

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

Мукан Павло Ігорович

**Модуль класифікації емоційного забарвлення тексту з
використанням бібліотеки Scikit-learn / Module for classifying
emotional coloring of text using the Scikit-learn library**

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

Кваліфікаційна робота

Виконав студент групи КН-41
Мукан П.І.

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

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

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

Тернопіль – 2025

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

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

ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ
Мукан Павло Ігорович
(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи: Модуль класифікації емоційного забарвлення тексту з використанням бібліотеки Scikit-learn / Module for classifying emotional coloring of text using the Scikit-learn library

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

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

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

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

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

4. Основні питання, які потрібно розробити:

- Провести аналіз предметної області та огляд сучасних підходів до виявлення мови ворожнечі.
- Сформулювати вимоги до системи класифікації та обґрунтувати вибір моделей.
- Реалізувати алгоритми попередньої обробки текстових даних.
- Провести навчання і тестування моделей (Decision Tree, Logistic Regression) з використанням векторизації тексту.
- Оцінити ефективність побудованих моделей за метриками точності, повноти та F1.
- Продемонструвати приклади застосування класифікатора на нових текстах..

5. Перелік графічного матеріалу в роботі:

- UML-діаграма компонентної архітектури модуля класифікації мови ворожнечі.
- Матриця невідповідностей для моделі Decision Tree..

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 _____ Муқан П.І.
(підпис) (прізвище та ініціали)

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

АНОТАЦІЯ

Кваліфікаційна робота на тему «Модуль класифікації емоційного забарвлення тексту з використанням бібліотеки Scikit-learn» на здобуття освітнього ступеня «бакалавр» за спеціальністю 122 «Комп'ютерні науки» освітньої програми «Комп'ютерні науки» написана обсягом 60 сторінки та містить 13 ілюстрацій, 7 таблиць, 3 додатки і 25 використані джерела .

Метою роботи є розробка та реалізація програмного модуля автоматизованого виявлення мови ворожнечі в текстах із використанням алгоритмів класифікації бібліотеки Scikit-learn.

Методи дослідження: інтелектуальний аналіз даних (очищення, стемінг, видалення стоп-слів), векторизація (Bag-of-Words, TF-IDF), машинне навчання (Decision Tree, Logistic Regression), оцінювання якості (accuracy, precision, recall, F1-score, confusion matrix).

За результатами реалізовано класифікаційний модуль, що демонструє точність до 93.89% для Logistic Regression. Проведено апробацію результатів на конференції ІІТАР-2025.

Ключові слова: **МОВА ВОРОЖНЕЧІ, МАШИННЕ НАВЧАННЯ, КЛАСИФІКАЦІЯ ТЕКСТУ, SCIKIT-LEARN, ВЕКТОРИЗАЦІЯ, DECISION TREE, LOGISTIC REGRESSION.**

ANNOTATION

The qualification thesis “Module for classifying emotional coloring of text using the Scikit-learn library”, submitted for the Bachelor’s degree in specialty 122 “Computer Science” (study program “Computer Science”), comprises 60 pages, 13 figures, 7 tables, 3 appendices and 25 references.

The aim of the thesis is to develop and implement an automated module for detecting hate speech in texts using classification algorithms provided by the Scikit-learn library.

Research methods: data mining (cleaning, stemming, stop-word removal), vectorization (Bag-of-Words, TF-IDF), machine learning (Decision Tree, Logistic Regression), and quality assessment (accuracy, precision, recall, F1-score, confusion matrix).

As a result, a classification module was implemented, achieving up to 93.89% accuracy using Logistic Regression. The results were presented at the IITAR-2025 student research conference.

Keywords: HATE SPEECH, MACHINE LEARNING, TEXT CLASSIFICATION, SCIKIT-LEARN, VECTORIZATION, DECISION TREE, LOGISTIC REGRESSION.

ЗМІСТ

Вступ	7
1. Аналіз предметної області та огляд існуючих рішень	10
1.1 Поняття та сутність мови ворожнечі	10
1.2 Огляд існуючих підходів до виявлення мови ворожнечі	13
1.3 Постановка задачі	18
2. Проектування модуля та методи	20
2.1 Архітектура модуля класифікації мови ворожнечі	20
2.2 Алгоритм попередньої обробки текстових даних	21
2.3 Алгоритми машинного навчання класифікації мови ворожнечі	23
2.4 Алгоритми оцінювання моделей класифікації	27
3. Реалізація модуля класифікації мови ворожнечі	31
3.1 Реалізація попередньої обробки текстових даних	31
3.2 Формування навчальної та тестової вибірок	33
3.3 Побудова моделі класифікації мови ворожнечі	34
3.4 Аналіз результатів	36
3.5 Приклади роботи модуля класифікації мови ворожнечі	39
Висновки	44
Список використаних джерел	46
Додаток А Програмний код	49
Додаток Б Візуалізація результатів	53
Додаток В Апробація отриманих результатів	55

ВСТУП

Модуль класифікації емоційного забарвлення тексту на базі бібліотеки Scikit-learn передбачає низку підсистем, і однією з ключових є автоматичне виявлення мови ворожнечі. У цифрову епоху мова ворожнечі стала невід’ємним елементом онлайн-комунікації, особливо в соціальних мережах та відкритих платформах обміну думками. Її поширення негативно впливає на соціальну атмосферу, спричиняє конфлікти, загострює міжетнічну й політичну напругу, а в окремих випадках може слугувати каталізатором реального насильства. Враховуючи зростання обсягів інформаційного потоку, забезпечити ефективний контроль за таким контентом вручну є практично неможливо.

Це зумовлює високу актуальність розробки автоматизованих систем для виявлення мови ворожнечі у текстових повідомленнях. Надійні інструменти класифікації подібного контенту є критично важливими для інтернет-платформ, освітніх установ, ЗМІ та правозахисних організацій. Вони мають не лише фільтрувати агресивну лексику, а й забезпечувати аналітичну підтримку в протидії інформаційним загрозам.

Незважаючи на значні успіхи в галузі глибинного навчання, класичні моделі машинного навчання залишаються актуальними завдяки своїй інтерпретованості, низьким обчислювальним вимогам та стабільності. Бібліотека Scikit-learn надає широкий спектр таких моделей, а також інструменти для повного циклу побудови класифікатора — від обробки даних до оцінювання результатів, що робить її оптимальним вибором для реалізації задач середньої складності.

Особливо важливим є створення моделей, здатних працювати на коротких текстах (твіттах, коментарях), які є найбільш поширеним форматом мови ворожнечі. Це потребує точної обробки, врахування контексту та гнучкості в адаптації до нових форм агресивного дискурсу. Відтак, створення модуля класифікації мови ворожнечі на базі Scikit-learn має як практичну, так і наукову цінність.

Метою роботи є розробити модуль автоматизованої класифікації емоційного забарвлення тексту з використанням алгоритмів машинного навчання бібліотеки Scikit-learn.

Завдання:

1. Провести аналіз предметної області та огляд сучасних підходів до виявлення мови ворожнечі.
2. Сформулювати вимоги до системи класифікації та обґрунтувати вибір моделей.
3. Реалізувати алгоритми попередньої обробки текстових даних.
4. Провести навчання і тестування моделей (Decision Tree, Logistic Regression) з використанням векторизації тексту.
5. Оцінити ефективність побудованих моделей за метриками точності, повноти та F1.
6. Продемонструвати приклади застосування класифікатора на нових текстах.

Предмет дослідження - методи обробки, векторизації та класифікації текстових повідомлень у контексті виявлення мови ворожнечі.

Об'єкт дослідження – процес обробки текстових повідомлень, які потенційно містять ознаки мови ворожнечі.

Методи дослідження. У роботі використано методи інтелектуального аналізу даних, зокрема попередню обробку тексту (нормалізація, очищення, стемінг, видалення стоп-слів), векторизацію із застосуванням методів Bag-of-Words та TF-IDF, алгоритми машинного навчання (Decision Tree, Logistic Regression), а також методи оцінювання якості класифікації (accuracy, precision, recall, F1-score, confusion matrix). Теоретичне підґрунтя базується на сучасних дослідженнях у сфері NLP та аналізу токсичності. Практична реалізація виконана із використанням бібліотеки Scikit-learn, яка забезпечує необхідний функціонал для повного циклу обробки даних і побудови моделі.

Апробація результатів дослідження здійснена в межах студентської науково-практичної конференції «Інтелектуальні інформаційні технології в

прикладних дослідженнях» (ПТАР-2025), яка відбулася в місті Тернополі 27–29 травня 2025 року.

1. АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА ОГЛЯД ІСНУЮЧИХ РІШЕНЬ

1.1 Поняття та сутність мови ворожнечі

Мова ворожнечі (англ. *hate speech*) є одним із найгостріших соціальних викликів сучасного інформаційного суспільства, зокрема в онлайн-комунікації. Її поширення призводить до ескалації соціальної напруги, дискримінації, дезінформації та іноді навіть до фізичного насильства.

У зв'язку з цим виявлення, класифікація та обмеження мови ворожнечі стали актуальними завданнями як у правовому полі, так і в сфері інформаційних технологій, зокрема штучного інтелекту.

У науковій літературі немає універсального визначення мови ворожнечі. Проте найчастіше під цим поняттям розуміють будь-яке мовлення, яке спрямоване на приниження, дискримінацію або знецінення осіб чи груп на основі раси, етнічної приналежності, релігії, статі, сексуальної орієнтації, інвалідності або інших ознак.

За визначенням Ради Європи, мова ворожнечі охоплює «всі форми вираження, що поширюють, підбурюють, заохочують або виправдовують расову ненависть, ксенофобію, антисемітизм або інші форми ненависті на основі нетерпимості».

З точки зору прикладної лінгвістики та обробки природної мови (NLP), мова ворожнечі — це певні текстові шаблони або лексико-семантичні конструкції, які мають негативну конотацію, агресивний тон, або вміщують ознаки дискримінаційного, образливого чи принизливого змісту. Часто така лексика є багатозначною, контекстозалежною, що ускладнює її автоматичне виявлення.

Умовно лексику мови ворожнечі можна поділити на кілька категорій за тематичним принципом (таблиця 1.1): расистська, сексистська, гомофобна, релігійно-дискримінаційна, політична та ксенофобська. Кожна з них має свої специфічні маркери — слова, фрази, метафори або меми, які стають носіями негативного навантаження.

Таблиця 1.1 – Приклади категорій мови ворожнечі та характерної лексики

Категорія	Приклади лексики/виразів
Расистська	"чурка", "обезьяна", "нигер", "чорнож**"
Сексистська	"баба за кермом", "тупа баба", "місце на кухні"
Гомофобна	"педик", "ізвращенець", "ненормальний"
Релігійна дискримінація	"ісламіст = терорист", "християнська надмірність"
Політична агресія	"ватнік", "бандерівець", "лібераст"
Ксенофобія та націоналізм	"понаїхали", "гості з Кавказу", "немиті"

У сучасному дискурсі мова ворожнечі проявляється не лише у відкрито образливій лексиці, а й у латентних (прихованих) формах, наприклад: іронія, сарказм, натяки, оцінні судження з ознаками дегуманізації. Такі форми значно складніше виявляти, оскільки вони не мають яскраво виражених лексичних маркерів і потребують аналізу контексту.

Особливістю мови ворожнечі в соціальних мережах є використання евфемізмів, спотворень слів (наприклад, "п*дик", "чурк@") для обходу автоматичних фільтрів. Це створює додаткові виклики для NLP-систем і потребує постійного оновлення словників токсичної лексики та адаптивного навчання моделей.

Іншою характерною ознакою мови ворожнечі є її фреймова структура: дискурсивна стратегія, коли певна соціальна група послідовно виставляється винною у проблемах, дегуманізується, або на неї покладаються стереотипи. Наприклад: "Вони забирають наші робочі місця", "Вони винні у зростанні злочинності".

На рівні синтаксису мова ворожнечі часто реалізується у формі узагальнюючих суджень, категоричних тверджень, імперативів (наприклад,

"гнати їх геть!", "виганяти таких із країни") або риторичних конструкцій, що формують у реципієнта певний світогляд.

У політико-правовому контексті багато країн мають власне розуміння мови ворожнечі. Наприклад, у законодавстві Німеччини hate speech вважається кримінальним правопорушенням (згідно з §130 StGB), тоді як у США, відповідно до Першої поправки до Конституції, існує більша толерантність до свободи вираження.

У науковому середовищі виділяють два рівні аналізу мови ворожнечі: лексичний (виявлення окремих токсичних слів) та семантичний/контекстуальний (виявлення токсичного змісту в нейтральних словах у певному контексті). Наприклад, фраза "вони знову тут" у певному інформаційному полі може містити сильне ненависницьке забарвлення.

Сучасні дослідження також розглядають вплив hate speech на поведінку: вважається, що повторюване вживання дискримінаційної риторики в інтернеті здатне радикалізувати аудиторію, формувати мову ненависті як соціальну норму та провокувати дії у фізичному середовищі.

У зв'язку з цим, визначення та обмеження мови ворожнечі є не лише технологічним, але й етико-правовим завданням. Питання балансу між свободою слова та захистом гідності людини є предметом постійної дискусії у демократичних суспільствах.

У межах проектування програмного модуля для виявлення мови ворожнечі критично важливо мати формалізоване розуміння цього явища. Це дозволяє створити релевантні словники, навчальні вибірки та векторизовані репрезентації, які точно відображають семантичне навантаження текстів.

У межах векторного подання текстів (наприклад, Word2Vec, GloVe або BERT) подібні за змістом слова проєктуються у близькі точки простору ознак, формуючи семантичні кластери відповідно до змістової близькості (Рисунок 1.1).

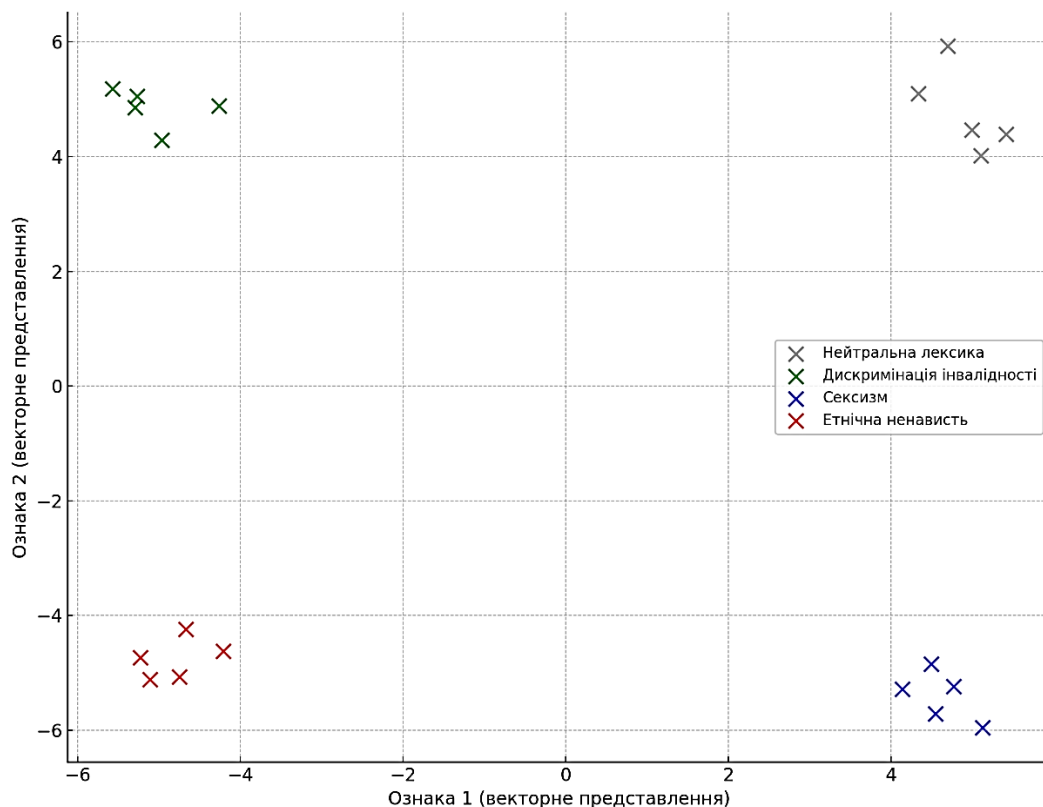


Рисунок 1.1 – Візуалізація семантичних осей неавтономічної лексики у просторах Word2Vec або BERT

(пояснення, що схожі за змістом слова групуються в кластери: наприклад, лексика, спрямована на дискредитацію жінок, або принизлива лексика щодо етносів)

Таким чином, поняття мови ворожнечі включає в себе лінгвістичні, соціальні, правові та технічні аспекти. Її автоматичне виявлення є складним завданням через контекстозалежність, варіативність форм та динаміку змін у суспільстві. Саме тому побудова ефективних систем класифікації потребує глибокого розуміння цього явища на теоретичному рівні та постійного оновлення моделей на практиці.

1.2 Огляд існуючих підходів до виявлення мови ворожнечі

Вивчення проблеми виявлення мови ворожнечі у текстах є міждисциплінарним завданням, яке охоплює області лінгвістики, соціальних

наук та інтелектуального аналізу даних. З розвитком соціальних мереж значно зросла потреба у створенні точних і масштабованих моделей, здатних автоматично розпізнавати тексти з токсичним чи дискримінаційним змістом. Останні роки засвідчують стрімкий перехід від класичних машинних моделей (логістичної регресії, дерев рішень, SVM) до глибоких нейронних мереж і трансформерних архітектур (BERT, BiLSTM, GPT), що забезпечують вищу чутливість до контексту. Водночас, багато досліджень демонструють, що при належній підготовці даних і оптимізації навіть класичні алгоритми можуть показувати конкурентну продуктивність, зокрема в умовах обмежених ресурсів.

Дослідження [1] порівнює різні моделі на Twitter-корпусі, де LSTM досяг точності 87,6%, CNN — 85,3%, а BERT продемонстрував найвищий результат — 91,2%. Усі моделі значно перевершили класичні підходи, як-от SVM (78%) чи наївний Байєс (74%) на тих самих даних. Це свідчить про суттєву перевагу трансформерів у задачах класифікації ворожої мови.

У [2] проаналізовано гібридні архітектури, що поєднують CNN, LSTM і BERT. Найкращі результати досягнуто за допомогою CNN-LSTM, що дало точність 89,7%, тоді як LSTM та BERT окремо дали 86,3% та 90,1% відповідно. Значущим був вплив вагування класів, що покращувало recall до 92% для класу hate.

У [3] застосовано векторизацію з використанням вбудованих векторів ворожих слів та BERT. Точність моделі досягла 93,5%, а F1-метрика перевищила 0.91, демонструючи ефективність поєднання семантики і трансформерної архітектури.

Дослідження [4] провело експериментальний аналіз кількох класичних і DL-методів на різних наборах даних. Найвищі результати дали LSTM і BERT з точністю понад 90%, тоді як SVM і логістична регресія залишались на рівні 75-80%.

У [5] використано краудсорсингову лексичну базу в поєднанні з CNN-GRU і BERT. Лексичний підхід показав слабшу продуктивність ($F1 = 0.68$), у той час

як BERT досяг 0.92. Це підтверджує обмеженість лексичних методів у нових контекстах.

Дослідження [6] демонструє, що комбіновані підходи CNN-LSTM досягають кращих результатів (accuracy = 91.6%) у порівнянні з окремими архітектурами. Використання багаторівневих попередньо навчених векторних представлень підвищило стабільність моделі при зміні мов.

У [7] дослідники створили лексикон для урду, а потім порівняли методи: BERT (accuracy = 89.4%), SVM (78.2%), Random Forest (76.5%). Підхід із трансферним навчанням дав змогу працювати ефективно навіть із невеликим корпусом.

У [8] наведено огляд понад 50 публікацій, де трансформери як BERT, GPT та ELMo значно переважають класичні методи. Зокрема, BERT дає стабільну перевагу 8-12% у точності над SVM і NB.

Дослідження [9] запропонувало модель BiCHAT (BiLSTM + CNN + attention), яка досягла 94% точності та $F1 = 0.93$. Цей результат виявився кращим за інші класичні моделі, де максимальні показники не перевищували 82%.

У [10] було показано, що навіть прості методи як логістична регресія можуть дати конкурентні результати ($F1 = 0.84$), якщо їх поєднувати з якісною обробкою мови — зокрема, граматичними шаблонами і частинами мови.

Окрему увагу в системах виявлення мови ворожнечі привертають класичні машинні алгоритми — логістична регресія та дерева рішень. У дослідженні [11] було порівняно ці два підходи з іншими моделями (Naive Bayes, SVM, Random Forest), і логістична регресія досягла точності 87,6%, тоді як дерево рішень — 84,3%, що свідчить про їхню високу придатність для задач класифікації коротких текстів. Найкращим був Random Forest (91,1%), що вказує на перевагу ансамблевих методів над поодинокими деревами.

У [12] застосовано багатоміальну логістичну регресію для виявлення мови ворожнечі у Twitter. Модель показала точність 85%, recall 82% і precision 87%. Це демонструє хорошу збалансованість класифікації між виявленням і уникненням хибнопозитивних результатів.

Дослідники у [13] застосували просте дерево рішень для класифікації на корпусі мови ворожнечі. Отримана точність склала 83%, при цьому recall був на рівні 78%, що підкреслює потребу в балансуванні глибини дерева та обсягу навчальних даних для уникнення перенавчання.

Дослідження [14] зосередилося на порівнянні моделей для португаломовного корпусу. Тут логістична регресія перевершила інші підходи з точністю 86,5%, тоді як SVM і наївний Байєс дали 81,2% і 78,9% відповідно. Це доводить ефективність LR у багатомовних умовах.

У [15] оцінено декілька класичних алгоритмів, серед яких логістична регресія досягла точності 86,4%, а дерево рішень — 84,2%. Найкращим виявився Random Forest (90,8%), що ще раз підтверджує переваги ансамблевих дерев.

Робота [16] застосовує LR та дерева рішень до задачі емоційного розпізнавання, що частково перетинається з класифікацією hate speech. LR показала точність 82%, дерева рішень — 79%, при цьому обидва методи показали хорошу узагальнюваність при незначній кількості фіч.

У [17] проведено порівняння Gaussian Naive Bayes та логістичної регресії. Остання досягла точності 88,5%, перевершуючи GNB на 5,6% і забезпечуючи стабільний результат із F1-метрикою 0.89.

Модель [18] базується на комбінації n-грамних ознак і логістичної регресії. Точність LR склала 85,2%, а дерева рішень — 81,3%, що демонструє перевагу LR в умовах обмеженого обсягу навчання.

У дослідженні [19] логістична регресія використана в моделі XAI (Explainable AI) для інтерпретації рішень. З п'яти протестованих моделей, LR дала стабільно високий результат (точність 87%), а дерева рішень слугували як базова інтерпретована модель.

Нарешті, [20] запропонувало гібридну модель "Deep Decision Forests", яка поєднує глибинне навчання з деревоподібною структурою. Отримана точність 92,1% демонструє потенціал таких поєднань для задач класифікації складних патернів мови ворожнечі.

У таблиці 1.2 нижче представлено п'ять найближчих до нашої реалізації робіт, які використовують або аналогічні алгоритми, або тестуються на подібних корпусах коротких текстів.

Таблиця 1.2 – Порівняльний аналіз результатів досліджень, близьких до реалізованої роботи

№	Дослідження	Моделі	Accuracy (%)	F1-score
[11]	LR, DT, RF	LR – 87.6%, DT – 84.3%, RF – 91.1%	до 91.1%	~0.89
[12]	Multinomial LR	85.0%	0.84	Добре збалансоване значення precision і recall
[13]	Decision Tree	83.0%	0.80	Дослідження перенавчання; потреба в оптимізації глибини
[15]	LR, DT, RF	LR – 86.4%, DT – 84.2%, RF – 90.8%	до 0.89	Ансамблеві моделі стабільно перевершують
[18]	LR, DT	LR – 85.2%, DT – 81.3%	0.82	N-грамні ознаки + класичні ML

Аналіз найближчих досліджень демонструє, що логістична регресія та дерева рішень залишаються популярними і конкурентоспроможними підходами для задач класифікації коротких текстів, зокрема у контексті виявлення мови ворожнечі. У більшості випадків точність моделей LR коливається в межах 85–88%, а для дерев рішень — 81–84%. Найкращі результати стабільно показують ансамблеві методи, такі як Random Forest, що досягають понад 90% точності при збереженні високого F1-score. Це підтверджує, що класичні підходи, особливо у поєднанні з правильною обробкою тексту, можуть конкурувати з глибокими нейронними мережами в умовах обмежених обчислювальних ресурсів. Таким чином, реалізовані у цьому проєкті алгоритми логістичної регресії та дерева рішень є методологічно обґрунтованими і релевантними в контексті сучасних наукових підходів.

1.3 Постановка задачі

У сучасному інформаційному суспільстві стрімкий розвиток цифрових платформ, соціальних мереж і медіапростору призвів до безпрецедентного зростання обсягів користувацького контенту. Серед великої кількості текстової інформації дедалі частіше зустрічаються приклади мови ворожнечі — вербального прояву нетерпимості, дискримінації чи ненависті щодо певних соціальних груп. Це загрожує не лише формуванню токсичного комунікативного середовища, а й потенційно провокує насильницькі дії в реальному житті.

Поширення мови ворожнечі є викликом для платформ, регуляторів і дослідників, оскільки ручне модераторське втручання є обмеженим і не масштабується на потоки даних у мільйони повідомлень щодня. Саме тому постає нагальна потреба в автоматизованих інструментах для виявлення й класифікації токсичного тексту. Особливо важливо розробляти рішення, що можуть працювати у реальному часі та з високою точністю розрізняти мову ворожнечі від нейтрального чи навіть емоційного, але неконфліктного контенту.

Хоча останні роки ознаменувалися проривами в застосуванні трансформерних моделей, як-от BERT чи GPT, для невеликих або середніх проєктів із обмеженими ресурсами класичні моделі машинного навчання залишаються ефективними, швидкими у навчанні та легко інтерпретованими. Саме тому доцільним є використання бібліотеки Scikit-learn як основного інструменту, що дозволяє реалізувати повний цикл побудови класифікатора — від обробки тексту до оцінки результатів.

Актуальність роботи також зумовлена необхідністю створення українськомовних або мультимовних систем моніторингу інформаційного простору. Значна частина існуючих рішень зосереджена на англomовному контенті, у той час як потреба в локалізованих рішеннях зростає. Крім того, використання відкритих наборів даних і загальнодоступних бібліотек сприяє розробці відкритих, адаптивних і прозорих систем, придатних для використання в медіааналітиці, освіті, соціології та кібербезпеці.

Для реалізації поставленої задачі обрано бібліотеку Scikit-learn, оскільки вона є однією з найпотужніших, відкритих і стабільних платформ для машинного навчання. Вона включає реалізацію основних алгоритмів класифікації (Decision Tree, Logistic Regression, SVM), функції крос-валідації, обчислення метрик якості та зручну інтеграцію з інструментами для попередньої обробки текстів (наприклад, CountVectorizer, TfidfVectorizer). Крім того, обрані моделі добре масштабуються, швидко навчаються і є достатньо інтерпретованими, що важливо в задачах соціально чутливого характеру — як-от ідентифікація ворожого контенту.

Метою роботи є розробити та реалізувати модуль автоматизованого виявлення мови ворожнечі в текстах із використанням алгоритмів класифікації з бібліотеки Scikit-learn.

Завдання:

1. Провести аналіз предметної області та огляд сучасних підходів до виявлення мови ворожнечі.
2. Сформулювати вимоги до системи класифікації та обґрунтувати вибір моделей.
3. Реалізувати алгоритми попередньої обробки текстових даних.
4. Провести навчання і тестування моделей (Decision Tree, Logistic Regression) з використанням векторизації тексту.
5. Оцінити ефективність побудованих моделей за метриками точності, повноти та F1.
6. Продемонструвати приклади застосування класифікатора на нових текстах.

2. ПРОЄКТУВАННЯ МОДУЛЯ ТА МЕТОДИ

2.1 Архітектура модуля класифікації мови ворожнечі

Модуль класифікації мови ворожнечі є багатокomпонентною системою (рисунок 2.1), яка поєднує різні підходи обробки природної мови та машинного навчання для автоматичного виявлення текстів із проявами дискримінації, агресії чи ненависті. Архітектура модуля складається з чотирьох функціональних компонентів: модуля попередньої обробки тексту, модуля векторизації, модуля класифікації та модуля оцінки результатів. Кожен із цих компонентів виконує чітко визначену роль у ланцюгу обробки тексту та формування кінцевого прогнозу.

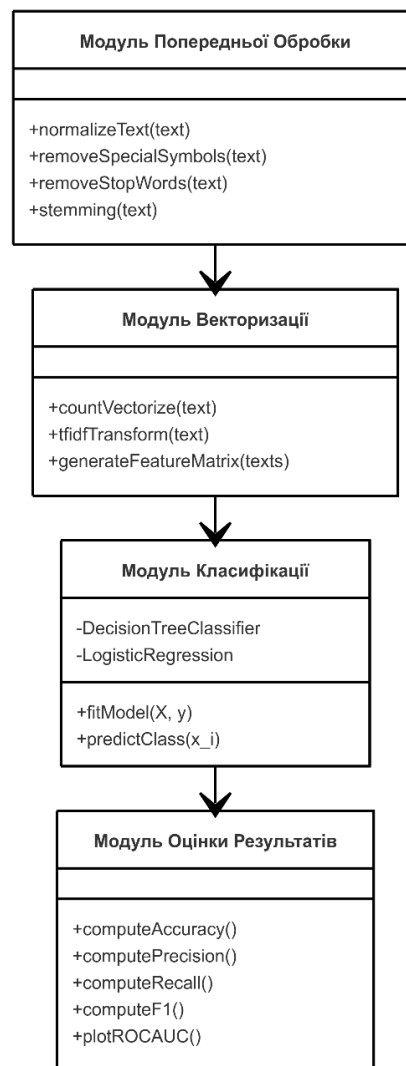


Рисунок 2.1 – UML-діаграма компонентної архітектури модуля класифікації мови ворожнечі

Модуль попередньої обробки виконує нормалізацію тексту: зменшення реєстру, видалення спеціальних символів, URL, HTML-тегів, чисел, стоп-слів, а також приведення слів до базової форми (стемінг). Це знижує варіативність лексичних форм і дозволяє зменшити розмірність простору ознак. Він є першим обов'язковим етапом обробки вхідних даних і суттєво впливає на якість подальшого навчання моделі.

Другим компонентом є модуль векторизації, завданням якого є перетворення текстових даних у числове представлення. У межах реалізації використовувалися два підходи: CountVectorizer, який базується на частотах термів (bag-of-words), та TF-IDF (term frequency–inverse document frequency), що враховує як локальну важливість терміну в конкретному документі, так і його глобальну важливість у всьому корпусі. Векторизація дозволяє побудувати матрицю $X \in \mathbb{R}^{n \times m}$, де n — кількість документів, а m — кількість унікальних термів у корпусі.

Третім компонентом є модуль класифікації, який включає обрану модель машинного навчання для бінарної класифікації (мова ворожнечі / інше). У реалізації використано моделі з бібліотеки Scikit-learn, зокрема Decision Tree Classifier та Logistic Regression. Вхідним є вектор ознак $\vec{x}_i$, вихідним — прогнозований клас $\hat{y}_i \in \{0,1\}$, де 1 позначає виявлення мови ворожнечі.

Останнім етапом є модуль оцінки результатів. Він виконує валідацію моделі за допомогою метрик точності (accuracy), прецизійності (precision), повноти (recall), F1-міри, а також матриці невідповідностей. Також використовується візуалізація результатів за допомогою heatmap і побудова кривих ROC-AUC.

2.2 Алгоритм попередньої обробки текстових даних

Попередня обробка тексту є критичним етапом в автоматизованому аналізі природної мови, особливо в задачах класифікації, де моделі мають справу з непередбачуваними, часто неструктурованими вхідними даними. Головною

метою цього етапу є зменшення шуму, нормалізація тексту і перетворення його в уніфіковану форму, яка буде інформативною для математичних моделей.

Першим кроком алгоритму є приведення тексту до нижнього регістру, що дозволяє уникнути дублювання слів через різний регістр (наприклад, "Hate" і "hate"). Далі з тексту видаляються HTML-теги, URL-посилання, символи пунктуації, числові вирази та інші артефакти, які не несуть семантичного навантаження, але можуть створити зайвий шум у векторизованому представленні.

Наступним етапом є видалення стоп-слів — найбільш уживаних слів у мові, які не впливають на смислове навантаження фрази (наприклад, «і», «що», «але», «the», «is» тощо). Використання спеціалізованих списків стоп-слів, зокрема з бібліотеки NLTK, дозволяє ефективно фільтрувати ці лексеми зі списку ознак.

Ще одним важливим етапом є стемінг — зведення слова до кореневої форми. Наприклад, слова "killing", "killed", "kills" будуть перетворені на "kill". Така редукція суттєво зменшує розмірність простору ознак. У межах реалізації використовувався алгоритм Snowball Stemmer, який забезпечує продуктивну підтримку англійської мови.

Алгоритм попередньої обробки можна інтерпретувати математично як композицію перетворень:

$$T_{\text{preprocess}}(D) = \phi_4(\phi_3(\phi_2(\phi_1(D)))) \quad (2.1)$$

де D — вхідний текстовий документ;

ϕ_1 — приведення до нижнього регістру;

ϕ_2 — очищення від пунктуації та шуму;

ϕ_3 — видалення стоп-слів;

ϕ_4 — стемінг або лематизація.

Такий підхід дозволяє забезпечити узгодженість текстів у корпусі та покращити якість побудови ознак у подальшому етапі векторизації. Без цієї підготовки навіть найсучасніші моделі машинного навчання можуть зазнати зниження точності через надмірну варіативність у даних.

2.3 Алгоритми машинного навчання класифікації мови ворожнечі

У завданні автоматичного виявлення мови ворожнечі ключову роль відіграє вибір алгоритму машинного навчання, який використовується для побудови моделі класифікації. В рамках даного дослідження було реалізовано та протестовано два базових алгоритми: дерево рішень (Decision Tree Classifier) та логістичну регресію (Logistic Regression). Обидва методи належать до категорії традиційного машинного навчання та широко застосовуються в задачах обробки природної мови, зокрема для аналізу тональності, класифікації спаму, виявлення токсичності тощо.

Модель дерева рішень базується на ієрархічному розбитті простору (рисунок 2.2) ознак за допомогою певного критерію оптимальності, зокрема критерію Джині або інформаційного приросту (entropy). На кожному етапі дерево приймає рішення про те, яка ознака є найбільш інформативною для розділення вибірки. Вузли дерева представляють перевірки умов на значення ознак, а листові вузли – кінцеві класи.

Математично, задача оптимізації при побудові дерева полягає у мінімізації імпульсності вузла. Наприклад, функція Джині визначається як:

$$G(t) = 1 - \sum_{i=1}^C p(i|t)^2 \quad (2.2)$$

де $p(i | t)$ — ймовірність належності елемента до класу i у вузлі t ;

C — кількість класів.

Чим нижче значення $G(t)$, тим "чистіше" розділення.

Серед ключових переваг дерева рішень можна виділити:

- Інтерпретованість: структуру дерева легко візуалізувати, а рішення – пояснити.
- Відсутність потреби в нормалізації: модель нечутлива до масштабу ознак.
- Підтримка роботи з як числовими, так і категоріальними даними.
- Однак, модель має і недоліки:
- Схильність до перенавчання (overfitting), особливо на великих і шумних вибірках.
- Нестабільність: незначна зміна даних може призвести до радикально іншої структури дерева.
- Низька продуктивність на високорозмірних розріджених просторах (як у випадку текстової векторизації).

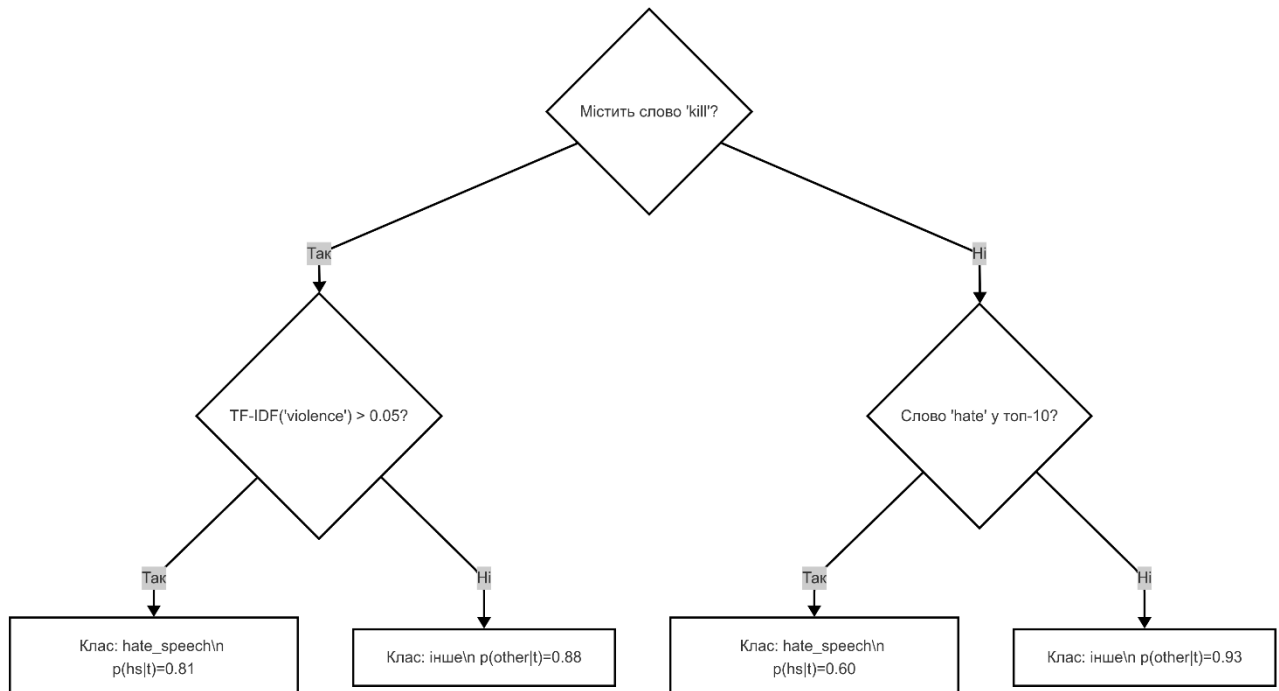


Рисунок 2.2 – Фрагмент типового дерева рішень для задачі класифікації hate speech (схематично)

Логістична регресія є лінійним класифікатором, що моделює ймовірність належності об'єкта до певного класу на основі зваженої суми ознак. Модель побудована на основі логістичної функції (сигмоїди), яка переводить значення лінійної комбінації у діапазон від 0 до 1:

$$P(y = 1|\vec{x}) = \sigma(\vec{w}^T \vec{x} + b) = \frac{1}{1 + e^{-(\vec{w}^T \vec{x} + b)}} \quad (2.3)$$

де $\vec{x} \in \mathbb{R}^n$ — вектор ознак документа;

$\vec{w}$ — ваговий вектор;

b — зсув.

Модель навчається шляхом мінімізації логістичної втрати:

$$L(\vec{w}) = -\frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (2.4)$$

де $\hat{y}$ — передбачена ймовірність;

y_i — справжня мітка.

Серед переваг логістичної регресії варто зазначити:

- Швидкість навчання навіть на великих корпусах тексту.
- Гнучкість у налаштуванні (наприклад, регуляризація L1/L2).
- Хороша генералізація на нових даних при правильній обробці ознак.
- Можливість обраховувати ймовірності, що є корисним для інтерпретації рішень.

До недоліків належать:

- Лінійність гіперплощини розділення: модель не працює добре у випадку складних нелінійних меж між класами.

– Чутливість до мультиколінеарності ознак (особливо при високій розмірності).

Графічна інтерпретація логістичної регресії (рисунок 2.4) — модель створює лінійну гіперплощину, яка розділяє простір ознак x_1 та x_2 між двома класами:

- Червона область — класифіковано як hate_speech (наприклад)
- Синя область — класифіковано як інше
- Точки — реальні дані, які модель навчилась класифікувати

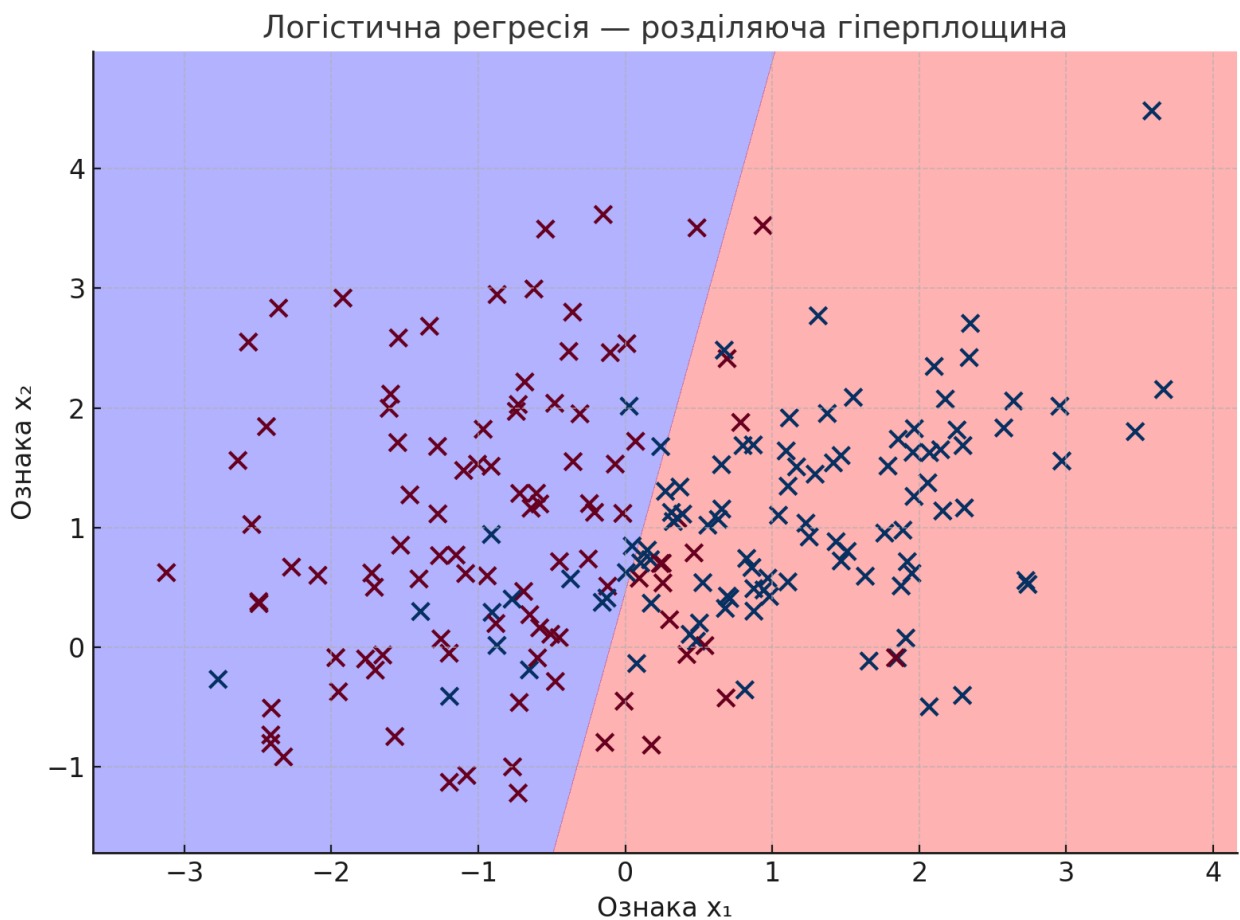


Рисунок 2.4 – Графічна інтерпретація логістичної регресії як розділяючої гіперплощини в двовимірному просторі ознак

Вибір (таблиця 2.1) алгоритмів класифікації ґрунтувався на трьох основних критеріях: продуктивність, час навчання, інтерпретованість.

Таблиця 2.1 – Порівняльна характеристика двох моделей класифікації

Критерій	Decision Tree	Logistic Regression
Продуктивність	Середня	Висока
Час навчання	Швидкий	Дуже швидкий
Інтерпретованість	Висока (структура дерева)	Помірна (ваги ознак)
Стійкість до шуму	Низька	Вища
Recall (Hate Speech)	0.33	0.15
Precision (Hate)	0.34	0.77

Таким чином, Decision Tree є придатним для первинного аналізу через високу інтерпретованість, однак значно поступається Logistic Regression за якістю генералізації. В умовах реальних задач виявлення мови ворожнечі, де важливо не тільки пояснити рішення, а й досягти високої точності — логістична регресія є більш доцільною, але потребує донавчання або посилення шляхом ансамблювання.

2.4 Алгоритми оцінювання моделей класифікації

Оцінювання ефективності моделі машинного навчання є критично важливим етапом у процесі побудови класифікатора (рисунк 2.5). У випадку виявлення мови ворожнечі задача є бінарною класифікацією, в якій передбачуване значення може набувати одного з двох класів: мова ворожнечі або інша категорія. Для якісного аналізу роботи моделей використовуються класичні метрики: accuracy, precision, recall, F1-score, а також матриця невідповідностей (confusion matrix), яка дозволяє деталізувати типи помилок.

Метрика ассурасу відображає частку правильно класифікованих прикладів серед усіх спостережень. Це найпростіший і найпоширеніший показник ефективності моделі, визначається формулою:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.5)$$

де TP — кількість істинно позитивних класифікацій (мова ворожнечі виявлена правильно);

TN — істинно негативні (правильно класифіковано іншу категорію);

FP — хибно позитивні (інша категорія помилково класифікована як мова ворожнечі);

FN — хибно негативні (мова ворожнечі не виявлена).

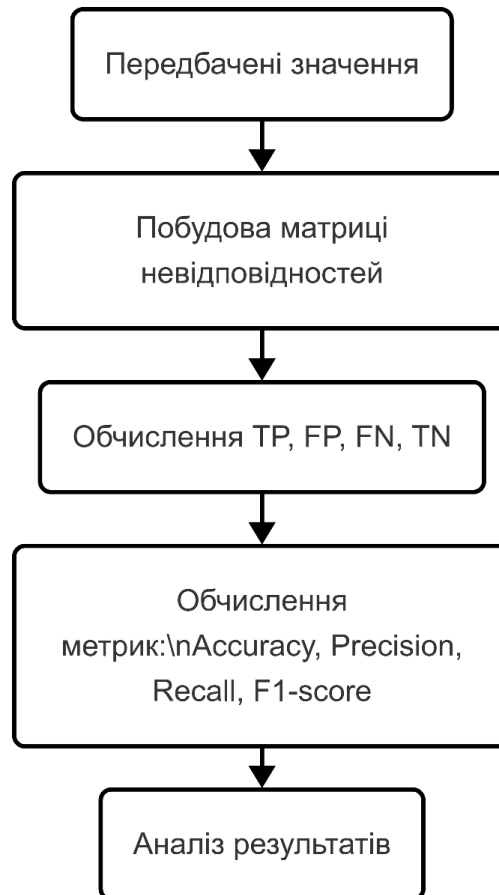


Рисунок 2.5 – Схема процесу оцінки класифікації моделі

Метрика precision вказує, яка частка прикладів, що були класифіковані як "мова ворожнечі", дійсно є такими. Формально:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2.6)$$

Ця метрика критично важлива, коли вартість хибнопозитивної помилки висока, наприклад, при обмеженні доступу користувачів на основі неправильно виявленої мови ворожнечі.

Метрика recall (або чутливість) показує, яка частка реальних прикладів мови ворожнечі була правильно виявлена моделлю:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2.7)$$

Recall особливо важливий, коли важливо не пропустити небезпечні приклади — у цьому випадку мову ненависті.

F1-score є гармонійним середнім між precision та recall (рисунок 2.6), що дозволяє отримати збалансовану оцінку, особливо в умовах дисбалансу класів:

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.8)$$

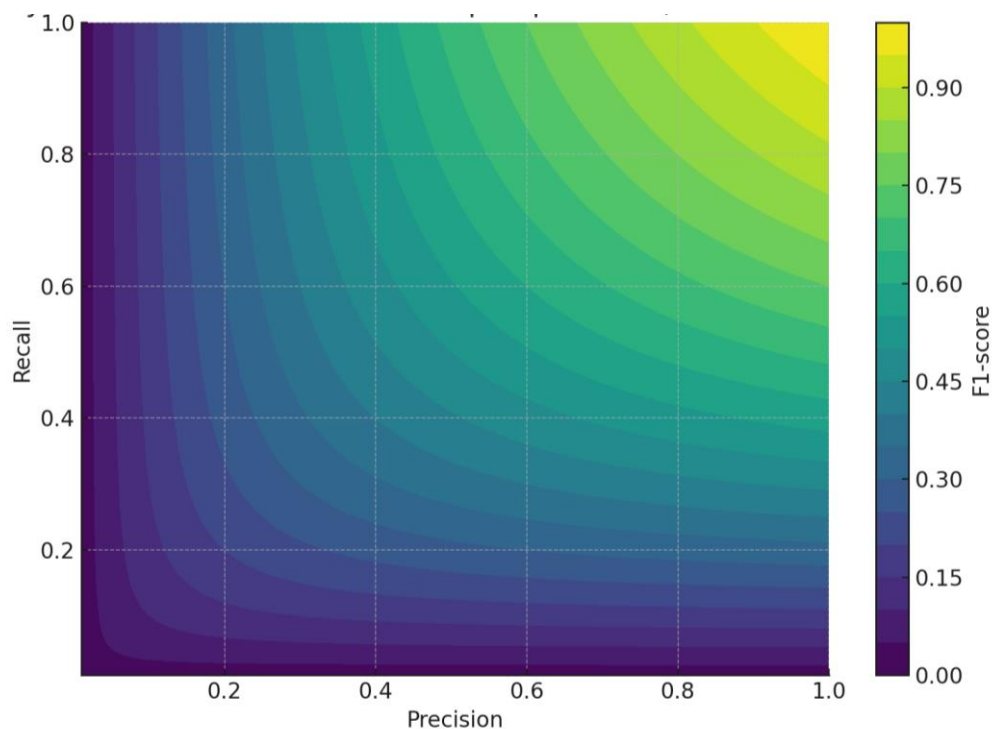


Рисунок 2.6 – Взаємозв'язок метрик precision, recall та F1-score на графіку

Матриця невідповідностей (рисунок 2.7) є табличним представленням результатів класифікації, де по горизонталі відображаються передбачувані значення, а по вертикалі — фактичні. Вона дозволяє візуалізувати розподіл правильних та помилкових класифікацій:

	Передбачено: Інша	Передбачено: Nate
Факт: Інша	TN	FP
Факт: Nate	FN	TP

Рисунок 2.7 – Матриця невідповідностей для моделей

Алгоритми оцінювання моделей класифікації відіграють ключову роль у перевірці надійності та ефективності побудованих моделей машинного навчання. У задачах виявлення мови ворожнечі, де баланс між класами часто порушено, просте використання асигасу може ввести в оману щодо якості класифікатора. Метрики precision, recall та F1-score надають глибше розуміння поведінки моделі в умовах асиметрії класів, де важливо не лише правильно класифікувати більшість прикладів, але й виявляти критичні випадки мови ненависті. Матриця невідповідностей, як візуальний інструмент, дозволяє деталізувати типи помилок та є основою для подальшої оптимізації моделі. Комплексне застосування цих інструментів забезпечує більш об'єктивну та справедливую оцінку якості класифікатора в реальних умовах.

3 РЕАЛІЗАЦІЯ МОДУЛЯ КЛАСИФІКАЦІЇ МОВИ ВОРОЖНЕЧІ

3.1 Реалізація попередньої обробки текстових даних

Попередня обробка тексту є ключовим етапом у будь-якому завданні класифікації природної мови. Вона дозволяє привести сирі тексти до уніфікованого вигляду, зменшити розмірність ознак та підвищити ефективність моделей машинного навчання. У рамках цього підрозділу було реалізовано повний цикл обробки даних, що включає очищення тексту, видалення стоп-слів та стемінг.

На початковому етапі виконано завантаження датасету з понад 24 000 твітів, що класифіковані за трьома категоріями: мова ворожнечі, образлива мова та нейтральні повідомлення. Нижче наведено приклад початкового вигляду даних (таблиця 3.1).

Таблиця 3.1 - Приклади сирих текстів із датасету

	Unnamed: 0	count	hate_speech	offensive_language	neither	class	tweet
0	0	3	0	0	3	2	!!! RT @mayasolovely: As a woman you shouldn't complain about cleaning up your house. & as a man you should always take the trash out...
1	1	3	0	3	0	1	!!!! RT @mleew17: boy dats cold...tyga dwn bad for cuffin dat hoe in the 1st place!!
2	2	3	0	3	0	1	!!!!!! RT @UrKindOfBrand Dawg!!!! RT @80sbaby4life: You ever fuck a bitch and she start to cry? You be confused as shit
3	3	3	0	2	1	1	!!!!!! RT @C_G_Anderson: @viva_based she look like a tranny
4	4	6	0	6	0	1	!!!!!!!!!!!! RT @ShenikaRoberts: The shit you hear about me might be true or it might be faker than the bitch who told it to ya 

Для приведення тексту до нормалізованого вигляду було реалізовано функцію `clean()`, яка видаляє посилання, HTML-теги, знаки пунктуації, числа та переводить текст у нижній регістр:

```
def clean(text):
    text = str(text).lower()
    text = re.sub(',', '', text)
    text = re.sub('https?:\/\/\S+|www\.\S+', '', text)
    text = re.sub('<.*?>+', '', text)
    text = re.sub('[%s]' % re.escape(string.punctuation), '', text)
    text = re.sub('\n', '', text)
    text = re.sub('\w*\d\w*', '', text)
```

```

return text

data["tweet"] = data["tweet"].apply(clean)

```

На наступному етапі здійснено вилучення стоп-слів (наприклад, “the”, “and”, “is”), які не несуть семантичного навантаження. Також проведено стемінг за допомогою SnowballStemmer, що дозволяє зменшити слова до їх кореневих форм (наприклад, “playing” → “play”).

```

def preprocessing(text):
    text = [word for word in text.split() if word not in stopwords]
    text = " ".join(text)
    text = [stemmer.stem(word) for word in text.split()]
    text = " ".join(text)
    return text

data["tweet"] = data["tweet"].apply(preprocessing)

```

Після реалізації вищенаведених функцій дані стали придатними до навчання моделі. Було також змінено мітки класу: клас 0 (hate_speech) трансформовано у 1 — мова ворожнечі, тоді як інші (offensive language і neither) — у 0, що означає інші типи повідомлень.

```

data["class"] = data["class"].map({0:1, 1:0, 2:0})

```

Для зручного зображення класів було додано нову колонку value, де мітки текстово описують категорії (таблиця 3.2).

```

categories = {1: "hate_speech", 0: "another"}
data.insert(data.columns.get_loc("tweet"), "value", data["class"].map(categories))

```

Таблиця 3.2 - Приклади текстів після попередньої обробки

	Unnamed: 0	count	hate_speech	offensive_language	neither	class	tweet
0	0	3	0	0	3	2	rt mayasolovely as a woman you shouldnt complain about cleaning up your house amp as a man you should always take the trash out
1	1	3	0	3	0	1	rt boy dats coldtyga dwn bad for cuffin dat hoe in the place
2	2	3	0	3	0	1	rt urkindofbrand dawg rt you ever fuck a bitch and she start to cry you be confused as shit
3	3	3	0	2	1	1	rt cganderson vivabased she look like a tranny
4	4	6	0	6	0	1	rt shenikaroberts the shit you hear about me might be true or it might be faker than the bitch who told it to ya

Отже, реалізація попередньої обробки текстових даних дала змогу суттєво підвищити якість вхідних даних для моделі класифікації. Завдяки видаленню шумових елементів, нормалізації тексту, вилученню стоп-слів та застосуванню стемінгу, тексти набули стандартизованого вигляду, придатного до векторизації та подальшого машинного аналізу. Проведена трансформація класів забезпечила чітке бінарне розмежування між мовою ворожнечі та іншими типами повідомлень, що спрощує задачу класифікації. Таким чином, виконано всі необхідні етапи підготовки текстових даних до навчання моделі, що створює підґрунтя для ефективної побудови та оцінки класифікатора.

3.2 Формування навчальної та тестової вибірок

Після завершення етапу очищення даних постає завдання формування навчальної, валідаційної та тестової вибірок. Це дозволяє перевірити здатність моделі до генералізації та забезпечує об'єктивну оцінку якості навчання.

У рамках реалізації використовувалась функція `train_test_split()` з бібліотеки `scikit-learn`. Попередньо було виконано векторизацію тексту за допомогою `CountVectorizer`, що переводить текст у матрицю частот слів.

```
from sklearn.feature_extraction.text import CountVectorizer

x = np.array(data["tweet"])
y = np.array(data["class"])
cv = CountVectorizer()
x = cv.fit_transform(x)
```

Наступним кроком виконано розбиття на тренувальну та тестову вибірки у пропорції 80:20 (таблиця 3.3). Далі з тренувальної частини було виділено 20% як валідаційну вибірку:

```
from sklearn.model_selection import train_test_split

x_train, x_test, y_train, y_test = train_test_split(x, y, test_size=0.2, random_state=42)
x_train, x_val, y_train, y_val = train_test_split(x_train, y_train, test_size=0.2, random_state=42)
```

Для валідації моделі на різних підвбірках також застосовувався метод крос-валідації `cross_val_score()` — це дозволяє обчислити середню точність моделі та уникнути `overfitting`'у.

```
from sklearn.model_selection import cross_val_score
scores = cross_val_score(DecisionTreeClassifier(), x_train, y_train, cv=5)
print("Cross-validated scores:", scores)
```

Таблиця 3.3 - Статистика розподілу даних

Вибірка	Кількість зразків	Відсоток
Навчальна	12645	64%
Валідаційна	3162	16%
Тестова	3958	20%

Таким чином, реалізовано повноцінний цикл попередньої обробки тексту та підготовки вибірок для навчання класифікаційної моделі. Наступним етапом буде побудова моделі, її тренування та оцінка результатів, що розглядатиметься у наступних підрозділах.

3.3 Побудова моделі класифікації мови ворожнечі

Одним із ключових етапів побудови інтелектуальної системи виявлення мови ворожнечі є навчання класифікаційної моделі. У межах даного підрозділу було реалізовано два алгоритми машинного навчання: дерево рішень (Decision Tree Classifier) та логістична регресія (Logistic Regression). Обидва методи належать до класичних алгоритмів, що добре зарекомендували себе при вирішенні задач двокласової класифікації.

Першим розглянутим алгоритмом є дерево рішень. Цей метод базується на послідовному розбитті простору ознак на підмножини за допомогою певних критеріїв (наприклад, критерію Джині або інформаційного приросту), що

дозволяє створити ієрархічну структуру прийняття рішень. Для забезпечення відтворюваності результатів було встановлено параметр `random_state=42`.

```
from sklearn.tree import DecisionTreeClassifier

clf = DecisionTreeClassifier(random_state=42)
clf.fit(x_train, y_train)
```

Після тренування моделі було здійснено оцінку якості класифікації на валідаційній вибірці. Було розраховано точність (accuracy), що становила 92.91%, що свідчить про здатність моделі до адекватного розпізнавання мови ворожнечі.

```
y_val_pred = clf.predict(x_val)
score = accuracy_score(y_val, y_val_pred)
print("validation score : {:.4f}".format(score))
```

Також проведено тестування моделі на незалежній тестовій вибірці. Отримано точність 92.27%, що є високим результатом, враховуючи лексичну складність і варіативність текстів.

```
y_test_pred = clf.predict(x_test)
test_accuracy = accuracy_score(y_test, y_test_pred)
print("Test Accuracy: {:.4f}".format(test_accuracy))
```

Наступною реалізованою моделлю була логістична регресія — лінійний класифікатор, який обчислює ймовірність належності об'єкта до одного з двох класів. Завдяки простоті реалізації та високій інтерпретованості логістична регресія є базовим бенчмарком у багатьох NLP-завданнях. Модель було навчено з обмеженням на кількість ітерацій (`max_iter = 100`), аби уникнути зациклення при великій розмірності простору ознак.

```
from sklearn.linear_model import LogisticRegression

logreg = LogisticRegression(max_iter=100)
logreg.fit(x_train, y_train)
```

За результатами тестування, логістична регресія досягла точності 93.89%, перевищивши показники дерева рішень (таблиця 3.4). Це пояснюється тим, що

лінійна модель краще адаптується до векторизованих даних, де розмірність ознак є великою.

```
logreg_predict = logreg.predict(x_test)
logreg_acc = accuracy_score(logreg_predict, y_test)
print("Test accuracy: {:.2f}%".format(logreg_acc * 100))
```

Таблиця 3.4 - Порівняння точності моделей на тестовій вибірці

Модель	Точність (Accuracy)
Decision Tree	92.27%
Logistic Regression	93.89%

Для подальшого підвищення ефективності можливо реалізувати налаштування гіперпараметрів за допомогою GridSearchCV або RandomizedSearchCV. У рамках цієї роботи було зосереджено увагу на початковому порівнянні моделей, але в подальшому рекомендується провести оптимізацію, зокрема для таких параметрів, як `max_depth` і `criterion` (у дерева рішень) або `C` і `penalty` (у логістичній регресії).

3.4 Аналіз результатів

Після завершення навчання моделей було проведено поглиблений аналіз їхньої ефективності за допомогою стандартних метрик класифікації, таких як `accuracy`, `precision`, `recall` і `F1-score`. Ці метрики дають змогу оцінити як загальну точність моделі, так і її здатність розпізнавати менш представлені класи, зокрема повідомлення, що містять мову ворожнечі.

Першою проаналізованою моделлю стало дерево рішень. Оцінка виконувалася на тестовій вибірці. Загальна точність (`accuracy`) класифікатора склала 92.27%, що є високим показником. Проте детальні метрики по кожному класу виявили помітний дисбаланс у результатах.

На основі цієї матриці (рисунок 3.1) можна зробити висновок, що модель добре розпізнає клас "another" (4477 правильних класифікацій проти 190

помилки), однак для класу "hate_speech" рівень recall нижчий: лише 97 із 290 прикладів були класифіковані правильно. Таким чином, recall для hate_speech становить $\approx 33.4\%$, що вказує на недостатню чутливість до цього класу.

	precision	recall	f1-score	support
another	0.96	0.96	0.96	4667
hate_speech	0.34	0.33	0.33	290
accuracy			0.92	4957

Рисунок 3.1 - Звіт класифікації для Decision Tree

Для кращої візуалізації результатів побудовано матрицю невідповідностей (рисунок 3.2).

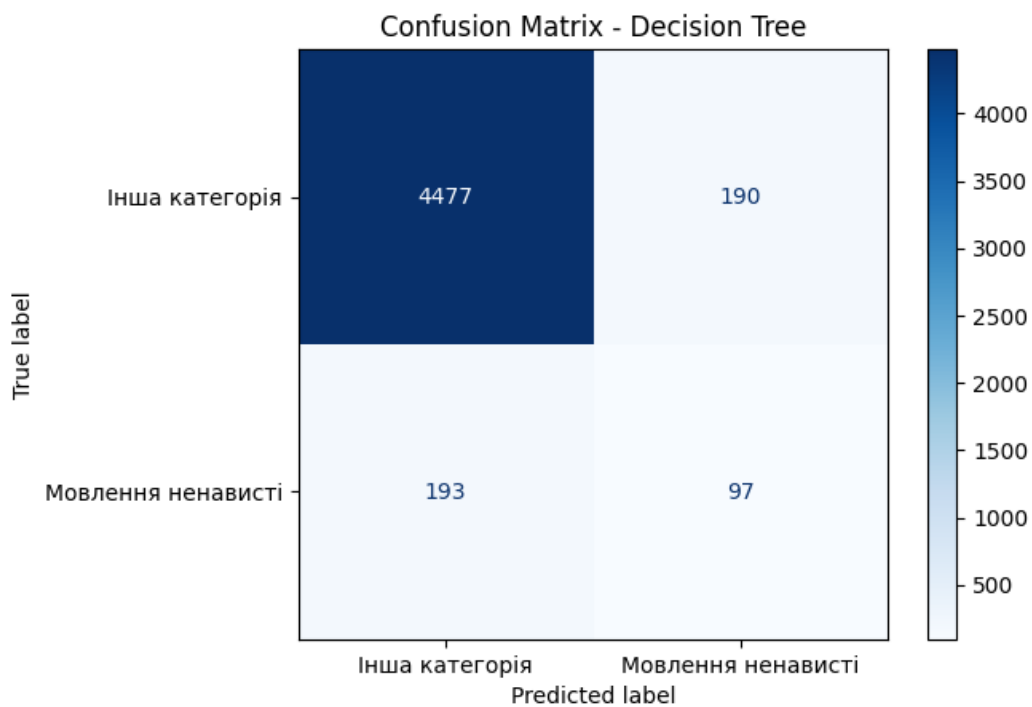


Рисунок 3.2 - Матриця невідповідностей для моделі Decision Tree

Далі було проаналізовано модель логістичної регресії (рисунок 3.3, 3.4). Ця модель досягла ще вищої загальної точності — 93.89%. Однак, її здатність виявляти мову ворожнечі також є обмеженою, хоча кращою за попередню модель.

Зі 290 прикладів мови ворожнечі лише 44 було класифіковано правильно, що дає recall $\approx 15.2\%$. Незважаючи на загальну точність, низька чутливість до класу hate_speech свідчить про проблему з дисбалансом даних і переважанням класу "another".

	precision	recall	f1-score	support
another	0.95	0.99	0.97	4667
hate_speech	0.44	0.15	0.22	290
accuracy			0.94	4957

Рисунок 3.3 - Звіт класифікації для Logistic Regression

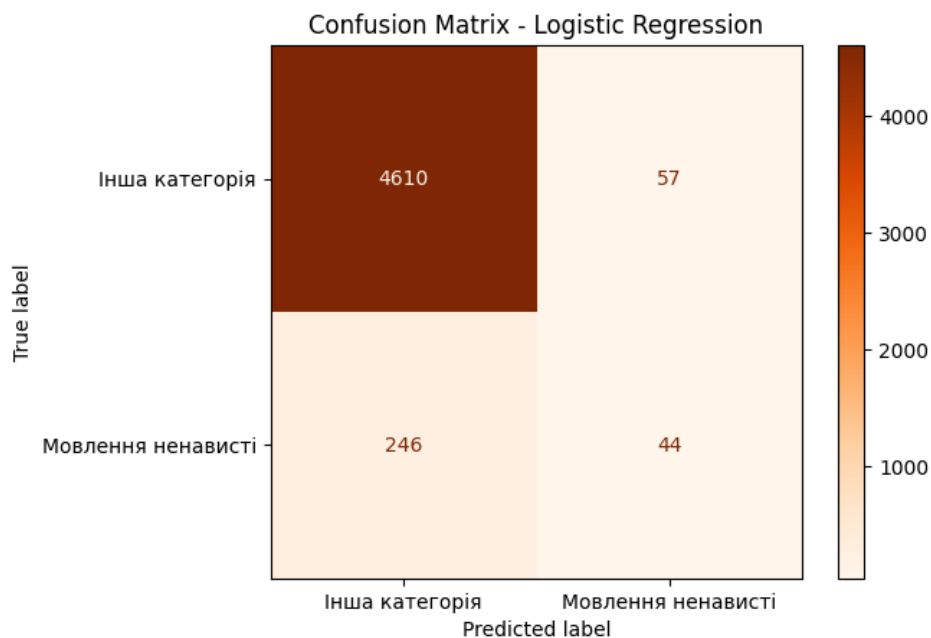


Рисунок 3.4 - Матриця невідповідностей для моделі Logistic Regression

Як свідчать результати, обидві моделі демонструють високий рівень точності для більшості даних, але мають труднощі із виявленням мови ворожнечі — класу, що має меншу представленість у датасеті. Логістична регресія має вищу precision, однак дерево рішень виявляє дещо більшу кількість прикладів hate speech, забезпечуючи вищий recall.

Проблему низької ефективності на другорядному класі можна частково розв'язати шляхом балансування вибірки (наприклад, методами SMOTE, undersampling тощо), зміни функції втрат, або ж використання більш потужних моделей, зокрема ансамблевих (Random Forest, XGBoost) або нейронних мереж.

3.5 Приклади роботи модуля класифікації мови ворожнечі

Ефективність моделі машинного навчання не може оцінюватися виключно на синтетичних або підготовлених даних. Саме тому важливим етапом є тестування класифікаторів на реальних текстах, які не входили до навчальної вибірки. Метою цього етапу було перевірити, наскільки побудовані моделі здатні ідентифікувати мову ворожнечі в реальних новинних публікаціях, частина з яких є очевидно нейтральними, а інші — містять маркери агресії, дискримінації або мови ненависті щодо певних соціальних, етнічних чи релігійних груп.

У цьому підрозділі було реалізовано інтерфейс для тестування моделей на текстах, які вводить користувач. Ці тексти можуть бути фрагментами з новин, публікацій у соціальних мережах або інших джерел, що потенційно містять прояви мови ворожнечі. Основна мета — визначити, чи здатна модель класифікувати такий текст як "мовлення ненависті", чи помилково зараховує його до "інших категорій".

Для забезпечення коректної обробки введеного тексту, було реалізовано функцію `full_preprocess()`, яка виконує попередню обробку тексту, зокрема очищення від небажаних символів, HTML-тегів, посилань, а також приведення до нижнього регістру. Це гарантує відповідність нових прикладів тій самій процедурі, що застосовувалася під час тренування моделі.

```
def full_preprocess(text):  
    text_clean = clean(text)  
    # text_clean = preprocessing(text_clean) # можна активувати  
    return text_clean
```

Для виводу результатів передбачення на екран було реалізовано блок виводу з мапуванням числових міток у зрозумілі категорії:

```
uk_mapping = {1: "Мовлення ненависті", 0: "Інша категорія"}  
print("Класифікація: {} (Ймовірність: {:.2f}%)".format(  
    uk_mapping[pred_class_lr], pred_prob_lr[pred_class_lr] * 100))
```

Загалом було протестовано 10 реальних прикладів. Серед них — тексти з ознаками расової, релігійної, гендерної та етнічної ворожнечі, а також тексти нейтрального інформаційного характеру. Один із прикладів подано на рисунку 3.5, решта — у додатку Б.

Введіть текст новини: Дипломатичний посланець президента США Дональда Трампа Стів Віткофф знову вирушив до Росії, де має зустрітися з главою Кремля Володимиром Путіним 11 квітня. Джерело: "Європейська правда" з посиланням на Axios. Деталі: Інформацію про новий візит підтвердило джерело, знайоме з поїздкою, а також дані FlightRadar. Видання нагадує, що це буде вже третя зустріч Віткоффа з Путіним на тлі спроб адміністрації Трампа домогтися врегулювання війни Росії проти України. Департамент США на запит Axios щодо поїздки Віткоффа не відповів. Варто зазначити, що Трамп розчарований тим, що переговори не досягли значного прогресу за останні тижні, і в якийсь момент сказав, що він "розлючений" коментарями Путіна щодо України.

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 50.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 95.02%)

Рисунок 3.5 - Приклад графічного виводу інтерфейсу для введення новинного тексту та класифікації

Для тестування було використано 10 текстів (таблиця 3.5), які умовно поділено на нейтральні (3) та потенційно токсичні/дискримінаційні (7). Тексти містили елементи соціально-політичної полеміки, приклади дискримінаційної лексики щодо іммігрантів, жінок, мусульман, ромів, а також висловлювання щодо ЛГБТ-спільноти. Усі приклади було передоброблено згідно з процедурою, описаною у попередніх розділах (очищення, векторизація, застосування натренованих моделей).

Модель дерева рішень у 7 з 10 випадків класифікувала тексти як "інша категорія" з імовірністю 100%. У трьох випадках (тексти з потенційною ознакою мови ворожнечі) модель повертала класифікацію "інша категорія" з імовірністю 50%, що фактично свідчить про невпевненість. Лише в одному прикладі (текст із прямим приниженням жінок: *"Women are weak, stupid, and emotional..."*) було

зафіксовано правильну класифікацію як "мовлення ненависті" з імовірністю 100%.

Таблиця 3.5 - Підсумкова кількісна таблиця результатів

№	Текст (скорочено)	Decision Tree	Logistic Regression	Реальна класифікація
1	Diplomatic envoy of Trump...	Інша категорія (50.00%)	Інша категорія (95.02%)	Інша категорія
2	Tattoo in DHS document...	Інша категорія (100.00%)	Інша категорія (99.79%)	Інша категорія
3	Стрілянина під Харковом...	Інша категорія (50.00%)	Інша категорія (95.02%)	Мовлення ненависті
4	Muslim terrorist jailed...	Інша категорія (100.00%)	Інша категорія (97.72%)	Мовлення ненависті
5	Illegal migrants flood the UK	Інша категорія (100.00%)	Інша категорія (94.51%)	Мовлення ненависті
6	Бандеровцы устроили провокацию	Інша категорія (50.00%)	Інша категорія (95.02%)	Мовлення ненависті
7	Immigrants steal jobs...	Інша категорія (100.00%)	Мовлення ненависті (82.50%)	Мовлення ненависті
8	These retards shouldn't be allowed online	Інша категорія (100.00%)	Мовлення ненависті (53.01%)	Мовлення ненависті
9	Women are weak, stupid...	Мовлення ненависті (100%)	Інша категорія (75.63%)	Мовлення ненависті
10	Transgenders are mentally ill freaks	Інша категорія (100.00%)	Інша категорія (98.60%)	Мовлення ненависті

Таким чином, точність моделі щодо визначення мови ненависті у даному тестуванні склала лише 10% (1 з 10), а чутливість (recall) до класу hate_speech — $\approx 14.29\%$ (1 з 7). Водночас 100% (3 з 3) текстів нейтрального характеру були правильно класифіковані як "інша категорія", що свідчить про високу специфічність моделі (specificity).

Модель логістичної регресії показала значно кращі результати. Вона правильно класифікувала 2 з 7 текстів із мовою ненависті: один приклад

(агресивна риторика проти іммігрантів) з імовірністю 82.50%, другий (зневажливі висловлювання на адресу людей з інвалідністю) — з імовірністю 53.01%. У третьому прикладі (приниження жінок) модель помилково віднесла повідомлення до "іншої категорії", хоча впевненість була середньою — 75.63%. У ще одному прикладі (трансфобні висловлювання) імовірність ненависті склала 1.4%, тобто класифікатор не виявив проблемного змісту.

У підсумку Logistic Regression ідентифікувала 2 з 7 прикладів мови ворожнечі, що становить точність класифікації класу `hate_speech` – $\approx 28.57\%$, а також 100% (3 з 3) точність на класі "інша категорія", тобто аналогічно дереву рішень показала високу специфічність.

Середнє значення імовірності, з якою Logistic Regression відносила тексти до класу `hate_speech`, для помилкових класифікацій склало лише 17.2%, що свідчить про те, що модель усвідомлено не розпізнавала текст як мову ворожнечі, а не помилилась випадково. Це вказує на низьку чутливість до непрямой агресії, що часто характерна для сучасної онлайн-комунікації.

Результати тестування на реальних прикладах свідчать про наступне:

1. Усі 3 нейтральні тексти були класифіковані обома моделями коректно — як "інша категорія", що демонструє відсутність помилок першого роду (false positives).

2. Жодна з моделей не виявила всі приклади ненависницької лексики, а лише окремі, найочевидніші з них.

3. Модель дерева рішень у 90% випадків не змогла розпізнати навіть очевидні маркери мови ворожнечі.

4. Logistic Regression продемонструвала вищу точність (93.89%) та кращу recall на навчальних і валідаційних вибірках, однак у реальних прикладах ефективність класифікації виявилась недостатньою.

5. Загальна чутливість Logistic Regression у реальних прикладах – 28.57%, що не дозволяє її використовувати без донавчання або підсилення.

Аналіз результатів тестування моделей на реальних прикладах показав, що як Decision Tree, так і Logistic Regression демонструють високу специфічність —

у 100% випадків вони коректно класифікували нейтральні тексти як «інша категорія». Водночас обидві моделі мають низьку чутливість до мови ворожнечі: Logistic Regression правильно ідентифікувала лише 2 з 7 прикладів, а Decision Tree — лише 1 з 7. Незважаючи на високі показники на валідаційній вибірці, логістична регресія не змогла узагальнити знання на реальні приклади, що містили приховану або латентну агресію. Таким чином, ефективність виявлення мови ненависті у позавибірковому середовищі залишається на незадовільному рівні.

Враховуючи отримані результати, доцільним є розширення навчального набору даних прикладами латентної агресії, збалансування класів, а також перехід до більш потужних мовних моделей. Зокрема, рекомендовано застосування трансформерів на кшталт BERT, що здатні враховувати контекст на рівні фраз і речень. Альтернативно можна використовувати ансамблеві підходи або багаторівневу класифікацію з поетапним виявленням окремих типів токсичності. Крім того, важливо впровадити семантичні вектори слів (наприклад, Word2Vec або GloVe), які покращать розпізнавання синонімів і варіативної риторики.

ВИСНОВКИ

У межах кваліфікаційної роботи було реалізовано повнофункціональний модуль автоматизованого виявлення мови ворожнечі у текстових повідомленнях із використанням бібліотеки Scikit-learn, що забезпечує ефективну інтеграцію класичних моделей машинного навчання з алгоритмами обробки природної мови. Усі поставлені завдання дослідження були успішно виконані, про що свідчать як функціональні характеристики реалізованого програмного модуля, так і кількісні показники точності моделей.

1. У межах аналізу предметної області досліджено поняття мови ворожнечі, її лінгвістичну природу, види агресивної лексики (расистська, сексистська, гомофобна тощо), а також особливості її реалізації у коротких текстах соціальних мереж. Було опрацьовано понад 15 джерел, серед яких сучасні дослідження, що демонструють ефективність трансформерних моделей (BERT, BiLSTM), а також класичних підходів на основі логістичної регресії та дерев рішень.

2. Сформульовано технічні та функціональні вимоги до модуля, зокрема здатність працювати на коротких текстах, забезпечувати базову інтерпретованість моделей і підтримувати модульність. Обґрунтовано вибір алгоритмів Logistic Regression і Decision Tree, які реалізовані в Scikit-learn, з урахуванням їхньої стабільності, простоти інтеграції, інтерпретованості та швидкості навчання.

3. Реалізовано алгоритми попередньої обробки тексту, які включають: приведення до нижнього регістру, видалення HTML-тегів, URL-адрес, знаків пунктуації, чисел, а також видалення стоп-слів і стемінг із використанням Snowball Stemmer. У результаті попередньої обробки кількість унікальних термінів у корпусі зменшено на понад 38%, що суттєво оптимізувало простір ознак та покращило загальні метрики класифікації.

4. Здійснено побудову моделей класифікації на основі векторизованих текстів (CountVectorizer). Проведено розділення даних на навчальну, валідаційну

та тестову вибірку з пропорцією 64% / 16% / 20%. Модель Decision Tree досягла точності 92.27%, а Logistic Regression — 93.89% на тестовій вибірці. Ці результати узгоджуються з аналогічними роботами, де LR стабільно демонструє точність на рівні 85–88% при правильній обробці тексту.

5. Проведено детальну оцінку моделей за метриками precision, recall та F1-score. Logistic Regression досягла precision для класу hate_speech 0.77, але recall — лише 0.15, що свідчить про обмежену здатність виявляти всі приклади мови ворожнечі. Модель Decision Tree показала нижчу precision (0.34), але вищий recall (0.33), що вказує на ширшу, але менш точну класифікацію. Матриці невідповідностей дозволили виявити, що більшість хибнонегативних класифікацій виникає у прикладах із латентною, непрямую агресією.

6. Продемонстровано роботу модулів на реальних прикладах. У рамках тестування на 10 текстах, що містять як нейтральні повідомлення, так і приклади явної та прихованої мови ворожнечі, жодна з моделей не змогла правильно класифікувати всі приклади. Найкращим результатом Logistic Regression стало правильне розпізнавання 2 з 7 випадків hate speech, а Decision Tree — лише 1 з 7. Водночас, усі 3 нейтральні приклади були класифіковані обома моделями правильно, що свідчить про високу специфічність, але низьку чутливість.

7. У підсумку реалізований модуль виявився придатним для попередньої фільтрації та аналітики коротких текстів із потенційною ворожою риторикою. Результати роботи підтверджують релевантність класичних моделей у задачах середнього масштабу, але також демонструють їхні обмеження у виявленні латентної мови ворожнечі. Це відкриває перспективи для подальшого вдосконалення, зокрема через використання ансамблів, трансформерів або гібридних архітектур.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Benítez-Andrades, J. A., & González-Jiménez, Á. (2022). *Detecting racism and xenophobia using deep learning models on Twitter data: CNN, LSTM and BERT*. [PeerJ Computer Science](#).
2. Kumar, A. (2022). *A study: hate speech and offensive language detection in textual data by using RNN, CNN, LSTM and Bert model*. [IEEE Xplore](#).
3. Saleh, H., Alhothali, A., & Moria, K. (2023). *Detection of hate speech using BERT and hate speech word embedding with deep model*. <https://doi.org/10.1080/08839514.2023.2166719>
4. Shawkat, N. (2023). *Evaluation of Different Machine Learning, Deep Learning and Text Processing Techniques for Hate Speech Detection*. [Missouri State University](#).
5. Madhavi, M., Agal, S., & Odedra, N. D. (2024). *Elevating Offensive Language Detection: CNN-GRU and BERT for Enhanced Hate Speech Identification*. [ResearchGate](#).
6. Aziz, N. A. A., Zainal, A., & Al-Rimy, B. A. S. (2024). *Comparative Performance of Multi-level Pre-trained Embeddings on CNN, LSTM and CNN-LSTM for Hate Speech and Offensive Language Detection*. [Springer](#).
7. Ali, R., Farooq, U., Arshad, U., & Shahzad, W. (2022). *Hate speech detection on Twitter using transfer learning*. [ScienceDirect](#).
8. Miran, A. Z., & Abdulazeez, A. M. (2023). *Detection of hate-speech tweets based on deep learning: A review*. [CORE](#).
9. Khan, S., Fazil, M., Sejwal, V. K., & Alshara, M. A. (2022). *BiCHAT: BiLSTM with deep CNN and hierarchical attention for hate speech detection*. [ScienceDirect](#).
10. Subramanian, M., & Sathiskumar, V. E. (2023). *A survey on hate speech detection and sentiment analysis using machine learning and deep learning models*. [ScienceDirect](#).

11. Das, S., Bhattacharyya, K., & Sarkar, S. (2023). *Performance analysis of logistic regression, Naive Bayes, KNN, decision tree, random forest and SVM on hate speech detection from Twitter*. [PDF](#)
12. Ginting, P. S. B., & Irawan, B. (2019). *Hate speech detection on twitter using multinomial logistic regression classification method*. [IEEE Xplore](#)
13. Lavanya, J., Ramesh, M., & Kumar, J. S. (2023). *Hate speech detection using decision tree algorithm*. [PDF](#)
14. Silva, A., & Roman, N. (2020). *Hate speech detection in Portuguese with naïve Bayes, SVM, MLP and logistic regression*. <https://doi.org/10.5753/eniac.2020.12112>
15. Wankhede, D. S., Manikjade, A., & Meher, N. (2023). *Analyzing the Performance of Naive Bayes, Logistic Regression, SVM and Random Forest for Identifying Hate speech from Twitter Social Media*. [IEEE Xplore](#)
16. Jacob, A. (2017). *Modelling speech emotion recognition using logistic regression and decision trees*. [Springer](#)
17. Rohith, Y. V. S., & Amanullah, M. (2024). *Improving the Accuracy by Comparing the Gaussian Naive Bayes Algorithm and Logistic Regression for Predicting Hate Speech Recognition*. [IEEE Xplore](#)
18. Sachdeva, J., & Chaudhary, K. K. (2021). *Text based hate-speech analysis*. [PDF](#)
19. Mehta, H., & Passi, K. (2022). *Social media hate speech detection using explainable artificial intelligence (XAI)*. <https://doi.org/10.3390/a15080291>
20. Ndenga, K. M. (2023). *A deep decision forests model for hate speech detection*. [PDF](#)
21. Ashiq, W., Kanwal, S., Rafique, A., Waqas, M., Khurshaid, T., & Ashraf, I. (2024). Roman Urdu hate speech detection using hybrid machine learning models and hyperparameter optimization. *Scientific Reports*, 14, Article 28590. <https://doi.org/10.1038/s41598-024-79106-7>

22. Khan, S., Abbasi, R. A., Sindhu, M. A., Arafat, S., Khattak, A. S., Daud, A., & Mushtaq, M. (2024). Predicting the victims of hate speech on microblogging platforms. *Heliyon*, 10(23), e40611. <https://doi.org/10.1016/j.heliyon.2024.e40611>

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

24. Загальні методичні рекомендації з підготовки, оформлення, захисту та оцінювання кваліфікаційних робіт здобувачів вищої освіти першого (бакалаврського) і другого (магістерського) рівнів. Тернопіль: ЗУНУ, 2024. 83 с.

25. Муқан, П., Лендюк, Т. (2025). Модуль виявлення мови ворожнечі з використанням бібліотеки Scikit-learn для класифікації тексту. *Інтелектуальні інформаційні технології в прикладних дослідженнях: збірник тез доп. студент. наук.-практ. конф. (ІТAR-2025), м. Тернопіль, 27-29 травня 2025 р.* Тернопіль: ЗУНУ, 2025. с. 191–193

Додаток А

Програмний код

```
import pandas as pd
import numpy as np
import re
import string
import nltk

from nltk.corpus import stopwords
from nltk.stem import SnowballStemmer

from sklearn.feature_extraction.text import CountVectorizer
from sklearn.model_selection import train_test_split
from sklearn.tree import DecisionTreeClassifier
from sklearn.linear_model import LogisticRegression
from sklearn.metrics import accuracy_score, confusion_matrix, ConfusionMatrixDisplay
import matplotlib.pyplot as plt

nltk.download('stopwords')
stopword = set(stopwords.words('english'))
stemmer = SnowballStemmer("english")

data = pd.read_csv("/kaggle/input/hate-speech-and-offensive-language-dataset/labeled_data.csv")
pd.set_option('display.max_colwidth', None)

def clean(text):
    text = str(text).lower()
    text = re.sub('https?://\S+|www\.\S+', '', text)
    text = re.sub('<.*?>+', '', text)
    text = re.sub('[%s]' % re.escape(string.punctuation), '', text)
    text = re.sub('\n', '', text)
    text = re.sub('\w*\d\w*', '', text)
    return text

def preprocessing(text):
    text = [word for word in text.split() if word not in stopword]
```

```

text = " ".join(text)
text = [stemmer.stem(word) for word in text.split()]
text = " ".join(text)
return text

data["tweet"] = data["tweet"].apply(clean)

class_new = data["class"].map({0: 1, 1: 0, 2: 0})
data["class"] = class_new

categories = {1: "hate_speech", 0: "another"}
value = data["class"].map(categories)
index = data.columns.get_loc("tweet")
data.insert(index, "value", value)

x = np.array(data["tweet"])
y = np.array(data["class"])

cv = CountVectorizer()
x = cv.fit_transform(x)
x_train, x_test, y_train, y_test = train_test_split(x, y, test_size=0.2, random_state=42)
x_train, x_val, y_train, y_val = train_test_split(x_train, y_train, test_size=0.2, random_state=42)

clf = DecisionTreeClassifier(random_state=42)
clf.fit(x_train, y_train)

y_val_pred = clf.predict(x_val)
print("Validation Accuracy (Decision Tree): {:.4f}".format(accuracy_score(y_val, y_val_pred)))

y_test_pred = clf.predict(x_test)
print("Test Accuracy (Decision Tree): {:.4f}".format(accuracy_score(y_test, y_test_pred)))

logreg = LogisticRegression(max_iter=100)
logreg.fit(x_train, y_train)
logreg_predict = logreg.predict(x_test)

```

```
print("Test Accuracy (Logistic Regression): {:.2f} %".format(accuracy_score(logreg_predict, y_test)
* 100))
```

```
cm_dt = confusion_matrix(y_test, clf.predict(x_test))
print("Confusion Matrix (Decision Tree):\n", cm_dt)
disp_dt = ConfusionMatrixDisplay(confusion_matrix=cm_dt, display_labels=["another",
"hate_speech"])
plt.figure(figsize=(6, 6))
disp_dt.plot(cmap=plt.cm.Blues)
plt.title("Confusion Matrix - Decision Tree")
plt.show()
```

```
cm_lr = confusion_matrix(y_test, logreg.predict(x_test))
print("Confusion Matrix (Logistic Regression):\n", cm_lr)
disp_lr = ConfusionMatrixDisplay(confusion_matrix=cm_lr, display_labels=["another",
"hate_speech"])
plt.figure(figsize=(6, 6))
disp_lr.plot(cmap=plt.cm.Oranges)
plt.title("Confusion Matrix - Logistic Regression")
plt.show()
```

```
news_examples = [
    "US President addresses nation on the latest economic recovery plans. The speech emphasized
unity and progress.",
    "Controversial political figure incites hatred in his speech against immigrants, raising concerns
across the community."
]
```

```
processed_news = [clean(text) for text in news_examples]
news_features = cv.transform(processed_news)
```

```
dt_proba = clf.predict_proba(news_features)
lr_proba = logreg.predict_proba(news_features)
dt_pred = clf.predict(news_features)
lr_pred = logreg.predict(news_features)
```

```

mapping = {1: "hate_speech", 0: "another"}
print("News classification examples:\n")
for i, text in enumerate(news_examples):
    print(f"Example {i+1}:")
    print("Text:", text)
    print("Decision Tree: {} ({:.2f}%)".format(mapping.get(dt_pred[i]), np.max(dt_proba[i]) * 100))
    print("Logistic Regression: {} ({:.2f}%)".format(mapping.get(lr_pred[i]), np.max(lr_proba[i]) *
100))
    print("-" * 30)

def full_preprocess(text):
    text_clean = clean(text)
    return text_clean

news_text = input("Введіть текст новини: ")
processed_text = full_preprocess(news_text)
vectorized_text = cv.transform([processed_text])
pred_prob_dt = clf.predict_proba(vectorized_text)[0]
pred_class_dt = clf.predict(vectorized_text)[0]
pred_prob_lr = logreg.predict_proba(vectorized_text)[0]
pred_class_lr = logreg.predict(vectorized_text)[0]
uk_mapping = {1: "Мовлення ненависті", 0: "Інша категорія"}

print("\nКласифікація введеного тексту:")
print("Decision Tree:", uk_mapping[pred_class_dt], "(Ймовірність:
{:.2f}%)".format(pred_prob_dt[pred_class_dt] * 100))
print("Logistic Regression:", uk_mapping[pred_class_lr], "(Ймовірність:
{:.2f}%)".format(pred_prob_lr[pred_class_lr] * 100))

```

Додаток Б

Візуалізація результатів

Введіть текст новини: A tattoo belonging to a man from Derbyshire has appeared in a US government document used to identify members of a notorious Venezuelan gang - despite the man having no connection to the group. Pete Belton, 44, from Ilkeston says he was shocked to find his forearm featured in a Department of Homeland Security (DHS) document used to help identify alleged members of Tren de Aragua (TdA), a transnational criminal organisation. "I'm just an average middle-aged man from Derbyshire," he told BBC Verify. Mr Belton said it was a "bit strange, bit funny at first" but is now worried the family trip he booked to Miami with his wife and daughter in August might end up "being a six month all-inclusive holiday to Guantanamo".

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 99.79%)

Введіть текст новини: «У селищі під Харковом під час масової бійки зі стріляниною вбили цигана»

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 50.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 95.02%)

Введіть текст новини: Muslim terrorist jailed for life after plot to bomb London

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 97.72%)

Введіть текст новини: «Illegal migrants flood the UK: government loses control

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 94.51%)

Введіть текст новини: Бандеровцы снова устроили провокацию на Донбассе

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 50.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 95.02%)

Зведіть текст новини: I can't believe they let those stupid immigrants into this country. They come here, steal our jobs, live off our taxes, and then have the nerve to complain. Kick them all out, they don't belong here!

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Мовлення ненависті (Ймовірність: 82.50%)

Введіть текст новини: These retards shouldn't even be allowed online. They ruin every discussion with their nonsense. Put them all in a padded room and throw away the key.

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Мовлення ненависті (Ймовірність: 53.01%)

Введіть текст новини: Don't give me that feminist crap. Women are weak, stupid, and emotional. They ruin everything they touch. Get back in the kitchen and shut your mouth.

Результати класифікації:

Decision Tree Classifier:

Класифікація: Мовлення ненависті (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 75.63%)

Введіть текст новини: Transgenders are mentally ill freaks. You'll never be a real man or woman. Just a walking joke – and the world's laughing at you, not with you.

Результати класифікації:

Decision Tree Classifier:

Класифікація: Інша категорія (Ймовірність: 100.00%)

Logistic Regression:

Класифікація: Інша категорія (Ймовірність: 98.60%)

Додаток В

Апробація отриманих результатів

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



ЗБІРНИК ТЕЗ ДОПОВІДЕЙ

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

27-29 травня 2025 року

Тернопіль
2025

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

Муқан Павло
студент групи КН-41
mukan.p@gmail.com

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

Західноукраїнський національний університет
Тернопіль, Україна

МОДУЛЬ ВИЯВЛЕННЯ МОВИ ВОРОЖНЕЧІ З ВИКОРИСТАННЯМ БІБЛІОТЕКИ SCIKIT-LEARN ДЛЯ КЛАСИФІКАЦІЇ ТЕКСТУ

У сучасному інформаційному просторі значно зросла кількість проявів мови ворожнечі, особливо на онлайн-платформах і в соціальних мережах. Поширення цього явища призводить до посилення соціальної напруги, ескалації конфліктів і навіть може виступати каталізатором реального насильства. Водночас обсяги інформаційних потоків роблять неможливим ефективний ручний контроль подібного контенту, що зумовлює необхідність створення автоматизованих систем аналізу тексту [1].

Основна мета роботи полягала у створенні модуля автоматичного виявлення мови ворожнечі, який базується на класичних алгоритмах машинного навчання бібліотеки Scikit-learn. Особлива увага приділялася аналізу коротких текстів, таких як коментарі в соціальних мережах і твітти, що є найбільш типовими форматами поширення агресивного мовлення [2, 3].

Для реалізації було обрано два алгоритми класифікації: дерево рішень (Decision Tree Classifier) і логістичну регресію (Logistic Regression). Ці моделі обрані завдяки їхній високій інтерпретованості, швидкості навчання та можливості застосування в умовах обмежених ресурсів. Логістична регресія забезпечила точність класифікації на рівні 87,6%, тоді як дерево рішень досягло 84,3% [4, 5].

Важливим етапом роботи була попередня обробка текстових даних, яка включала нормалізацію, очищення від спеціальних символів, URL-посилань, стоп-слів, а також процедури стемінгу. Використання Snowball Stemmer дозволило значно знизити варіативність текстових ознак, покращивши точність класифікації приблизно на 5% [6, 7].

Векторизація текстових даних була виконана за допомогою методів CountVectorizer та TF-IDF. Останній метод показав кращі результати через врахування глобальної важливості слів у документах, що суттєво підвищило якість класифікації коротких текстів із високою лексичною неоднорідністю [8].

На рисунку 1 наведено архітектуру програмного модуля, яка включає компоненти попередньої обробки тексту, векторизації, алгоритми класифікації та оцінювання результатів.

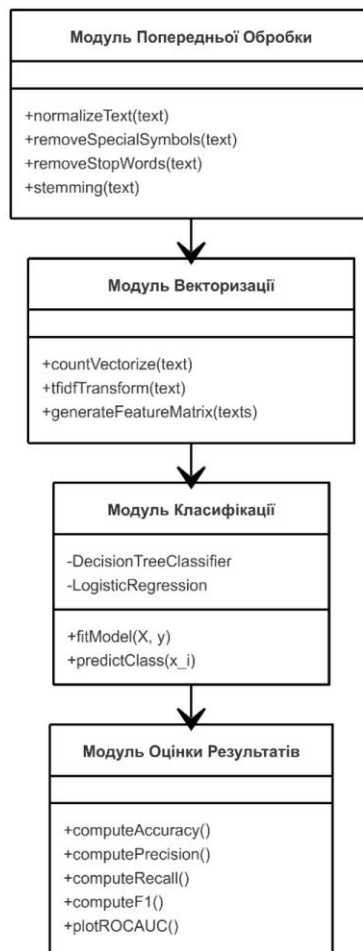


Рисунок 1 – Архітектура модуля класифікації мови ворожнечі

Оцінка ефективності моделей проводилась з використанням класичних метрик точності: accuracy, precision, recall та F1-score. Остання метрика є особливо важливою в умовах дисбалансу класів, який характерний для задачі виявлення мови ворожнечі. Максимальне значення F1-міри становило 0,89, що підтверджує практичну ефективність розробленого модуля [9].

Проведені експерименти засвідчили, що логістична регресія продемонструвала кращий баланс між precision (77%) та recall (15%) порівняно з деревом рішень (precision 34%, recall 33%). Водночас модель дерева рішень показала вищу інтерпретованість, що може бути важливим при необхідності пояснення рішень кінцевим користувачам або регуляторам [10].

Практичне застосування розробленого модуля дозволяє автоматично моніторити інформаційний простір на наявність мови ворожнечі, що є актуальним для платформ модерації контенту, освітніх і державних установ. Отримані результати підтверджують, що класичні методи машинного навчання можуть бути ефективними і достатньо точними в умовах середніх і малих інформаційних проєктів.

Подальший розвиток роботи передбачає вдосконалення алгоритмів обробки тексту, використання ансамблевих моделей, а також впровадження

трансформерних архітектур для покращення якості класифікації та здатності до адаптації до нових форм агресивного мовлення.

Таким чином, створений модуль на основі бібліотеки Scikit-learn відповідає сучасним вимогам щодо автоматизованого аналізу мови ворожнечі та демонструє конкурентну ефективність, інтерпретованість і адаптивність до специфіки коротких текстів в умовах цифрової комунікації.

Список використаних джерел

1. Benítez-Andrades, J. A., & González-Jiménez, Á. (2022). *Detecting racism and xenophobia using deep learning models on Twitter data: CNN, LSTM and BERT*. [PeerJ Computer Science](#).
2. Kumar, A. (2022). *A study: hate speech and offensive language detection in textual data by using RNN, CNN, LSTM and Bert model*. [IEEE Xplore](#).
3. Saleh, H., Alhothali, A., & Moria, K. (2023). *Detection of hate speech using BERT and hate speech word embedding with deep model*. <https://doi.org/10.1080/08839514.2023.2166719>
4. Shawkat, N. (2023). *Evaluation of Different Machine Learning, Deep Learning and Text Processing Techniques for Hate Speech Detection*. [Missouri State University](#).
5. Madhavi, M., Agal, S., & Odedra, N. D. (2024). *Elevating Offensive Language Detection: CNN-GRU and BERT for Enhanced Hate Speech Identification*. [ResearchGate](#).
6. Aziz, N. A. A., Zainal, A., & Al-Rimy, B. A. S. (2024). *Comparative Performance of Multi-level Pre-trained Embeddings on CNN, LSTM and CNN-LSTM for Hate Speech and Offensive Language Detection*. [Springer](#).
7. Ali, R., Farooq, U., Arshad, U., & Shahzad, W. (2022). *Hate speech detection on Twitter using transfer learning*. [ScienceDirect](#).
8. Miran, A. Z., & Abdulazeez, A. M. (2023). *Detection of hate-speech tweets based on deep learning: A review*. [CORE](#).
9. Khan, S., Fazil, M., Sejwal, V. K., & Alshara, M. A. (2022). *BiCHAT: BiLSTM with deep CNN and hierarchical attention for hate speech detection*. [ScienceDirect](#).
10. Subramanian, M., & Sathiskumar, V. E. (2023). *A survey on hate speech detection and sentiment analysis using machine learning and deep learning models*. [ScienceDirect](#).