

М. І. КОНИК

аспірант кафедри інформаційних систем та мереж  
Національний університет «Львівська політехніка»  
ORCID: 0009-0001-1692-4262

## ПОРІВНЯЛЬНИЙ АНАЛІЗ ПІДХОДІВ ТА МЕТОДІВ ВИЗНАЧЕННЯ ТОНАЛЬНОСТІ ТЕКСТУ В КОНТЕКСТІ ОПРАЦЮВАННЯ ВІДГУКІВ МЕШКАНЦІВ МІСТА

У статті представлено ґрунтовний порівняльний аналіз сучасних підходів і методів, що застосовуються для визначення емоційного тону тексту, з особливим акцентом на україномовний контент. Актуальність дослідження зумовлена зростаючою потребою в ефективних інструментах для обробки зворотного зв'язку від громадян, зібраного через цифрові платформи, мобільні застосунки та соціальні мережі, які є цінними джерелами інформації для вдосконалення міського управління та якості надання послуг. Метою дослідження є подолання методологічного та лінгвістичного розриву в аналізі тональності українською мовою, яка, на відміну від англійської, досі залишається недостатньо забезпеченою лексичними базами, корпусами та попередньо навченими моделями.

У статті систематизовано методи аналізу тональності в межах чотирьох основних парадигм: лексиконно-орієнтованого, підходу на основі правил, підходу машинного навчання та гібридного. Для кожного підходу розглянуто його теоретичне підґрунтя, характерні алгоритми, лінгвістичні інструменти та приклади практичного застосування. Лексиконно-орієнтовані методи, зокрема ті, що використовують словник NRC EtoLex та корпусні інструменти на кшталт Sketch Engine, відзначаються простотою використання та адаптивністю до середовищ із обмеженими ресурсами. Підходи на основі правил, такі як VADER і LIWC, виокремлюються здатністю враховувати синтаксичну структуру, підсилювачі та заперечення, забезпечуючи кращу інтерпретованість, хоча й обмежену мовну універсальність.

У розділі, присвяченому машинному навчанню, розглянуто як традиційні алгоритми класифікації – зокрема моделі на основі граничних гіперплощин, ймовірного підходу та деревоподібних структур, – так і сучасні архітектури глибокого навчання, включно з багатопаровими нейронними мережами. Також проаналізовано підхід максимальної ентропії як приклад статистичного моделювання з мінімальними припущеннями щодо вхідних ознак.

Особливу увагу приділено сучасним глибоким нейронним мережам – згортковим нейронним мережам (CNN), рекурентним мережам LSTM та трансформерним архітектурам, таким як BERT, RoBERTa і GPT. Представлені емпіричні результати з останніх досліджень підтверджують високу ефективність трансформерних моделей у багатомовних умовах, зокрема при аналізі тональності українських текстів, де моделі XLM-RoBERTa та Ukr-RoBERTa досягли точності понад 91 %.

У завершальній частині статті розглянуто гібридні моделі, що поєднують переваги різних підходів з метою підвищення стійкості, точності та адаптивності до доменної специфіки. Ось перефразований варіант:

Запропонована класифікаційна структура, розроблена на основі результатів дослідження, забезпечує цілісне уявлення про існуючі методи та становить методологічне підґрунтя для розробки інтелектуальних систем підтримки прийняття рішень в управлінні містом із залученням громадськості.

У висновках окреслено перспективи подальших досліджень, серед яких – локалізація моделей глибокого навчання для української мови, інтеграція аналізу тональності з виділенням тем і іменованих сутностей, а також застосування методів напівконтрольованого та активного навчання для покращення результатів у разі обмеженості анотованих даних. Загалом, запропонована таксономія не лише відображає поточний стан технологій аналізу тональності, а й задає вектори їх подальшого розвитку в умовах багатомовності та соціально орієнтованих застосувань.

**Ключові слова:** аналіз тональності, лексикон, підхід на основі правил, гібридні методи, трансформерні моделі, природна мова.

М. І. KONYK

Postgraduate Student at the Department of Information Systems and Networks  
National University "Lviv Polytechnic"  
ORCID: 0009-0001-1692-4262

## COMPARATIVE ANALYSIS OF APPROACHES AND METHODS FOR SENTIMENT ANALYSIS OF TEXT IN THE CONTEXT OF PROCESSING CITY RESIDENTS' FEEDBACK

The article presents a comprehensive comparative analysis of contemporary approaches and methods used to determine the emotional tone of text, with a particular emphasis on Ukrainian-language content. The relevance of this study stems from the growing need for effective tools to process citizen feedback collected via digital platforms, mobile applications,

and social media – all of which serve as valuable sources of information for improving urban governance and service delivery. The research aims to address both methodological and linguistic gaps in sentiment analysis for the Ukrainian language, which, unlike English, remains under-resourced in terms of lexicons, corpora, and pretrained models.

The study systematizes sentiment analysis methods into four main paradigms: those relying on predefined lexicons, those governed by linguistic rules, data-driven learning techniques, and integrated models that combine multiple strategies. For each paradigm, the theoretical foundations, typical algorithms, linguistic tools, and practical application examples are discussed. Lexicon-based methods, such as those utilizing the NRC EmoLex dictionary and corpus tools like Sketch Engine, are noted for their simplicity and adaptability to low-resource environments. Rule-based systems, including VADER and LIWC, stand out for their ability to account for syntactic structure, intensifiers, and negations, offering better interpretability, albeit with limited language generalizability.

The section on machine learning explores both traditional classification algorithms – including models based on decision boundaries, probabilistic inference, and tree-like structures – and modern deep learning architectures, such as multilayer neural networks. The maximum entropy approach is also examined as a representative of statistical modeling that requires minimal assumptions about input features.

Particular attention is given to recent advancements in deep neural networks, namely convolutional neural networks (CNN), recurrent LSTM networks, and transformer-based architectures such as BERT, RoBERTa, and GPT. Empirical results from recent studies confirm the high effectiveness of transformer models in multilingual contexts, especially in sentiment analysis of Ukrainian texts, where models like XLM-RoBERTa and Ukr-RoBERTa have achieved accuracy levels exceeding 91 %.

The final part of the article discusses hybrid models that combine the strengths of different paradigms to enhance robustness, accuracy, and adaptability to domain-specific data. The proposed classification framework, developed on the basis of the conducted research, provides a coherent overview of existing methods and serves as a methodological foundation for the development of intelligent decision-support systems in urban governance with active citizen participation.

The conclusions outline directions for further research, including the localization of deep learning models for Ukrainian, integration of sentiment analysis with topic modeling and named entity recognition, and the application of semi-supervised and active learning techniques to improve outcomes in scenarios with limited annotated data. Overall, the proposed taxonomy not only reflects the current state of sentiment analysis technologies but also outlines future development trajectories in multilingual and socially oriented applications.

**Key words:** sentiment analysis, lexicon, rule-based approach, hybrid methods, transformer models, natural language.

### Постановка проблеми

У сучасних умовах швидкої урбанізації та зростання чисельності міського населення ефективне управління міською інфраструктурою стає важливим завданням для органів місцевого самоврядування. Відгуки громадян, зібрані через різноманітні цифрові платформи, соціальні мережі, мобільні застосунки та інші канали комунікації, є цінним джерелом інформації про якість міських послуг, стан інфраструктури, екологічні проблеми та інші аспекти життєдіяльності міст. Аналіз відгуків мешканців дозволяє не лише виявляти існуючі проблеми, але й прогнозувати тенденції громадської думки, що сприяє підвищенню якості прийняття рішень у міському управлінні. Застосування методу сентимент-аналізу щодо обробки текстових даних допомагає оцінити ставлення мешканців до різних аспектів міського середовища.

Актуальність дослідження методів визначення емоційного забарвлення тексту зумовлена необхідністю оперативного виявлення проблемних аспектів у зворотному зв'язку від громадськості, формування аналітичних звітів для органів державної та муніципальної влади, а також підвищення якості прийняття управлінських рішень [1]. Емоційний аналіз тексту (sentiment analysis) є складною задачею природної мовної обробки, яка потребує врахування як семантичних, так і синтаксичних особливостей мови, контексту висловлювання, полісемії та іронії.

На відміну від англійської мови, для якої розроблено велику кількість ресурсів, лексиконів, корпусів та моделей машинного навчання для аналізу емоцій, українськомовні тексти досі є недостатньо забезпеченими відповідними інструментами та лінгвістичними ресурсами. Це створює додаткові труднощі при розробці автоматизованих систем аналізу емоційного тону україномовного контенту та підвищує вимоги до вибору методологічного підходу до вирішення цього завдання.

### Зв'язок із важливими науковими чи практичними завданнями

Розв'язання задачі визначення емоційного забарвлення україномовного тексту має важливе значення як з наукової, так і з практичної точки зору. З наукового боку це пов'язано з розширенням лінгвістичних та інтелектуальних технологій для обробки природної мови, розробкою моделей аналізу тональності, адаптованих до морфологічних та синтаксичних особливостей української мови, а також з удосконаленням методів машинного та глибинного навчання для мультимовного середовища.

З прикладної точки зору, результати таких досліджень є критично важливими для створення інтелектуальних систем підтримки прийняття рішень, зокрема у сфері муніципального управління, де аналіз звернень громадян дозволяє здійснювати моніторинг громадських настроїв, виявляти негативні сигнали та формувати обґрунтовані

управлінські дії на основі емоційного тону зворотного зв'язку. Застосування таких підходів сприяє розвитку паратисипативних механізмів управління та підвищує рівень прозорості, відкритості й довіри між владою та громадянами [2]. У цьому контексті постає завдання порівняльного аналізу сучасних підходів і методів визначення емоційного забарвлення текстів, їх класифікації та визначення придатності до обробки україномовного матеріалу.

### Формулювання мети дослідження

Метою роботи є систематизація сучасних підходів і методів визначення тональності текстових даних, а також розроблення узагальненої класифікаційної моделі, яка враховує специфіку завдань аналізу емоційного тону у контексті інтелектуальної платформи для підтримки прийняття рішень управління містом із залученням громадськості. Особливу увагу зосереджено на можливостях застосування відповідних методів до україномовного тексту, що надходить від мешканців у вигляді звернень, пропозицій та відгуків.

### Викладення основного матеріалу дослідження

У сучасній практиці аналізу тональності сформувалося кілька концептуальних підходів: лексиконно-орієнтований, rule-based, машинного навчання та гібридний. У цій роботі здійснено їх порівняльний аналіз із прикладами реалізації, оцінкою ефективності та перспективністю застосування в україномовному інформаційному просторі.

#### 1. Підхід на основі лексикону (Lexicon-Based)

Лексиконно-орієнтовані методи базуються на словниках, які пов'язують слова з полярністю або емоціями. У межах лексиконного методу виділяють два підходи. Перший – це dictionary-based підхід, який формує лексикони від базових «seed words» та використовує онлайн-словники для їх розширення. Прикладом є NRC EmoLex, який пов'язує понад 14 тис. слів з восьмима емоціями та має українську версію. У роботі [3] EmoLex було використано для аналізу економічних текстів до/під час COVID-19. Для частотного аналізу застосовано Sketch Engine, а загальну тональність – Lingmotif 2. Дослідження виявило зміни емоційного фону публікацій у кризовий період. Другий – Corpus-based підхід, що аналізує контекст вживання слів у корпусах. У дослідженні [4] поєднання цього підходу з SVM дало високу точність аналізу, а також підтвердило залежність якості від повноти лексикону.

#### 2. Підхід на основі правил (Rule-Based)

Підхід на основі правил поєднує лексикон з набором лінгвістичних правил, що дозволяють враховувати граматику, заперечення, підсилювачі та контекстуальні зміни полярності. На відміну від лексиконного підходу, де кожне слово має фіксовану емоційну оцінку, Rule-Based система може адаптувати цю оцінку залежно від мовної структури речення.

Одним із найпоширеніших інструментів є VADER, який враховує інтенсивність емоцій через пунктуацію, капіталізацію та модифікатори. У дослідженні [5] за допомогою VADER проаналізовано понад мільйон твітів про COVID-19 у Нігерії. Результати показали: позитивна тональність – 39,8 %, нейтральна – 31,3 %, негативна – 28,9 %, що перевищило точність TextBlob у виявленні полярності.

Інший приклад – LIWC (Linguistic Inquiry and Word Count), який класифікує слова за психологічними категоріями. У роботі [6] LIWC порівнювали з SVM при класифікації індонезійських твітів на теми етнічності й релігії. LIWC був ефективнішим для позитивної тональності ( $F1 = 0.69$ ), тоді як SVM краще визначав негативну ( $F1 = 0.72$ ). Додавання нових термінів до словника покращило середнє значення  $F1$  до 0.64.

Хоча rule-based методи поступаються за гнучкістю глибокому навчанню, вони цінні завдяки прозорості логіки, інтерпретованості та відносній простоті адаптації. Вони також можуть бути інтегровані як компоненти гібридних систем, підвищуючи точність в умовах обмежених ресурсів або у специфічних доменах.

#### 3. Підхід машинного навчання

Машинне навчання (ML) є одним з провідних підходів до аналізу тональності тексту, адже дає змогу виявляти приховані закономірності у великих обсягах текстових даних без необхідності ручного формалізування лінгвістичних правил. Залежно від типу навчання, ці підходи поділяються на контрольоване навчання, неконтрольоване навчання та напівконтрольоване навчання [7].

**Контрольоване навчання** є найбільш розповсюдженим підходом у задачах аналізу тональності тексту, що базується на використанні розмічених корпусів, де кожному фрагменту тексту присвоєна мітка емоційної полярності. Метою є побудова моделі класифікації, яка здатна узагальнювати залежності між текстовими ознаками та відповідною тональністю.

Серед лінійних моделей ефективність демонструє метод опорних векторів (SVM), що добре працює з малими наборами даних. У дослідженні [8] SVM перевершив Random Forest за точністю (80,39 % проти 78,56 %), підтвердивши свою ефективність у задачах бінарної класифікації з чітким розмежуванням класів.

Логістична регресія, незважаючи на простоту, також показує конкурентні результати. Зокрема, у роботі [9] вона досягла точності 84,92 % при класифікації англійських твітів, продемонструвавши її доцільність для багатомовних даних.

Штучні нейронні мережі (ANN) використовуються для моделювання складних нелінійних залежностей. У роботі [10] покращена модель з механізмом уваги перевершила базові ANN та SVM, забезпечивши зростання точності та  $F1$ -міри на корпусах із соціальних мереж.

Імовірнісні методи, як-от наївний байєс (NB), завдяки простоті реалізації та високій швидкості, залишаються популярними у прикладних задачах. Наприклад, у [11] застосували NB для класифікації скарг громадян в Індонезії, досягнувши 92 % точності, що демонструє його ефективність для адміністративних цілей.

Метод максимальної ентропії (MaxEnt) дозволяє працювати з неоднорідними ознаками без припущення про їхню незалежність. У дослідженні [12], присвяченому аналізу відгуків про застосунок Shopee, модель MaxEnt досягла точності 97,32 %, продемонструвавши свою високу ефективність для класифікації користувацьких оцінок.

Серед сучасних глибоких моделей CNN та LSTM показують високу продуктивність: відповідно 90,2 % на Yelp [13] та 86,85 % на IMDB [14]. LSTM ефективно працює з довгими послідовностями, а CNN – з локальними патернами у тексті.

Значний прорив забезпечили трансформерні моделі. У роботі [15] BERT досяг точності до 89 % на Twitter, F1-міра для позитивного класу склала 0.88. У контексті української мови, робота [16] порівняла кілька трансформерних архітектур, серед яких XLM-RoBERTa досягла 91,32 %, а Ukr-RoBERTa – 91,18 % із меншою потребою в ресурсах.

Окремо варто відзначити GPT. У дослідженні [17] модель SentimentGPT досягла точності 94,1 % та дозволила не лише класифікувати, а й генерувати пояснення до прийнятих рішень. Архітектура поєднує prompt-based взаємодію, fine-tuning GPT на розмічених корпусах та інтеграцію ембедінгів у класичні ML-моделі. Підсумовуючи, методи контрольованого навчання демонструють різний баланс між точністю, інтерпретованістю та ресурсомісткістю. SVM і логістична регресія є ефективними при роботі з невеликими наборами даних, тоді як дерева рішень забезпечують зручність інтерпретації, але менш стабільні. Глибокі нейронні мережі, особливо трансформери, значно підвищують якість класифікації, хоча вимагають значних обчислювальних ресурсів. У практичних застосуваннях вибір методу часто обумовлений потребою в швидкості, доступності ресурсів та вимогах до пояснюваності.

**Неконтрольоване навчання** використовується для аналізу текстів без наявності міток, що дозволяє виявляти приховані закономірності та структури в даних. Цей підхід є особливо корисним для кластеризації громадських звернень, тематичного моделювання та виявлення трендів у суспільних настроях. До основних методів належать алгоритми кластеризації, такі як K-Means [18] і DBSCAN [19], які групують схожі за змістом документи, тематичне моделювання за допомогою Latent Dirichlet Allocation (LDA) [20]. K-Means підходить для кластерів з чіткою формою, тоді як DBSCAN краще справляється з кластерами довільної форми. LDA є ефективним для тематичного моделювання, хоча ігнорує синтаксис і потребує ретельного налаштування кількості тем. Ці методи доцільні при первинному аналізі великих обсягів неструктурованих текстів, особливо у випадках виявлення нових тем або трендів.

**Напівконтрольоване навчання** поєднує обидва підходи, використовуючи як мічені, так і не мічені дані. Це дозволяє використовувати великий масив громадських звернень, навіть якщо лише частина з них має вказану тональність. У цьому контексті широко застосовуються методи самонавчання (self-training) та ко-тренування [21], що дозволяють поступово розширювати обсяг мічених прикладів. Зокрема, в гібридних системах поєднуються генеративні та графові підходи (GNN) [22], що дозволяє зберегти баланс між узагальненням і точністю. У мультимедійних підходах текст представлено в різних векторних просторах (словникових, синтаксичних, ембедінгових), що дозволяє моделі синтезувати знання з різних джерел та формувати більш стійке рішення. Таблиця 3 відображає підсумковий опис переваг та недоліків методів напівконтрольованого навчання. Отже, Self-training є найпростішим, але схильним до накопичення помилок. Co-training зменшує цей ефект, однак вимагає наявності кількох незалежних представлень даних. Новітні підходи, як-от GNN, демонструють перспективи в обробці складних графових структур, хоча їх результативність істотно визначається тим, наскільки точно змодельовано взаємозв'язки між об'єктами.

#### 4. Гібридний підхід

Гібридний підхід у завданнях аналізу тональності тексту являє собою інтеграцію методів, що належать до різних концептуальних парадигм, зокрема лексиконного, rule-based та машинного навчання. Основна ідея полягає у поєднанні переваг кожного з підходів з метою підвищення точності класифікації, стійкості до шумів, гнучкості у роботі з різними мовами та доменами. Така комбінація дає змогу ефективніше обробляти неоднозначні або контекстно залежні висловлювання, які є складними для аналізу за допомогою окремих підходів. У гібридних системах лексикони можуть використовуватись для початкової фільтрації або вагового моделювання ознак, що подаються до класифікатора машинного навчання. З іншого боку, rule-based механізми можуть слугувати джерелом апіорного знання або бути інтегрованими як евристики при формуванні навчальної вибірки. Така синергія дозволяє компенсувати обмеження окремих підходів і покращити загальну ефективність системи. Так, дослідження [23], представило гібридний ансамблевий підхід для аналізу тональності коротких текстів, інтегруючи глибоке навчання з традиційними методами машинного навчання. Використовуючи комбінацію моделей LSTM, Random Forest та логістичної регресії, дослідники досягли точності 83.24 %, що перевищує результати окремих моделей. Практичним підсумком проведеного аналізу є сформована класифікаційна схема (Рис. 1) яка систематизує основні напрями в межах цієї задачі.

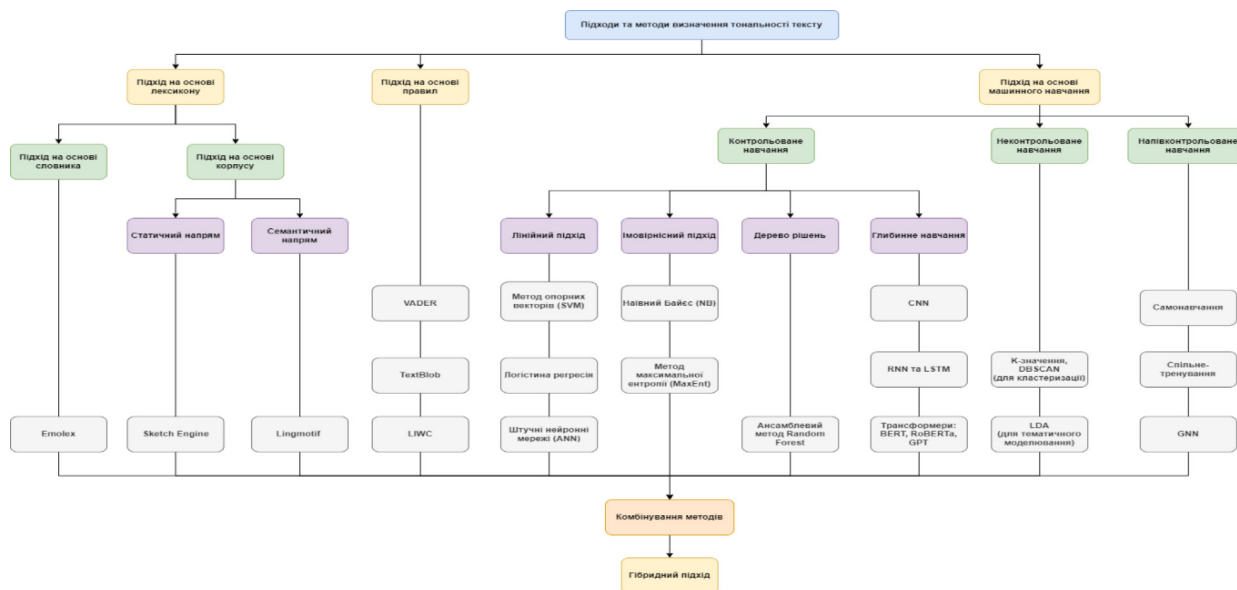


Рис. 1. Класифікаційна схема підходів та методів визначення тональності тексту

### Висновки та перспективи

У межах дослідження сформовано класифікаційну схему підходів до визначення тональності тексту, яка узагальнює сучасні методи та враховує специфіку українськомовного контенту. Запропонована структура є складовою інтелектуальної системи підтримки прийняття рішень на основі аналізу звернень мешканців, що дозволяє виявляти суспільні настрої та критичні проблеми.

Подальші дослідження передбачають локалізацію глибоких моделей (BERT, RoBERTa, GPT), створення гібридних підходів із залученням онтологічних знань, аналіз динаміки емоційного фону, а також інтеграцію тонального аналізу з тематичним, просторовим і суб'єктивним. Особливу увагу слід приділити методам активного та напівконтрольованого навчання для ефективної роботи з частково розміченими зверненнями.

Класифікація, запропонована у роботі, окреслює орієнтири подальшого розвитку аналітичних систем у сфері цифрового врядування та громадської участі в управлінні містом.

### Список використаної літератури

1. Cambria, E., Schuller, B., Xia, Y., & Havasi, C. (2013). New avenues in opinion mining and sentiment analysis. *IEEE Intelligent Systems*, 28(2), 15–21.
2. Kotsyba, N., Romanyshyn, O., & Shevchuk, O. (2021). Challenges in Sentiment Analysis for the Ukrainian Language. *CEUR Workshop Proceedings*, 2917, 118–126.
3. Vargas-Sierra, C., & Orts, M. (2023). *Sentiment and emotion in financial journalism: a corpus-based, cross-linguistic analysis of the effects of COVID*. Humanities and Social Sciences Communications, 10. <https://doi.org/10.1057/s41599-023-01725-8>
4. Abdulla, N. A., Ahmed, N. A., Shehab, M. A., & Al-Ayyoub, M. (2013). *Arabic sentiment analysis: lexicon-based and corpus-based*. IEEE Jordan Conference on Applied Electrical Engineering and Computing Technologies (AEECT).
5. Abiola, O., Abayomi-Alli, A., Tale, O.A. et al. (2023). *Sentiment analysis of COVID-19 tweets from selected hashtags in Nigeria using VADER and Text Blob analyser*. Journal of Electrical Systems and Information Technology, 10, 5. <https://doi.org/10.1186/s43067-023-00070-9>
6. Karyawati, A. E., Utomo, P. A., & Wibawa, I. G. A. (2022). Comparison of SVM and LIWC for Sentiment Analysis of SARA. *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 16(1), 45–54.
7. Birjali, M., Kasri, M., & Beni-Hssane, A. (2021). A comprehensive survey on sentiment analysis: Approaches, challenges and trends. *Knowledge-Based Systems*, 226, 107134.
8. Khan, M. J., Abbas, Q., & Hussain, M. (2024). *Comparative study of supervised learning techniques for sentiment classification*. Computers, Materials & Continua, 78(1), 1–16. <https://doi.org/10.32604/cmc.2023.043301>
9. Aliman, N. M., Mustaffa, M., Jambari, H., & Jusoh, M. S. (2022). *Performance evaluation of supervised learning algorithms for sentiment analysis: A comparative study*. Journal of Engineering and Applied Sciences, 17(5), 143–150.
10. Lu, W., Wang, J., Wang, Y., Zhou, Y., & Qin, H. (2021). *Sentiment analysis of social media texts with deep learning models and attention mechanism*. Information, 12(10), 392. <https://doi.org/10.3390/info12100392>

11. Wabang, G. S., Ahmad, T., & Wijaya, D. E. (2022). *Application of the Naive Bayes Classifier Algorithm to Classify Community Complaints*. Journal of Physics: Conference Series, 2180(1), 012045. <https://doi.org/10.1088/1742-6596/2180/1/012045>
12. Rhohmawati, A., Sari, R. F., & Puspitasari, R. (2019). *Sentiment analysis using maximum entropy for Shopee reviews*. In 2019 4th International Conference on Information Technology, Information Systems and Electrical Engineering (ICITISEE) (pp. 237–242). IEEE.
13. Kim, J., & Jeong, Y. (2019). *Sentiment classification using convolutional neural networks*. International Journal of Advanced Computer Science and Applications, 10(5), 303–308.
14. Tholusuri, H., Gadde, P. R., & Sista, S. R. (2019). *Sentiment analysis using LSTM with IMDB dataset*. In 2019 4th International Conference on Communication and Electronics Systems (ICCES) (pp. 1192–1195). IEEE.
15. Roccabruna, S., Nesi, P., & Pantaleo, G. (2022). *A comparison of BERT-based models for sentiment analysis in social media*. Applied Sciences, 12(6), 3016. <https://doi.org/10.3390/app12063016>
16. Prytula, Y. (2024). *Evaluation of transformer-based models for sentiment analysis of Ukrainian texts*. Proceedings of the International Conference on Computational Linguistics and Intelligent Systems (COLINS), 2024.
17. Kheiri, M., & Karimi, M. (2023). *SentimentGPT: Prompt-based sentiment analysis using generative pre-trained transformers*. Journal of Big Data, 10, Article 42. <https://doi.org/10.1186/s40537-023-00708-4>
18. Riaz, S., Fatima, M., Kamran, M., & Nisar, M. W. (2019). *Opinion mining on large scale data using sentiment analysis and k-means clustering*. Cluster Computing, 22, 7149–7164.
19. Andriyani, F., & Puspitarani, Y.. (2022). *Performance Comparison of K-Means and DBScan Algorithms for Text Clustering Product Reviews*. Sinkron : Jurnal Dan Penelitian Teknik Informatika, 6(3), 944–949. <https://doi.org/10.33395/sinkron.v7i3.11569>
20. Xue, J., Chen, J., Chen, C., Zheng, C., Li, S., & Zhu, T. (2020). *Public discourse and sentiment during the COVID-19 pandemic: Using Latent Dirichlet Allocation for topic modeling on Twitter*. PloS one, 15(9), e0239441.
21. Riaz, S., Fatima, M., Kamran, M., & Nisar, M. W. (2019). *Opinion mining on large scale data using sentiment analysis and k-means clustering*. Cluster Computing, 22, 7149–7164.
22. Andriyani, F., & Puspitarani, Y.. (2022). *Performance Comparison of K-Means and DBScan Algorithms for Text Clustering Product Reviews*. Sinkron : Jurnal Dan Penelitian Teknik Informatika, 6(3), 944–949. <https://doi.org/10.33395/sinkron.v7i3.11569>
23. Xue, J., Chen, J., Chen, C., Zheng, C., Li, S., & Zhu, T. (2020). *Public discourse and sentiment during the COVID-19 pandemic: Using Latent Dirichlet Allocation for topic modeling on Twitter*. PloS one, 15(9), e0239441.

## References

1. Cambria, E., Schuller, B., Xia, Y., & Havasi, C. (2013). *New avenues in opinion mining and sentiment analysis*. IEEE Intelligent Systems, 28(2), 15–21.
2. Kotsyba, N., Romanyshyn, O., & Shevchuk, O. (2021). *Challenges in Sentiment Analysis for the Ukrainian Language*. CEUR Workshop Proceedings, 2917, 118–126.
3. Vargas-Sierra, C., & Orts, M. (2023). *Sentiment and emotion in financial journalism: a corpus-based, cross-linguistic analysis of the effects of COVID*. Humanities and Social Sciences Communications, 10. <https://doi.org/10.1057/s41599-023-01725-8>
4. Abdulla, N. A., Ahmed, N. A., Shehab, M. A., & Al-Ayyoub, M. (2013). *Arabic sentiment analysis: lexicon-based and corpus-based*. IEEE Jordan Conference on Applied Electrical Engineering and Computing Technologies (AEECT).
5. Abiola, O., Abayomi-Alli, A., Tale, O.A. et al. (2023). *Sentiment analysis of COVID-19 tweets from selected hashtags in Nigeria using VADER and Text Blob analyser*. Journal of Electrical Systems and Information Technology, 10, 5. <https://doi.org/10.1186/s43067-023-00070-9>
6. Karyawati, A. E., Utomo, P. A., & Wibawa, I. G. A. (2022). *Comparison of SVM and LIWC for Sentiment Analysis of SARA*. IJCCS (Indonesian Journal of Computing and Cybernetics Systems), 16(1), 45–54.
7. Birjali, M., Kasri, M., & Beni-Hssane, A. (2021). *A comprehensive survey on sentiment analysis: Approaches, challenges and trends*. Knowledge-Based Systems, 226, 107134.
8. Khan, M. J., Abbas, Q., & Hussain, M. (2024). *Comparative study of supervised learning techniques for sentiment classification*. Computers, Materials & Continua, 78(1), 1–16. <https://doi.org/10.32604/cmc.2023.043301>
9. Aliman, N. M., Mustafa, M., Jambari, H., & Jusoh, M. S. (2022). *Performance evaluation of supervised learning algorithms for sentiment analysis: A comparative study*. Journal of Engineering and Applied Sciences, 17(5), 143–150.
10. Lu, W., Wang, J., Wang, Y., Zhou, Y., & Qin, H. (2021). *Sentiment analysis of social media texts with deep learning models and attention mechanism*. Information, 12(10), 392. <https://doi.org/10.3390/info12100392>
11. Wabang, G. S., Ahmad, T., & Wijaya, D. E. (2022). *Application of the Naive Bayes Classifier Algorithm to Classify Community Complaints*. Journal of Physics: Conference Series, 2180(1), 012045. <https://doi.org/10.1088/1742-6596/2180/1/012045>

12. Rhohmawati, A., Sari, R. F., & Puspitasari, R. (2019). *Sentiment analysis using maximum entropy for Shopee reviews*. In 2019 4th International Conference on Information Technology, Information Systems and Electrical Engineering (ICITISEE) (pp. 237–242). IEEE.
13. Kim, J., & Jeong, Y. (2019). *Sentiment classification using convolutional neural networks*. International Journal of Advanced Computer Science and Applications, 10(5), 303–308.
14. Tholusuri, H., Gadde, P. R., & Sista, S. R. (2019). *Sentiment analysis using LSTM with IMDB dataset*. In 2019 4th International Conference on Communication and Electronics Systems (ICCES) (pp. 1192–1195). IEEE.
15. Roccabruna, S., Nesi, P., & Pantaleo, G. (2022). *A comparison of BERT-based models for sentiment analysis in social media*. Applied Sciences, 12(6), 3016. <https://doi.org/10.3390/app12063016>
16. Prytula, Y. (2024). *Evaluation of transformer-based models for sentiment analysis of Ukrainian texts*. Proceedings of the International Conference on Computational Linguistics and Intelligent Systems (COLINS), 2024.
17. Kheiri, M., & Karimi, M. (2023). *SentimentGPT: Prompt-based sentiment analysis using generative pre-trained transformers*. Journal of Big Data, 10, Article 42. <https://doi.org/10.1186/s40537-023-00708-4>
18. Riaz, S., Fatima, M., Kamran, M., & Nisar, M. W. (2019). Opinion mining on large scale data using sentiment analysis and k-means clustering. Cluster Computing, 22, 7149–7164.
19. Andriyani, F., & Puspitarani, Y.. (2022). Performance Comparison of K-Means and DBScan Algorithms for Text Clustering Product Reviews. Sinkron : Jurnal Dan Penelitian Teknik Informatika, 6(3), 944–949. <https://doi.org/10.33395/sinkron.v7i3.11569>
20. Xue, J., Chen, J., Chen, C., Zheng, C., Li, S., & Zhu, T. (2020). Public discourse and sentiment during the COVID-19 pandemic: Using Latent Dirichlet Allocation for topic modeling on Twitter. PloS one, 15(9), e0239441.
21. Riaz, S., Fatima, M., Kamran, M., & Nisar, M. W. (2019). Opinion mining on large scale data using sentiment analysis and k-means clustering. Cluster Computing, 22, 7149–7164.
22. Andriyani, F., & Puspitarani, Y.. (2022). Performance Comparison of K-Means and DBScan Algorithms for Text Clustering Product Reviews. Sinkron : Jurnal Dan Penelitian Teknik Informatika, 6(3), 944–949. <https://doi.org/10.33395/sinkron.v7i3.11569>
23. Xue, J., Chen, J., Chen, C., Zheng, C., Li, S., & Zhu, T. (2020). Public discourse and sentiment during the COVID-19 pandemic: Using Latent Dirichlet Allocation for topic modeling on Twitter. PloS one, 15(9), e0239441.