ability to accurately classify emotional nuances. The analysis of the text will be assigned to one of the following thematic headings, which were determined by experts from the security service of state bodies:

- Military and political leadership of Ukraine at all levels;
- Law enforcement agencies of Ukraine;
- Armed Forces of Ukraine;
- Socio-political situation in the regions of Ukraine (attitudes towards mobilization, socio-economic stability, etc.);
- Pro-Russian religious organizations on the territory of Ukraine;
- Pro-Russian movements, formation of the concept of "Russian peace";
- International image of Ukraine in the EU (English, German, Polish, Romanian, French, Hungarian, Ukrainian, Russian languages);
- International image of Ukraine in the USA, Canada and Great Britain (English language);
- International image of Ukraine in African countries (English, French, Arabic languages);
- International image of Ukraine in Asian countries (Chinese, Russian, Turkish, Arabic, Georgian, Kazakh, Farsi, Kyrgyz, Tajik, Uzbek, Japanese, Korean languages);
- Ukraine in the information space of the Russian Federation;
- Ukraine in the information space of the Republic of Belarus;

- Socio-economic, political, military situation in the Russian Federation (attitudes towards the top management, mobilization, deterioration of the economic situation, etc.).

The main functionality of the web interface (Fig. 3) [15] allows users to insert url links of news sites for analysis. After submitting the text, the system uses an algorithm to determine the emotional color, displaying the results in the form of an understandable report. The report includes quantitative indicators of the presence of positive and negative words, as well as an overall text tone index.

As a result of the analysis of the tonality of the text [16], the rubric "Military and political leadership of Ukraine at all levels" was determined, where a positive tone of +80% was indicated. This reflects a high level of positive emotional coloring of the text of the article. Details are given about military aspects, including the acquisition or use of military equipment (the F-16 and Bayraktar are mentioned), and an optimistic attitude towards events related to the war in Ukraine is reflected.

Enter url

<https://www.8352.ua/news/3738470/ukraina-moze-otrimati-vitricuvaci-f-16-u-personu-plevici-2024-ro>

☐ Inversion

Apply

Рубрика: Військово-політичне керівництво України всіх рівнів

Тональність: +80% (позитивна)

Текст статті:

Винищувачі F-16, які надасть Данія, з'являться у небі України до літа цього року, як про це прем'єр-міністр Данії Метте Фредеріксен заявила у Львові на спільній з Прем'єр-міністра України Радією Європи прес-конференції. Вона також заявила, що Данія збирається надати ці винищувачі у співпраці з Нідерландами. Додатково цитування матеріалів без отримання попередньої згоди 8352.ua за умови в матеріалах з плашками "Примо", "Партнерський матеріал", "Партнерський спонсор", "П

Figure 3: Example of text analysis result

The Inversion function (Fig. 4) of the results was included for users who wish to analyze the opposite emotional tone of the text. This can be useful, for example, when analyzing texts containing sarcasm or irony, where the literal meaning of the words can be misleading. The inversion allows you to quickly see how the overall assessment of emotional coloring will change if these stylistic figures are taken into account.

Enter url

<https://tsn.ua/ato/specsluzhbi-ri-vlashtuvai-masshtabnu-kampaniyu-poshiryuyut-feyki-vid-imeni-sirsk>

☒ Inversion

Apply

Figure 4: Field for entering a link and using inversion

One of the most important aspects of the web interface (Fig. 5a) is the ability to edit keywords for each heading. Users can add new words to dictionaries, remove existing ones, which affects their significance in the analysis algorithm. This allows you to adapt the system to the specific needs of the user and increase the accuracy of determining the tonality for specific topics or writing styles.

In addition, the web interface includes a dictionary management module (Fig. 5c), which allows users to view and update the database of words on the basis of which the analysis is performed. This is especially important to take into account linguistic updates, social and cultural changes that can affect language and its emotional load.

a – key words

b – dictionary of emotions

Figure 5: Example of editing keywords and dictionary

A comparative analysis was conducted to assess the effectiveness of the system for determining the tonality of texts (100 news stories) by comparing the automated results of the system with experts' assessments. The analysis took into account such parameters as the assignment of the text to the appropriate rubric by the system and experts, the comparison of the tonality of the texts determined by both the system and experts, and the correspondence of these assessments. Using the URL as a unique identifier for each text ensured accurate tracking of results. In addition to the quantitative evaluations, the experts provided additional notes for a deeper understanding of the reasons for the discrepancies between the expert evaluations and the results of the system, which will contribute to the further improvement of its accuracy and reliability.

Based on the results of filling out Table 1, statistical indicators were calculated, namely the simple percentage of matches between the system and experts.

Table 1
Evaluation of System Efficiency

Statistical indicator	indicator value
Correspondence of rubrics	92%
Average tonality deviation	15%

The obtained results (see Table 1) indicate a fairly high efficiency of the system in terms of classification of texts by headings with a correspondence of 92%. This means that in most cases the system correctly identifies the thematic category of the text, which indicates its reliability in determining the context of news. However, the average tonality deviation of 15% is quite significant and may indicate some shortcomings in the work of the algorithms of the system for evaluating the emotional coloring of the text. This may be due to the incompleteness of the dictionaries used for sentiment analysis. However, dictionaries can be constantly updated, which is a significant advantage of the system, since the language is constantly evolving, and the context and use of the vocabulary can change. The ability to supplement dictionaries by users allows the system to adapt more quickly to novelties in the language and changes in the use of terms, especially in the field of news, where it is important to take into account not only the lexical meaning of words, but also their connotative influence.

Thus, the system demonstrates a high accuracy in the classification of thematic headings, but needs improvement in determining tonality. Constant addition and updating of dictionaries, with the possibility of making changes from users, is an important process for increasing the accuracy of the analysis of the emotional coloring of texts.

5. Conclusion

A method for determining the text sentiment by thematic rubrics is proposed based on an integrated approach that integrates natural language processing, machine learning and linguistic analysis for automatic classification of text data. This allows you to assess the emotional color of the text (positive, neutral or negative) and express it quantitatively in the form of a percentage scale from -100% to +100%.

The text sentiment analysis system implemented based on the method has a high accuracy (92%) of classification of thematic headings. This means that in most cases the system correctly identifies the thematic category of the text, which indicates its reliability in determining the context of news.

However, the average sentiment deviation of 15% is quite significant and may indicate some shortcomings in the work of the algorithms of the system for evaluating

the emotional coloring of the text. This may be due to the incompleteness of the dictionaries used for sentiment analysis.

In the future, authors are going to explore the methods [17-20] for improving the quality and performance of text sentiment analysis.

References

- [1] P. Ajitha, A. Sivasangari, R. Rajkumar, & S. Poonguzhali, Design of text sentiment analysis tool using feature extraction based on fusing machine learning algorithms. *Journal of Intelligent & Fuzzy Systems* 40 (2021), 6375-6383 <https://dx.doi.org/10.3233/jifs-189478>
- [2] S. Singh, K. Kumar, & B. Kumar, Sentiment analysis of twitter data using TF-IDF and machine learning techniques, in: *Proceedings of the 2022 IEEE 6th International Conference on Communication and Electronics Systems (ICCES)*, 2022, pp. 252-255. <https://dx.doi.org/10.1109/com-it-con54601.2022.9850477>
- [3] S. Garg, A. Saini, & N. Khanna, Is sentiment analysis an art or a science? Impact of lexical richness in training corpus on machine learning. in: *Proceedings of the 2016 IEEE International Conference on Advances in Computing, Communications and Informatics (ICACCI)*, 2016, pp. 2729-2735. <https://dx.doi.org/10.1109/ICACCI.2016.7732474>
- [4] M. Dhina, & S. Sumathi, An innovative approach to classify hierarchical remarks with multi-class using BERT and customized naïve bayes classifier. *International Journal of Engineering, Science and Technology*, 13 (2022), 32-45. <https://dx.doi.org/10.4314/ijest.v13i4.4>
- [5] Z. Liu, H. Liao, M. Li, Q. Yang, & F. A Meng, Deep learning-based sentiment analysis approach for online product ranking with probabilistic linguistic term sets, *IEEE Transactions on Engineering Management* (2023). <https://dx.doi.org/10.1109/tem.2023.3271597>
- [6] K. Brindha, S. Senthilkumar, A. K. Singh, & P. M. Sharma, Sentiment analysis with NLP on Twitter data, in: *Proceedings of the IEEE International Conference on Smart Generation Computing, Communication and Networking (SMART GENCON)*, 2022, pp. 1-5. <https://dx.doi.org/10.1109/SMARTGENCON56628.2022.10084036>
- [7] K. Darshan, J. Samuel, M. Swamy, P. Koparde, & N. Shivashankara, NLP - Powered sentiment analysis on the Twitter, *Saudi Journal of Engineering and Technology* 9, (2024) 1-11. <https://dx.doi.org/10.36348/sjet.2024.v09i01.001>
- [8] P. Srinivas, K. Gayathri, K. Bhavitha, Jahnavi & K. D. Sarath, BLIP-NLP model for sentiment analysis, in: *Proceedings of the 2023 2nd International Conference on Edge Computing and Applications (ICECAA)*, 2023, pp. 468-475. IEEE <https://dx.doi.org/10.1109/ICECAA58104.2023.10212253>
- [9] L. Bharadwaj, Sentiment analysis in online product reviews: mining customer opinions for sentiment classification 5, (2023). <https://dx.doi.org/10.36948/ijfmr.2023.v05i05.6090>

- [10] S. Voloshyn, V. Vysotska, O. Markiv, I. Dyyak, I. Budz and V. Schuchmann, Sentiment analysis technology of English newspapers quotes based on neural network as public opinion influences identification tool, in: Proceedings of the 2022 IEEE 17th International Conference on Computer Sciences and Information Technologies (CSIT), 2022, pp. 83-88, doi: 10.1109/CSIT56902.2022.10000627.
- [11] R. Gramyak, H. Lipyana-Goncharenko, A. Sachenko, T. Lendyuk, & D. Zahorodnia, Intelligent Method of a competitive product choosing based on the emotional feedbacks coloring, in: Proceedings of the IntelITSIS, 2021, pp. 246-257.
- [12] H. Lipyana, S. Sachenko, T. Lendyuk, & A. Sachenko, Targeting model of HEI video marketing based on classification tree, in: Proceedings of the ICTERI Workshops, October 2020, pp. 487-498.
- [13] H. Lipyana, S. Sachenko, T. Lendyuk, V. Brych, V. Yatskiv, & O. Osolinskiy, Method of detecting a fictitious company on the machine learning base, in: International Conference on Computer Science, Engineering and Education Applications, January, 2021, pp. 138-146. Cham: Springer International Publishing. https://doi.org/10.1007/978-3-030-80472-5_12
- [14] K. Lipianina-Honcharenko, T. Lendiuk, A. Sachenko, O. Osolinskiy, D. Zahorodnia, & M. Komar, An intelligent method for forming the advertising content of higher education institutions based on semantic analysis, in: International Conference on Information and Communication Technologies in Education, Research, and Industrial Applications, September 2021, pp. 169-182. Cham: Springer International Publishing. https://doi.org/10.1007/978-3-031-14841-5_11
- [15] Streamlit. URL: <https://nelczgkwcsghtwdkw7buxn.streamlit.app/>
- [16] Ukraine can receive F-16 fighter jets in the first half of 2024, - the premier of Denmark, 2024. URL: <https://www.0352.ua/news/3738470/ukraina-moze-otrimativinisivaci-f-16-u-persomu-pivricci-2024-roku-premerka-danii>
- [17] M. Komar, V. Golovko, A. Sachenko and S. Bezobrazov, Development of neural network immune detectors for computer attacks recognition and classification, in: Proceedings of the 2013 IEEE 7th International Conference on Intelligent Data Acquisition and Advanced Computing Systems (IDAACS), Berlin, Germany, 2013, pp. 665-668, doi: 10.1109/IDAACS.2013.6663008.
- [18] Zhengbing Hu, Yevgeniy V. Bodyanskiy, Nonna Ye. Kulishova, Oleksii K. Tyshchenko, A multidimensional extended neo-fuzzy neuron for facial expression recognition, International Journal of Intelligent Systems and Applications (IJISA) 9 (2017) 29-36. DOI: 10.5815/ijisa.2017.09.04
- [19] V. Vysotska, O. Markiv, S. Tchynetskyi, B. Polishchuk, O. Bratasyuk, V. Panasyuk, Sentiment analysis of information space as feedback of target audience for regional e-business support in Ukraine, in: CEUR Workshop Proceedings, vol-3426, 2023, 488-513.
- [20] Sutriawan, S., P. N. Andono, M. Muljono, & R. A. Pramunendar, Performance evaluation of classification algorithm for movie review sentiment analysis,

- [21] S. Voloshyn, O. Markiv, V. Vysotska, I. Dyyak, L. Chyrun and V. Panasyuk, Emotion recognition system project of English newspapers to regional e-business adaptation, in: Proceedings of the 2022 IEEE 17th International Conference on Computer Sciences and Information Technologies (CSIT), 2022, pp. 392-397, doi: 10.1109/CSIT56902.2022.10000527.
- [22] K. Mehta, & S. P. Panda, Sentiment analysis on e-commerce apparels using convolutional neural network, International Journal of Computing, vol. 21, issue 2, pp. 234-241, 2022. <https://doi.org/10.47839/ijc.21.2.2592>
- [23] R. Lynnyk, V. Vysotska, Y. Matseliukh, Y. Burov, L. Demkiv, A. Zaverbnyj, A. Sachenko, I. Shylinska, I. Yevseyeva, and O. Bihun, DDOS attacks analysis based on machine learning in challenges of global changes, in: 2020 CEUR Workshop Proceedings, 2020 vol. 2631, pp. 159-171.
- [24] O. Soprun, M. Bublyk, Y. Matseliukh, V. Andrunyk, L. Chyrun, I. Dyyak, A. Yakovlev, M. Emmerich, O. Osolinsky, A. Sachenko, Forecasting temperatures of a synchronous motor with permanent magnets using machine learning, in: 2020 CEUR Workshop Proceedings, 2020, vol. 2631, pp. 95-120.

An Intelligent Method for Recognizing Real and Generated Ukrainian-Language Voices

Khrystyna Lipianina-Honcharenko
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
xrustyia.com@gmail.com

Dmytro Lendiuk
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
dmytrolendiuk@gmail.com

Adam Ivaniush
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
adam.ivaniush@gmail.com

Gena Boguta
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
genaboguta7@gmail.com

Mariana Soia
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
m.soya@wunu.edu.ua

Khrystyna Yurkiv
Department for Information-Computing
Systems and Control
West Ukrainian National University
Ternopil, Ukraine
kh.yurkiv@wunu.edu.ua

Abstract—The modern development of speech synthesis technologies and convolutional neural networks (CNN) creates new opportunities and challenges in the field of voice recognition. One of the key tasks is to ensure the security and authenticity of audio information, especially in the context of information warfare and disinformation campaigns. In this context, the ability to recognize and identify generated voices becomes particularly important. This study aims to develop an intelligent method for recognizing real and generated Ukrainian-language voices using convolutional neural networks. The proposed method is based on the careful collection and preparation of data, the development and training of a CNN model, and its validation on test samples. Special attention is paid to ensuring the balance of the dataset, which is crucial for preventing model bias and improving its ability to generalize to unknown data. The model shows high performance in detecting real and generated voices, making this approach useful for local applications, particularly in Ukraine.

Keywords—voice recognition, generated voices, convolutional neural networks, Ukrainian language

INTRODUCTION

Modern development of speech synthesis and convolutional neural networks (CNN) technologies creates new opportunities and challenges in the field of voice recognition. On the one hand, these technologies contribute to improving the quality of synthesized audio recordings, which can be used in a wide range of applications – from virtual assistants to the generation of voice content. On the other hand, increasing the realism of generated voices poses new challenges for researchers working to ensure the safety and authenticity of audio information.

In particular, in the context of information warfare and disinformation campaigns, the ability to recognize and identify generated voices becomes especially relevant. Russia's war of aggression against Ukraine has highlighted the need for reliable tools to detect fake audio materials that can be used to manipulate public opinion, spread false information, and undermine trust in official sources.

The purpose of this study is to develop an intelligent method for recognizing real and generated Ukrainian voices using convolutional neural networks. The proposed method is based on careful data collection and preparation, development and training of the CNN model, as well as its validation on test samples.

Particular attention is paid to ensuring the balance of the dataset, which is key to preventing model bias and increasing its ability to generalize to unknown data.

This work is divided as follows. Section 2 deals with the analysis of related works, section 3 presents an intelligent method for recognizing real and generated Ukrainian-speaking voices. Section 4 presents the implementation of the method. Chapter 5 presents the conclusions to the study.

RELATED WORKS

The modern development of speech synthesis and deep learning technologies creates new opportunities and challenges in the field of voice recognition. Synthesized voices are becoming more and more realistic, which increases the importance of effective methods for their detection and identification. In particular, the ability to distinguish real voices from synthetic ones is critical for ensuring information security, especially in the context of disinformation campaigns and fake news.

Research [1] presents a FoR dataset containing more than 198,000 utterances from modern speech synthesizers and real language. It also describes the use of this dataset to train deep learning models that have achieved 99.96% accuracy in detecting synthetic language.

The study [2] describes the methods of synthetic language detection using deep neural networks, in particular the analysis of frequency characteristics of the language for training classifiers. The authors achieved high accuracy using dynamic acoustic features.

Research [3] analyzes the reliability of voice recognition systems using deep learning algorithms and convolutional neural networks to improve voice identification accuracy and security.

Research [4] proposes a solution to detect fake voice recordings created with Deep Voice and Imitation technologies using convolutional neural networks for high-accuracy audio classification.

The study [5] examines the ability of two-dimensional convolutional neural networks to generalize and detect false speech, achieving high performance even on new synthetic language systems.

Research [6] presents a methodology for detecting cough sounds using deep neural networks, where high accuracy was achieved in the classification of cough sounds, which can be applied to the diagnosis of diseases.

Research [7] analyzes the use of convolutional neural networks for remote language recognition in a large vocabulary. The authors achieved a significant improvement in the accuracy of voice recognition in noisy environments.

The study [8] presents an ensemble of convolutional neural networks for audio classification. Using this approach, the authors achieved high accuracy in the classification of various types of audio signals.

Research [9] analyzes the use of convolutional neural networks for voice emotion recognition. The authors have achieved high accuracy in determining emotional states, which makes this technology useful for many applications.

Below is a comparative Table 1, which contains the main characteristics and features of four studies related to the recognition of real and generated voices using convolutional neural networks.

TABLE 1. COMPARISON OF CLOSEST ANALOGUES

# of reference	Accuracy	Peculiarities of the study	Used language
[1]	99.96%	For dataset containing over 198,000 utterances. Using deep learning models for synthetic language detection.	English, synthetic
[2]	99.63%	Methods of detecting synthetic language using deep neural networks. Analysis of frequency characteristics of speech for training classifiers.	English
[3]	98.5%	Detection of fake voice recordings created using Deep Voice and imitation technologies. Using convolutional neural networks for audio classification.	English, synthetic

Therefore, this study differs from analogues (see Table 1) in the field of recognition of real and generated voices using convolutional neural networks in that it focuses on Ukrainian-speaking voices, using modern text-to-speech technologies to create generated recordings. The developed model achieves high accuracy in detecting real and generated voices due to the relevance and diversity of the dataset, which includes both real and generated Ukrainian voices. This makes this work unique and useful for local applications, providing new opportunities for the development of speech recognition technologies in Ukraine.

METHODS AND MATERIALS

Below is a detailed description of the intelligent method for recognizing real and generated Ukrainian voices, developed using convolutional neural networks. The method is presented in the form of step-by-step instructions, each step of which describes a separate aspect of the process and structure (Fig. 1).

Stage 1. Data collection. To develop and train a voice classification model, two main sources of data must be identified. An appropriate number of records was selected for each source, ensuring sufficient representativeness and diversity within the dataset. This allows for a balanced data

set, which is important to avoid biasing the model and improve its ability to generalize to unknown data.

1.1 Set of real Ukrainian-speaking voices. A dataset that contains recordings of real voices collected from open sources.

1.2 Set of generated Ukrainian voices. Created dataset including generated voices created using modern text-to-speech technologies.

1.2.1. Downloading and cleaning the news dataset. Downloading data from news sources can be described as a selection function D that selects a subset of texts T from the entire population of available texts N :

$$D(T) \subseteq N.$$

Cleaning texts from irrelevant elements such as links and advertising blocks can be expressed as a cleaning function C that transforms each text $t \in T$ into a clean text t' :

$$C(t) = t' \quad \forall t \in T,$$

where t' contains no unwanted elements.

1.2.2. Audio generation. The text-to-audio conversion process can be described using a synthesis function S , which converts plain text t' to audio a :

$$S(t') = a,$$

where S uses speech synthesis algorithms to create natural-sounding audio recordings.

1.2.3. Save audio. Saving audio recordings can be represented as a save function R that maps each audio recording a to a file system for later use:

$$R(a) = \text{audio file}.$$

This mathematical approach allows formalizing the process of data preparation for audio generation and ensures clarity and consistency in data manipulation, ensuring that each stage can be accurately reproduced and analyzed in scientific research.

Stage 2. Data transformation. A spectrogram is obtained from each audio file, using the short-time Fourier transform [10], which is described by the following formula:

$$S(f, t) = \int x(\tau) \omega(\tau - t) e^{-j2\pi f \tau} d\tau,$$

where:

- $x(\tau)$ – signal in time,
- $\omega(\tau - t)$ is a window function focusing the analysis around time t ,
- f – frequency,
- $S(f, t)$ is the complex amplitude value in frequency f and time t .

Stage 3 Preparation of training. All spectrograms must be scaled to the same size using a bilinear interpolation procedure that preserves key properties of the data during

resizing. This process can be expressed by the following mathematical formula [10]:

$$S'(u, v) = \sum_{t=1}^T \sum_{f=1}^F S(t, f) \cdot \text{interp}(u, v, t, f),$$

where:

- $S(t, f)$ – original spectrogram values,
- $S'(u, v)$ is the new value in the pixel of the resulting spectrogram, (u, v) ,
- interp is a bilinear interpolation function that takes into account the weight of neighboring pixels.

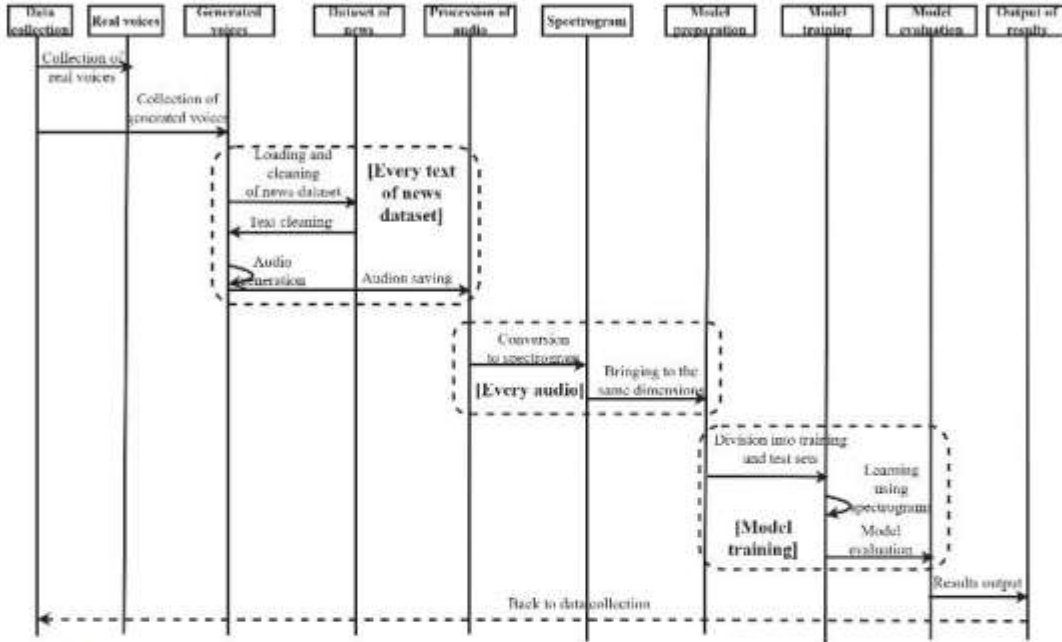


Fig. 1. The structure of the intelligent method of recognizing real and generated Ukrainian voices

Stage 4. Distribution of data. The next step is to divide the data into training and test sets in the proportion of 80% to 20%, respectively. In the training set, each spectrogram is labeled as a representation of a real or generated voice. This distribution allows to provide an adequate amount of data for effective training and validation of the model, reducing the possibility of its retraining.

Stage 5. Model development. For classification, a convolutional neural network [11] will be developed, which includes several convolutional layers designed to detect textural and structural features in spectrograms. The dimensions of the input layer of the network will correspond to the dimensions of the spectrograms. As an activation function in convolutional layers, ReLU is used, which is described as follows:

$$\text{ReLU}(x) = \max(0, x).$$

Stage 6. Model compilation and training. The next step is to compile the model using the Adam optimizer [12], which effectively minimizes the loss function. For this task, binary cross-entropy is used, designed to work with two classes, which is described as follows:

$$L(y, \hat{y}) = - \sum_{i=1}^C y_i \log(\hat{y}_i),$$

where:

- y – true class label,
- $\hat{y}$ – predicted probability for the class,
- C – number of classes (in this case, real or generated). $C = 2$

Stage 7. Evaluation of the model. The model is tested on a separate test sample, which contains an equal number of real and generated audio recordings. The accuracy of the model (accuracy [13]) is calculated using the following formula:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

where:

- TP (True Positives) – the number of correctly identified generated votes,

- TN (True Negatives) – the number of correctly identified real voices,
- FP (False Positives) – number of generated votes falsely identified as real,
- FN (False Negatives) – number of real votes falsely identified as generated.

In addition, indicators such as sensitivity and specificity can also be calculated, which are important for assessing the balance between the detection of real and generated voices:

$$\text{Sensitivity} = TP / TP + FN,$$

$$\text{Specificity} = TN / TN + FP.$$

Stage 8: Output of results. In the final step, the model is used to determine whether each audio recording is generated or real. Using the processed spectrograms as input, the model classifies the audio recordings and outputs results that are displayed through the user interface or stored for further analysis. This process allows us to evaluate the model's ability to accurately distinguish between real and generated voices, which is critical for its practical application in areas where the accuracy of audio authentication is key.

REALIZATION

Two main sources of data were used to implement the proposed method [14]. The first is a text dataset used to generate text-to-speech (TTS) generated audio that simulates human speech with different intonations and accents. The second source is the common_voice_16_1 audio dataset from the Mozilla Foundation, which includes audio recordings of real people exhibiting a wide range of vocal characteristics.

The collected datasets were subjected to thorough processing in order to create a single consolidated data set. The first stage of processing consisted in the standardization of audio file formats, which were unified to the WAV format with the same sampling rate and bit depth parameters. Each audio file has been labeled to indicate its origin: "generated" or "real".

The final stage of data preparation involved balancing the data set to ensure equal representation of generated and real records. This step is critical to prevent bias in the model and improve its ability to generalize to new, unknown data. By using carefully selected and prepared data, the dataset has become a valid tool for effective training and testing of deep learning algorithms aimed at audio recognition.

To ensure the clarity of the data preparation process and its impact on the training of the neural network, a visualization was developed that displays the relationship between the generated and real audio recordings in the training set. As shown in Fig. 2, the visualization includes two types of audio: generated using TTS technologies and real audio obtained from Mozilla's common_voice_16_1 dataset. The presented visualization illustrates the quantitative distribution of audio recording classification results. Two columns are shown: the first column in orange represents the number of audio tracks that were correctly identified as generated ("true") and has approximately 11,146 units. The second column in green shows the number of audio recordings that were falsely classified as not generated ("false"), showing a similar number of about 10,015 units. This mapping allows you to evaluate

the balance of the data set and the classification efficiency of the model.

The process of preparing audio data for machine learning begins by downloading the necessary libraries, such as Librosa and SciPy, which provide tools for reading, analyzing and visualizing audio data. Each audio file is converted to a spectrogram using a short-time Fourier transform (STFT), which allows visualization of the frequency changes of the sound over time, followed by normalization of the spectrogram dimensions to 128x128 pixels via bilinear interpolation to standardize the input data. These spectrograms, used as input to a convolutional neural network, allow the model to analyze patterns in audio signals, greatly increasing the accuracy of recognizing audio as real or generated, which is key to successful model training and validation.

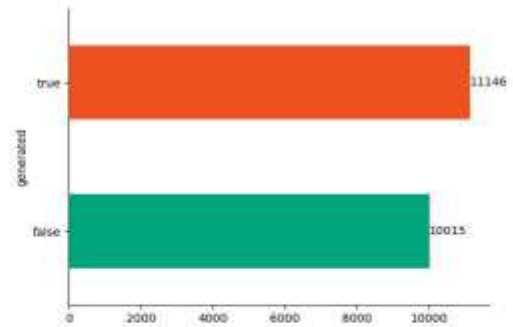


Fig. 2. Distribution of classification results of generated audio recordings

To recognize real and generated Ukrainian voices, a deep neural network was created using the TensorFlow library and the Keras API. The architecture of the model includes several key components: a sequence of convolutional layers (Conv2D), which serve to detect the main patterns in spectrograms, maximum pooling layers (MaxPooling2D), which reduce the dimensionality of the original data for more efficient processing, and fully connected layers (Dense), which allow networks to study nonlinear dependencies between patterns. Dropout layers are integrated between core layers to prevent overtraining, and activation of all hidden layers is done using ReLU. The final layer of the model uses a sigmoidal activation function to determine the probability that the audio belongs to the category of the generated records.

The model was compiled using the Adam optimizer, which is optimal for a wide range of deep learning problems due to its ability to adapt the learning rate based on the computed gradients. Binary cross-entropy is used as the loss function, which is the standard for binary classification.

To train the audio classification model, the total dataset was divided into training and validation samples in an 80/20 ratio, which provided a balanced training base and sufficient data to test the model's performance. Training was carried out for 10 epochs, which allowed the model to effectively adapt to the task of recognizing generated and real audio without the risk of significant overtraining.

After completing the training process of the model, its effectiveness was evaluated on a separate test and validation sample. Using the error matrix shown in Fig. 3, a quantitative analysis of the classification results was performed.

According to the matrix, the model correctly identified 11,130 real audio recordings and 9,964 generated audio recordings, while getting it wrong only 16 times for real and 51 times for generated audio recordings.

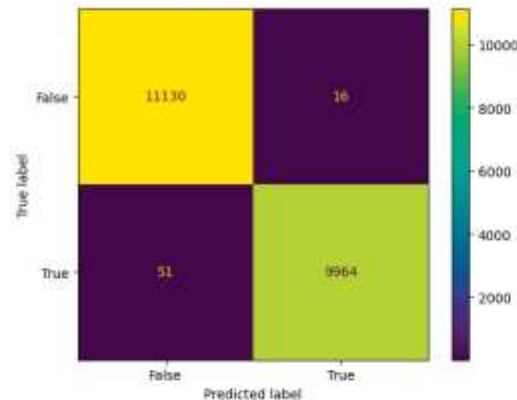


Fig. 3. Error matrix for evaluating the classification of audio recordings by the model

At the final evaluation stage, the model demonstrates impressive performance indicators, where the classification accuracy reaches a value of 0.9968. This means that the model successfully identified 99.68% of cases correctly on the validation and test samples, effectively distinguishing between real and generated audio recordings.

A. Case Study: An Analysis of Practical Model Validation on Audio from TikTok and YouTube Shorts

At the final stage of practical testing of the model's ability to distinguish between real and generated voices, three audio recordings from the TikTok and YouTube Shorts platforms were analyzed. The test results presented in Fig. 4 illustrate the high accuracy of the model in determining natural human speech and its interaction with background noise.

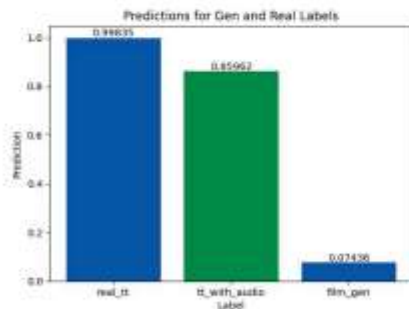


Fig. 4. Visualization of model predictions for audio recordings from TikTok and YouTube Shorts

Real tt
https://figshare.com/articles/software/audio_database/25451923 and
 Tt with audio
https://www.tiktok.com/@k0sta4kaa/video/7370695956372802822?is_from_webapp=1&sender_device=pc clips showed high prediction accuracy approaching 1.0, confirming the effectiveness of the model in recognizing real voices, including those with background noise. At the same time, the

Film_Gen
<https://www.tiktok.com/@vladasoloviova/video/7371392969301101829>, with a score of about 0.07, turned out to be identified as generated, since the score is close to 0, which indicates a positive accuracy of identifying a generated voice on a given clip.

The following model showed a high ability to recognize real voices, as demonstrated by the results for the real_tt and tt with audio samples with accuracies of 0.998353 and 0.859620, respectively. This testifies to the reliability of the model in recognizing natural human speech, and even in the presence of additional sound effects it performs quite well. The result for the film_gen sample with an indicator of 0.074357 indicates a significant accuracy of the model in classifying the generated voice in the clip provided for the test. This result can be interpreted as the imperfection of voice detection with background noise such as music, which is an indication of the need for further improvement.

CONCLUSIONS

This study presented a new approach to recognizing real and generated Ukrainian voices using convolutional neural networks (CNN). The main stages of development and implementation of the method included data collection, their transformation and preparation, model development, their training and evaluation. The use of mathematical formalizations made it possible to clearly outline each step of the process, which contributed to obtaining high-quality results.

The results of the experiments demonstrate that the developed CNN model achieves high accuracy in voice recognition. Using the Adam optimizer and the binary cross-entropy loss function ensured efficient model training. Validation of the model on a separate test sample showed a classification accuracy of 99.68%, which indicates a high ability of the model to distinguish between real and generated audio recordings.

Analysis of audio from the TikTok platform confirmed the effectiveness of the model in real conditions. The model successfully recognized natural human speech even in the presence of background noise. However, some challenges remain, especially regarding the recognition of modern synthesized voices that can mimic real human speech.

Thus, the proposed approach has significant potential for use in various fields where the accuracy of audio authentication is key. Further research will be aimed at studying the use of generated votes for fake research, which is especially relevant in the context of Russia's aggression against Ukraine. This will improve tools for combating disinformation and provide more reliable methods of authenticating audio materials in the difficult conditions of information warfare (<https://www.youtube.com/shorts/z6XLjrGr7PQ>).

REFERENCES

- [1] R. Reimao, & V. Tzerpos, *For: A Dataset for Synthetic Speech Detection*, York University, 2020. [Online]. Available at: https://bil.eecs.yorku.ca/wp-content/uploads/2020/01/For-Dataset_RR_VT_final.pdf
- [2] R. Reimao, *Synthetic Speech Detection Using Deep Neural Networks*, Master Thesis, York University, Canada. [Online]. Available at: <https://core.ac.uk/download/pdf/240138805.pdf>
- [3] W. Ibrahim, H. Candia, & H. Isyanto, "Voice recognition security reliability analysis using deep learning convolutional neural network

- algorithm," *Journal of Electrical Technology UMY*, vol. 6, issue 1, pp. 1-11, 2022. <https://doi.org/10.18196/jet.v6i1.14281>
- [4] D. M. Ballesteros, Y. Rodríguez-Ortega, D. Renza, & G. Arce, "Deep4SNet: deep learning for fake speech classification," *Expert Systems with Applications*, vol. 184, 115465, 2021. <https://doi.org/10.1016/j.eswa.2021.115465>
- [5] C. Papastergiopoulos, A. Vafeiadis, I. Papadimitriou, K. Votis, & D. Tzovaras, "On the generalizability of two-dimensional convolutional neural networks for fake speech detection," *Proceedings of the 1st International Workshop on Multimedia AI against Disinformation MAD'22*, 2022, pp. 3-9. <https://doi.org/10.1145/3512732.3533585>
- [6] J. Amoh, & K. Odame, "Deep neural networks for identifying cough sounds," *IEEE Transactions on Biomedical Circuits and Systems*, vol. 10, Issue 5, pp. 1003-1011, 2016. <https://doi.org/10.1109/TBCAS.2016.2598794>
- [7] P. Swietojanski, A. Ghoshal, & S. Renals, "Convolutional neural networks for distant speech recognition," *IEEE Signal Processing Letters*, vol. 21, issue 9, pp. 1120-1124, 2014. doi: <https://doi.org/10.1109/LSP.2014.2325781>
- [8] L. Nammi, G. Maguolo, S. Brahnam, & M. Paci, "An ensemble of convolutional neural networks for audio classification," *Applied Sciences*, vol. 11, issue 13, 5796, 2021. <https://doi.org/10.3390/app11135796>
- [9] A. B. A. Qayyum, A. Arefeen, & C. Shahnaz, "Convolutional neural network (CNN) based speech-emotion recognition," *Proceedings of the 2019 IEEE International Conference on Signal Processing, Information, Communication & Systems (SPICSCON'2019)*, November 2019, pp. 122-125. <https://doi.org/10.1109/SPICSCON48833.2019.9065172>
- [10] A. V. Oppenheim, "Speech spectrograms using the fast Fourier transform," *IEEE Spectrum*, vol. 7, issue 8, pp. 57-62, 1970. <https://doi.org/10.1109/MSPEC.1970.5213512>
- [11] V. Golovko, A. Kroschanka, E. Mikhno, M. Komar, A. Sachenko, "Deep convolutional neural network for detection of solar panels," In: Radivilova, T., Ageyev, D., Kryvinska, N. (eds) *Data-Centric Business and Applications. Lecture Notes on Data Engineering and Communications Technologies*, vol. 48, 2021, pp. 371-389. Springer, Cham. https://doi.org/10.1007/978-3-030-43070-2_17
- [12] Z. Zhang, "Improved Adam optimizer for deep neural networks," *Proceedings of the 2018 IEEE/ACM 26th International Symposium on Quality of Service (IWQoS)*, June 2018, pp. 1-2. <https://doi.org/10.1109/IWQoS.2018.8624183>
- [13] K. Lipianina-Honcharenko, C. Wolff, A. Sachenko, I. Kit, D. Zahorodnia, "Intelligent method for classifying the level of anthropogenic disasters," *Big Data and Cognitive Computing*, vol. 7, issue 3, 157, 2023. <https://doi.org/10.3390/bdcc7030157>
- [14] H. Lipyanina, S. Sachenko, T. Lendyuk, V. Brych, V. Yatskiv, O. Osolinskiy, "Method of detecting a fictitious company on the machine learning base," In: Hu, Z., Petoukhov, S., Dyehka, I., He, M. (eds) *Advances in Computer Science for Engineering and Education IV. ICCSEEA 2021. Lecture Notes on Data Engineering and Communications Technologies*, vol. 83, 2021, pp. 138-146. Springer, Cham. https://doi.org/10.1007/978-3-030-80472-5_12