

ОВЧАРЕНКО МАКСИМ

Національний технічний університет «Дніпровська політехніка»

<https://orcid.org/0009-0006-0730-0913>e-mail: ovcharenko.m.a@nmu.one**ІВАНЬКО АРТЕМ**

Національний технічний університет «Дніпровська політехніка»

<https://orcid.org/0009-0002-0491-5374>e-mail: ivanko.a.m@nmu.one**КАШТАН ВІТА**

Національний технічний університет «Дніпровська політехніка»

<https://orcid.org/0000-0002-0395-5895>e-mail: kashtan.v.yu@nmu.one

ВИЗНАЧЕННЯ ЕФЕКТИВНОСТІ НЕЙРОМЕРЕЖЕВИХ МОДЕЛЕЙ ДЛЯ АНАЛІЗУ ЕМОЦІЙНОГО СТАНУ КОРИСТУВАЧІВ У СОЦІАЛЬНИХ МЕРЕЖАХ

У роботі представлено дослідження ефективності нейромережеских моделей для аналізу емоційного стану користувачів у соціальних мережах на основі текстових даних із Twitter, пов'язаних з COVID-19. Розглянуто моделі нейронних мереж: Convolutional Neural Network (CNN), CNN з оптимізатором Adam, Random Forest, Extra-tree, XGBoost та BERT. Для кожної з моделей були оцінені показники точності, повноти, F1-Score та помилкових класифікацій, зокрема True Positives, True Negatives, False Positives та False Negatives. Результати дослідження показали, що модель BERT продемонструвала найвищу точність та ефективність для класифікації емоцій, досягнувши точності 86%. Моделі CNN та CNN з оптимізатором Adam показали результат 83%, що є нижчим, ніж у моделі BERT, але вищим порівняно з іншими моделями. Моделі Random Forest, Extra-tree та XGBoost показали значно нижчі показники точності та повноти, що свідчить про обмежену ефективність цих моделей для розпізнавання емоцій за текстовими наборами даних.

Ключові слова: аналіз емоцій, нейронні моделі, CNN, BERT, Random Forest, XGBoost, точність класифікації, оптимізатор Adam.

OVCHARENKO MAKSYM**IVANKO ARTEM****VITA KASHTAN**

Dnipro University of Technology

EVALUATION OF THE NEURAL NETWORK MODEL'S EFFICIENCY IN ANALYZING THE USERS' EMOTIONAL MOOD IN SOCIAL NETWORKS

This study investigates the effectiveness of various neural network models for sentiment analysis, specifically for detecting the emotional states of users based on textual data from social media platforms, focusing on Twitter posts related to COVID-19. The research evaluates the performance of six models: Convolutional Neural Network (CNN), CNN with the Adam optimizer, Random Forest, Extra-tree, XGBoost, and BERT. The models were assessed using key evaluation metrics such as accuracy, precision, recall, F1-score, and confusion matrices to measure their ability to classify positive and negative emotions in user-generated text.

The experimental results demonstrated that the BERT model achieved the highest performance, with % overall accuracy of 86%. This result highlights BERT's superior ability to accurately classify emotions with minimal misclassification, making it an optimal choice for sentiment analysis in social media data. BERT's performance can be attributed to its deep contextual understanding of the text, which is essential for accurately interpreting emotional tones in complex language.

The CNN and CNN with the Adam optimizer models achieved a general accuracy of 83%, slightly lower than BERT but still outperforming the traditional machine learning models such as Random Forest and XGBoost. Notably, including the Adam optimizer improved the F1-score for the positive class, indicating better performance in detecting positive sentiments. Despite this improvement, these models still lag behind BERT regarding overall accuracy and efficiency.

In contrast, Random Forest, Extra-tree, and XGBoost models showed significantly lower results. Specifically, Random Forest and XGBoost exhibited poor accuracy in detecting positive emotions, with 57% and 54% precision scores, respectively. It suggests limited effectiveness in distinguishing between positive and negative emotions in the analyzed texts. Extra-tree, while outperforming Random Forest and XGBoost, still failed to reach the level of accuracy demonstrated by CNN and BERT, with a general accuracy of 68%.

The findings of this study confirm that deep learning models, particularly BERT, are highly effective for sentiment analysis tasks, especially when working with complex social media data. These models can capture intricate patterns and contextual nuances that traditional machine learning models struggle to identify, making them the preferred choice for tasks related to emotion detection in text. The research underscores the need to select appropriate algorithms based on the specific characteristics of the dataset and task requirements.

Keywords: emotion analysis, neural models, CNN, BERT, Random Forest, XGBoost, classification accuracy, Adam optimizer.

Стаття надійшла до редакції / Received 11.04.2025

Прийнята до друку / Accepted 26.04.2025

Постановка проблеми

В епоху цифрових технологій соціальні мережі стали основними каналами для комунікації, обміну думками, рекомендаціями та отриманням інформації. Великі обсяги даних, що створюються користувачами в соціальних мережах, містять цінну інформацію про їхні настрої, думки та емоції [1].

Такий обсяг неструктурованої інформації, як правило, включає тексти, коментарі, зображення та відео, що може стати важливим джерелом для розуміння потреб користувачів, їхнього емоційного стану та взаємодії з певними подіями чи контентом. Аналіз цих емоційних реакцій може допомогти не тільки в розумінні настроїв і потреб користувачів, але й у створенні персоналізованих маркетингових стратегій, поліпшенні клієнтського сервісу та вдосконаленні продуктів.

Аналіз емоційного стану користувачів соціальних мереж є складним завданням через такі фактори як: по-перше, щодня в соціальних мережах генеруються мільярди повідомлень, коментарів та зображень, що ускладнює їхню обробку та аналіз. По-друге, більшість даних у соціальних мережах є неструктурованими, тобто вони не мають чіткої структури та формату, що ускладнює їхній аналіз [2]. По-третє, емоції користувачів є суб'єктивними та контекстуально залежними, що ускладнює їхнє автоматизоване визначення. Крім того, користувачі соціальних мереж використовують різні мови, діалекти та сленг, що ускладнює аналіз тексту. Нарешті, емоційні реакції користувачів часто виражаються не лише через текстові повідомлення, але й через міміку, зображення та відео, що вимагає інтеграції методів обробки зображень та відео з методами аналізу тексту.

Завдяки розвитку машинного навчання, обробки природної мови та технологій глибинного навчання, стало можливим створення неймережевих моделей, здатних ефективно класифікувати емоційні стани користувачів на основі їхніх повідомлень та реакцій у соціальних мережах [3]. Однак, незважаючи на значний прогрес, який досягнутий у цій сфері, існують суттєві труднощі у точності моделей через складність обробки неструктурованих даних, що складають основну частину інформації в Інтернеті.

Тому для вирішення вище зазначених проблем необхідно розробити більш ефективні методи аналізу емоційного стану користувачів соціальних мереж, що зможуть точно та швидко обробляти великі обсяги неструктурованої інформації, виявляти приховані закономірності та забезпечувати високу точність класифікації емоційних реакцій. Неймережеві моделі, завдяки своїй здатності до навчання на великих даних та виявлення складних залежностей, є перспективним інструментом для розв'язання цих завдань.

Аналіз останніх публікацій

Для аналізу емоційного стану користувачів соціальних мереж було запропоновано ряд методів на основі машинного навчання для класифікації та аналізу текстових даних, що генеруються в соціальних мережах. Дослідження, представлене в [4] присвячено розробці спеціальної платформи для аналізу настроїв користувачів соціальних мереж. Платформа автоматизує завдання оцінки емоційних реакцій на публікації, генеруючи звіти на основі результатів. Вона дозволяє виконувати аналіз настроїв щодо всіх видів активності користувачів, що допомагає адміністраторам приймати обґрунтовані рішення, а також ідентифікувати користувачів, які потребують особливої уваги, зокрема у випадках негативних емоцій. У роботі [5] дослідники обговорюють потенціал використання даних соціальних медіа для створення прибуткових бізнес-моделей. Вони наголошують на важливості застосування машинного навчання та методів обробки природної мови для аналізу думок і ставлень користувачів, зокрема у маркетингових цілях. Важливою частиною їхнього дослідження є аналіз перспектив глибокого навчання та оптимізаційних методів для вдосконалення моделей аналізу настроїв. Автори в роботі [6] досліджували застосування методів машинного навчання для аналізу настроїв у контексті даних з Twitter, пов'язаних з податком на товари та послуги (GST). За допомогою веб-збору було зібрано текстові набори, що висвітлюють позитивні, негативні та нейтральні емоції користувачів. Дослідники застосували методи візуалізації даних для пошуку закономірностей у настроях користувачів, що дозволяє прогнозувати майбутні тенденції на основі поточного використання платформи.

Традиційні методи для аналізу настроїв користувачів, як правило, базуються на ручному вилученні ознак та їх подальшій обробці, що є трудомістким процесом. Натомість, глибоке навчання, зокрема, використання нейронних мереж, стало потужною альтернативою традиційним технікам машинного навчання. Це дозволяє досягати значних покращень у продуктивності та точності результатів аналізу настроїв [7].

Методи, засновані на машинному навчанні, використовують тренування алгоритму на попередньо підготовленому наборі даних, після чого застосовують його до нових, раніше не відомих даних. У 2016 році було проведено порівняльне дослідження між підходами на основі лексики та методами машинного навчання [8]. У рамках цього порівняння були застосовані різні алгоритми, зокрема наївний Байєс (NB), машини опорних векторів (SVM) та максимальна ентропія, на наборі даних з Twitter. Результати експериментів показали, що точність наївного Байєса становить 74,44%, а точність методу SVM 77,73%. Таким чином, алгоритм SVM продемонстрував кращі результати порівняно з іншими методами.

У роботі [9] автори застосували методи глибокого навчання та алгоритм лінійного машинного навчання для класифікації емоційних настроїв. Крім того, були запропоновані два методи, які комбінують стандартні класифікатори з іншими підходами, що часто використовуються в аналізі настроїв [10]. Зокрема, дослідники розглянули два способи об'єднання глибоких навчальних методів з іншими техніками для інтеграції інформації з численних джерел даних. Окрім цього, класифікація

різних моделей була здійснена за категоріями, відповідно до наукової літератури. Для оцінки ефективності запропонованих підходів було проведено кілька тестів на продуктивність, що дозволило виміряти їхню ефективність.

У роботі [11] автори представили метод, що використовує згорткові нейронні мережі (CNN) та довготривалу короткочасну пам'ять (LSTM) для аналізу фрагментів тексту. Цей метод дозволяє обробляти дані на основі векторів, які представляють слова, та ефективно впоратися з довгостроковими залежностями в текстових наборах. Результати показали, що CNN дозволяє знижувати ризик надмірної інформації та мінімізувати ймовірність виникнення перезавантаження даних. Запропонована техніка виявила свою ефективність у виявленні емоційних реакцій та мінімізації витрат на обчислювальні ресурси. В роботі [12] вивчали комбінацію машинного навчання та підходів на основі лексикону, застосованих до даних з Twitter та Facebook. Проведені експерименти показали, що використання лексикографічних методів в поєднанні з попередньою обробкою даних значно підвищує точність аналізу. Авторі дійшли висновку, що методи на основі лексикону є дуже ефективними у порівнянні з підходами, заснованими на машинному навчанні, оскільки дозволяють враховувати аспекти семантики тексту для оцінки емоційного стану.

Розглянуті вищеописані методи аналізу емоційного стану користувачів соціальних мереж, включаючи як традиційні, так і новітні підходи на основі глибокого навчання відкривають перспективи для автоматизованого аналізу текстових даних. Але незважаючи на стрімкий розвиток технологій глибокого навчання, необхідність в оцінці ефективності різних нейромережевих моделей для аналізу емоцій користувачів соціальних мереж залишається актуальною.

Формулювання цілей статті

Метою роботи є проведення порівняльного аналізу ефективності нейромережевих моделей, а саме CNN, CNN+Adam, Random Forest, Extra-tree, XGBoost та BERT для визначення емоційного стану користувачів на основі текстових даних із соціальних мереж, зокрема Twitter, пов'язаних з COVID-19. Завданнями дослідження є оцінка точності та продуктивності кожної з моделей при класифікації емоційних настроїв (позитивних та негативних) в текстових даних; порівняння результатів роботи нейронних моделей; дослідження використання параметрів оптимізатора Adam на показники класифікації емоційних станів у нейромережевих моделях.

Виклад основного матеріалу

Для порівняльного аналізу нейромережевих моделей, що застосовуються для визначення емоційного стану користувачів соціальних мереж, було розроблено підхід, зосереджений на аналізі текстових наборів даних, що стосуються пандемії COVID-19 та вакцини проти неї. У цьому дослідженні зібрано текстові набори даних з соціальної мережі Twitter, що містить текстову інформацію, хештеги та інші елементи контенту, які відображають настрої користувачів з приводу пандемії та вакцинації. Текстові набори даних, отримані з соціальної мережі Twitter містять максимально 280 символів, що дозволяє отримувати короткі, але змістовні повідомлення про громадську думку. В роботі використано близько 80 000 текстових наборів про COVID-19, опублікованих протягом трьох місяців з березня 2020 року по травень 2020 року [13]. Хештеги, які є важливою частиною контенту у Twitter, також виступають у ролі маркерів для класифікації настроїв і корисні для фільтрації релевантних даних. Запропонована методологія аналізу емоційного стану користувачів на основі нейромережевих моделей показана на рис. 1.



Рис. 1. Структурна схема процесу аналізу настроїв користувачів на основі нейромережевих моделей

Дані, отримані з соціальної мережі Twitter, не можуть бути безпосередньо використані для підготовки моделей машинного навчання, так як такі дані часто містять друкарські та граматичні помилки, аббревіатури, сленгові слова, сарказм, смайлики, URL-адреси, зайві пробіли та багато інших

елементів, які ускладнюють точну інтерпретацію емоційного настрою користувачів. Тому попередня обробка даних є важливим етапом для аналізу настроїв користувачів. У цьому дослідженні для попередньої обробки даних був використаний вбудований модуль регулярних виразів Python «re», який має набір функцій для зіставлення шаблонів з регулярними виразами. Етапи попередньої обробки даних включають:

- видалення URL-адрес з текстових наборів, оскільки вони не несуть важливої емоційної інформації для аналізу настроїв;
- видалення емодзі та знаків пунктуації, які можуть створювати шум і спотворювати емоційний контекст;
- перетворення тексту в малі літери, що дозволяє уникнути дублювання слів, написаних з великої літери, та нормалізує текст для подальшої обробки;
- токенизація тексту, що дозволяє групувати текст на слова, фрази, що є необхідним для подальшого аналізу;
- видалення стоп-слів, які не несуть значущої інформації, були виключені з текстів;
- видалення хештегів так як це не містить важливої інформації для аналізу емоцій;
- видалення зайвих пробілів, що зменшує розмір тексту і полегшує подальшу обробку;
- перетворення емоційних позначень на текст, що дозволяє зберегти емоційну інформацію, яку вони містять.

Особливу увагу приділено обробці емодзі, оскільки їх використання в Twitter-повідомленнях є популярним. Перетворення емодзі на текстовий формат є важливим етапом для точного визначення емоційного стану користувачів, що неможливо правильно інтерпретувати без перетворення їх на текст. У таблиці 1 наведено фрагмент набору даних, що містить Twitter -повідомлення, їхні хештеги, а також маркери емоційного стану користувачів (позитивний, негативний), які були визначені за допомогою попередньої обробки та класифікації тексту. Ці дані служать основою для подальшого аналізу ефективності застосованих нейромережових моделей для визначення емоційного стану користувачів у соціальних мережах.

Таблиця 1

Фрагмент набору даних Twitter -повідомлень

Дата публікації	Текст твіту	Хештеги	Емоційний стан
02-03-2020	TRENDING: New Yorkers encounter empty supermarket shelves (pictured, Wegmans in Brooklyn), sold-out online grocers (FoodKick, MaxDelivery) as #coronavirus-fearing shoppers1 st	#coronavirus	Позитивний
02-03-2020	When I couldn't find hand sanitizer at Fred Meyer, I turned to #Amazon. But \$114.97 for a 2 pack of Purell?!! Check out how #coronavirus concerns are driving up prices.	#Amazon, #coronavirus	Негативний
02-03-2020	Find out how you can protect yourself and loved ones from #coronavirus. ?	#coronavirus	Позитивний
03-03-2020	Do you remember the last time you paid \$2.99 a gallon for regular gas in Los Angeles?	#coronavirus	Позитивний
03-03-2020	"@DrTedros " "We can't stop #COVID19 without protecting #healthworkers.	#COVID19, #coronavirus	Нейтральний
04-03-2020	HI TWITTER ! I am a pharmacist. I sell hand sanitizer for a living ! Or I do when any exists.	-	Позитивний
04-03-2020	Best quality couches at unbelievably low prices available to order.	#SuperTuesdsy, #Covid_19,	Позитивний

Після етапу попередньої обробки текстові дані необхідно перетворити на числовий формат, що дозволяє ефективно обробляти ці дані за допомогою нейромережових моделей. Цей процес називається векторизацією тексту. Одним із поширених методів векторизації є Word2Vec, який використовує нейронні мережі для створення векторних представлень слів у текстовому наборі даних. Слова, що часто зустрічаються в схожих контекстах, отримують близькі векторні представлення. В рамках аналізу настроїв, векторизація тексту за допомогою Word2Vec дозволяє моделі не лише розпізнавати окремі слова, а й враховувати їх контекст у конкретному тексті. Це дозволяє точніше визначати емоційний настрій тексту, оскільки аналізуються не лише окремі слова, але й їхні взаємозв'язки.

Наступним етапом є фільтрація ознак для підвищення продуктивності та точності нейромережових моделей. Процес фільтрації ознак передбачає усунення з набору даних нерелевантних та дублюючих атрибутів, які не впливають на ефективність і точність прогнозуючої моделі. Для фільтрації ознак застосовано статистичний метод CHI (хі-квадрат), який оцінює значення кожної

ознаки шляхом обчислення хі-квадрат для кожного класу. Метод аналізує залежність між термінами та класами: значення 0 вказує на відсутність залежності між терміном і класом, а значення 1 свідчить про наявність залежності. Застосування методу CHI дозволило виокремити найбільш важливі ознаки та забезпечити коректну оцінку різниці значущості між категоріям [14]:

$$CHI(t, c_i) = \frac{N \cdot (AD - BE)^2}{(A + E) \cdot (B + D) \cdot (A + B) \cdot (E + D)}, \quad (1)$$

$$CHI_{max}(t) = \max_i (CHI(t, c_i)), \quad (2)$$

де A – частота, коли слово t і клас c_i зустрічаються разом; B – кількість випадків, коли слово t з'являється без класу c_i ; C – кількість випадків, коли клас c_i з'являється без слова t ; D – частота, коли ні слово t , ні клас c_i не з'являються; N – загальна кількість слів у текстовому наборі.

Етап аналізу емоційного стану полягає у визначенні суб'єктивного ставлення користувачів до обговорюваних тем. У даному дослідженні виконано порівняння ефективності нейромережових моделей: CNN, CNN+Adam, Random Forest, Extra-tree, XGBoost та BERT. Конволюційна нейронна мережа CNN працює з текстом як з послідовністю символів або слів, де фільтри (які є аналогами ядер згортки) застосовуються для виділення важливих ознак. Модель CNN автоматично знаходить ключові шаблони в тексті, що робить її надзвичайно ефективною для задач класифікації текстів, таких як визначення емоційного стану користувачів. Використано п'ять різних розмірів згорткових фільтрів $h=(2,3,4,5)$ з кількістю фільтрів 50, 50, 50, 30, 20 для кожного розміру фільтра відповідно. Для кожного згорткового фільтра використовується функція активації ReLU, що дозволяє моделі ефективно виявляти нелінійні зв'язки в даних. Після застосування згорткових фільтрів на текстові дані проведено операцію максимального об'єднання (max-pooling). Це дозволило виділити основні ознаки тексту незалежно від їхнього просторового розташування. Всі максимальні значення, отримані після застосування кожного фільтра, об'єднуються в один вектор, який далі передається через повністю зв'язаний прихований шар з 30 нейронами. На фінальному етапі результати, що виходять з цього шару, проходять через сигмоподібну активаційну функцію для прогнозування кінцевих класів, що відповідають заданим емоційним категоріям. З метою зменшення перенавчання моделі CNN введено механізм випадкового відключення нейронів з ймовірністю $P=0.3$, а також після повністю зв'язаного шару з ймовірністю $P=0.1$. Це дозволило зменшити надмірну адаптацію моделі до тренувальних даних і покращити її здатність до узагальнення на нові, невідомі дані.

Після застосування моделі на основі конволюційних нейронних мереж (CNN), наступним етапом є використання алгоритму випадкового лісу (Random Forest). Основна ідея алгоритму полягає в створенні ансамблю дерев рішень, кожне з яких навчається на випадково вибраній підмножині даних. Для класифікації нового текстового набору кожне дерево приймає рішення щодо належності до певного емоційного класу (наприклад, позитивного або негативного), і остаточний результат визначається шляхом агрегування прогнозів усіх дерев [14]:

$$ni_j = w_j C_j - w_{left(j)} C_{left(j)} - w_{right(j)} C_{right(j)}, \quad (3)$$

де ni_j – ознака j після поділу даних на категорії емоцій; w_j – вага ознаки j у загальному наборі текстових даних; C_j – ознаки j у загальному наборі даних до поділу на категорії емоцій; $w_{left(j)}$ – вага ознаки j після поділу даних на групу, яка має певну емоцію (негативну або позитивну); $C_{left(j)}$ – ознаки j для негативних емоцій після аналізу тексту; $w_{right(j)}$ – вага ознаки j для позитивних емоцій; $C_{right(j)}$ – ознаки j для позитивних емоцій після аналізу тексту.

Наступною моделлю, що використовується в дослідженні, є Extra-tree, яка є варіантом алгоритму випадкового лісу і має суттєві відмінності в процесі побудови дерев рішень. На відміну від Random Forest, де при побудові кожного дерева вибираються випадкові ознаки для кожного вузла, Extra-tree використовує випадковий вибір порогу для розподілу даних на кожному вузлі дерева. Це означає, що під час вибору розподілу в кожному вузлі модель не виконує детального аналізу всіх можливих порогових значень, а натомість випадковим чином обирає одне з них із множини допустимих варіантів. Така випадковість сприяє отриманню більш різноманітних дерев у ансамблі, що покращує результативність класифікації. Алгоритм Extra-Trees, подібно до стандартних дерев рішень, використовує ентропію для визначення найінформативнішого критерію класифікації категоріальних змінних. Важливим моментом є те, що класифікація ознак в Extra-tree здійснюється на основі випадкових значень, що дозволяє додатково знизити кореляцію між деревами і підвищити якість ансамблю.

Наступним етапом є використання моделі XGBoost, що використовує підхід градієнтного бустингу: кожне дерево додається з метою усунення помилок, допущених попередніми, що дозволяє суттєво підвищити точність класифікації. На початковому етапі кожне дерево формується з однаковими вагами для всіх значень вхідних даних. Потім ці ваги поступово оновлюються в залежності від помилок, допущених попередніми деревами. Після кожної побудови дерева рішень, проводиться модифікація ваг для кожної ознаки з метою мінімізації похибок у прогнозах. Важливість кожної ознаки коригується на кожному етапі навчання. Процес навчання XGBoost математично описано формулою

(4), що визначає класифікаційну модель на кожному етапі побудови дерева [15]:

$$\hat{y}_i^{(t)} = \sum_{k=1}^t f_k(x_i), \quad (4)$$

де $\hat{y}_i^{(t)}$ – це передбачення для елемента i після побудови t -го дерева; $f_k(x_i)$ – це вихід k -го дерева для елемента i , який класифікує емоцію; x_i – це вхідний елемент (текстові дані); t – кількість дерев, побудованих на поточному етапі навчання.

Наступною моделлю є BERT (Bidirectional Encoder Representations from Transformers), заснована на трансформерах, яка здатна обробляти контекст тексту з обох напрямків, що робить її надзвичайно ефективною для задач, пов'язаних з аналізом природної мови. Завдяки двосторонньому кодуванню слів, BERT здатна враховувати контекст навколо кожного слова, що дозволяє точніше інтерпретувати значення слів у тексті. Ця особливість робить BERT надзвичайно ефективним інструментом для задач класифікації текстів, зокрема для аналізу емоційного настрою, де важливий не лише контекст окремих слів, але й їх взаємозв'язки.

Після застосування вищеприписаних нейромережових моделей для аналізу емоційного стану користувачів на основі текстових даних з соціальних мереж, пов'язаних з COVID-19, необхідно провести оцінку їх ефективності та порівняти отримані результати. Для цього використано метрики класифікації: точність (precision), повнота (recall), f1-score та точність [14]. Точність (Precision) є однією з основних метрик для оцінки ефективності класифікаційних моделей, зокрема в задачах, де важливо мінімізувати кількість хибнопозитивних результатів. Точність вимірює здатність моделі коректно класифікувати позитивні випадки серед всіх випадків, які вона класифікує як позитивні. Повнота (Recall) це метрика, яка показує здатність моделі знаходити всі позитивні приклади, включаючи ті, що могли бути пропущені. Для поєднання точності та повноти в одну метрику використовують F1-Score – гармонійне середнє між точністю та повнотою, що дає можливість комплексно оцінити загальну ефективність моделі, враховуючи обидва параметри.

В таблиці 2 наведено результати порівняльного аналізу нейромережових моделей (CNN, CNN з оптимізатором Adam, Random Forest, Extra-tree та BERT) для визначення емоційного стану користувачів на основі текстових даних із соціальних мереж.

Таблиця 2

Метрики оцінки нейромережових моделей для визначення емоційного стану користувачів

Модель	Метрики				Емоційний стан
	Точність	Повнота	F1-Score	Загальна точність	
CNN	0.84	0.85	0.84	0.83	Позитивний
	0.82	0.81	0.82		Негативний
CNN з оптимізатором Adam	0.85	0.84	0.88	0.83	Позитивний
	0.82	0.81	0.82		Негативний
Random Forest	0.57	0.52	0.54	0.61	Позитивний
	0.63	0.68	0.65		Негативний
Extra-tree	0.66	0.59	0.62	0.68	Позитивний
	0.69	0.75	0.72		Негативний
XGBoost	0.54	0.52	0.53	0.59	Позитивний
	0.62	0.64	0.63		Негативний
BERT	0.87	0.83	0.85	0.86	Позитивний
	0.85	0.88	0.86		Негативний

Аналіз результатів, представлених у таблиці 2, свідчить про перевагу моделі BERT у порівнянні з іншими моделями. Ця модель досягла найвищих показників точності для класифікації як позитивного, так і негативного емоційного стану. Зокрема, для позитивного емоційного стану точність склала 0.87, повнота – 0.83, а F1-Score – 0.85. Для негативного емоційного стану точність досягла 0.85, повнота – 0.88, а F1-Score – 0.86, що вказує на високу здатність моделі до коректної класифікації обох емоційних категорій при мінімальних помилках. Застосування оптимізатора Adam до моделі CNN призвело до покращення її продуктивності в задачі аналізу настроїв користувача. Точність для позитивного емоційного стану збільшилася до 0.85, що вказує на більш високу ефективність моделі при класифікації позитивних настроїв порівняно з базовою версією. Для негативного емоційного стану точність залишилася на рівні 0.82, що свідчить про стабільну ефективність моделі для цього класу емоцій. Натомість моделі Random Forest, Extra-tree та XGBoost, показали значно нижчі результати. Моделі XGBoost і Random Forest продемонстрували точність лише 0.54 і 0.57 для позитивного емоційного стану, що вказує на їх низьку ефективність у класифікації позитивних емоцій. Хоча ці

моделі досягли кращих результатів для негативного емоційного стану, їх загальна ефективність поступалася нейронним мережам. Модель Extra-tree показала кращі результати порівняно з XGBoost та Random Forest, але не змогла досягти таких високих показників для позитивного емоційного стану, як CNN або BERT. Для негативного емоційного стану точність Extra-tree становила 0.69, а повнота – 0.75.

Матриці помилок, наведені на рисунку 2, відображають результати ефективності кожної з нейронних моделей для аналізу емоційного стану користувачів. Кожна з представлених матриць помилок містить ключові метрики, таких як кількість справжніх позитивних результатів (True Positives, TP), справжніх негативних результатів (True Negatives, TN), помилкових позитивних результатів (False Positives, FP) та помилкових негативних результатів (False Negatives, FN). Ці метрики використано для оцінки здатності моделей до точної класифікації настроїв користувачів.

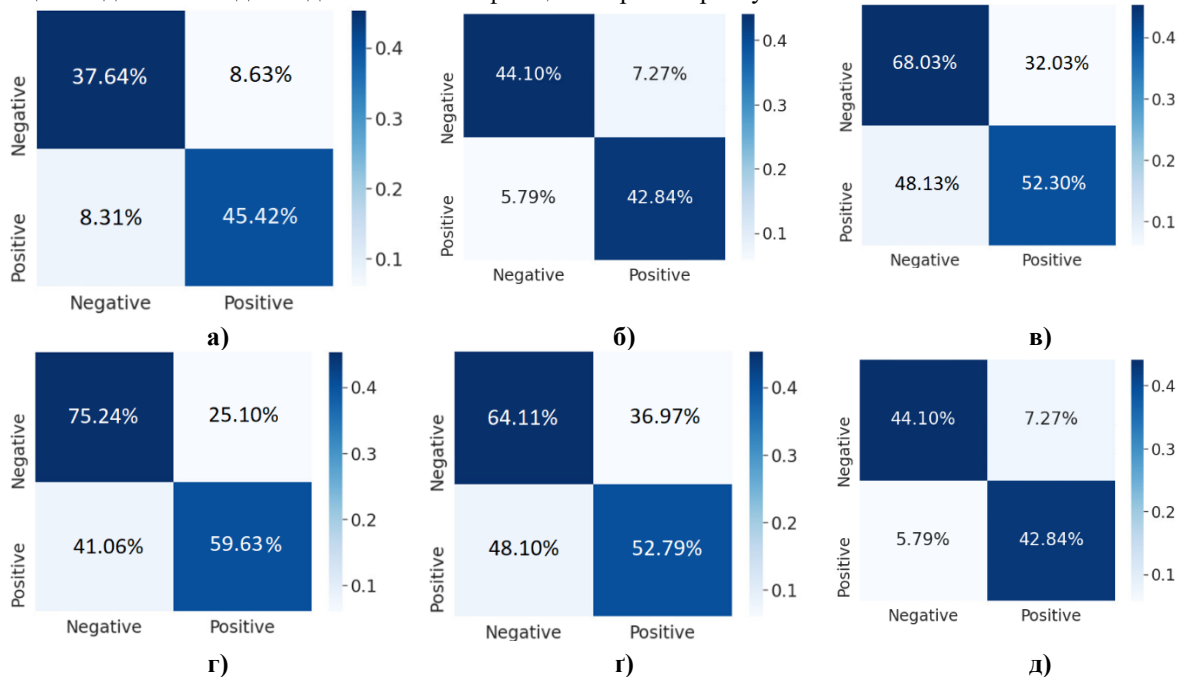


Рис. 2. Матриця помилок нейромережових моделей: а) – CNN; б) – CNN з оптимізатором Adam; в) – Random Forest; г) Extra-tree; р) – XGBoost; д) – BERT

Модель CNN (рис. 2а) продемонструвала високу ефективність у розпізнаванні позитивних і негативних емоцій, зокрема вона краще ідентифікує позитивні емоції. Використання оптимізатора Adam (рис. 2б) значно покращило результати моделі, зокрема знизивши кількість хибнопозитивних та хибнонегативних класифікацій. Це дозволило підвищити здатність моделі точно ідентифікувати позитивні емоції, зробивши її більш збалансованою та знижуючи рівень помилок у класифікації.

Моделі Random Forest (рис. 2в), Extra-tree (рис. 2г) та XGBoost (рис. 2р) продемонстрували нижчі показники ефективності за результатами матриць помилок. Моделі Random Forest та XGBoost характеризуються високим відсотком хибнопозитивних результатів, що свідчить про їхню обмежену здатність точно розрізняти позитивні та негативні емоції в текстах. Модель Extra-tree, хоча й перевершила Random Forest і XGBoost за точністю, але не змогла досягти рівня точності, продемонстрованого моделями BERT та CNN.

Висновки

У результаті проведеного дослідження було оцінено ефективність нейромережових моделей, зокрема CNN, CNN з оптимізатором Adam, Random Forest, Extra-tree, XGBoost та BERT для аналізу емоційного стану користувачів на основі текстових даних із соціальної мережі Twitter, пов'язаних з темою COVID-19.

Результати експерименту показали, що модель BERT досягла загальної точності 86%. Це свідчить про її високу здатність точно класифікувати емоції в текстах, з мінімальною кількістю помилкових класифікацій, що робить її оптимальним інструментом для аналізу настроїв користувачів. Висока точність BERT може бути пояснена її здатністю враховувати контекстні залежності в тексті, що є важливим для точного визначення емоційного стану. Моделі CNN та CNN з оптимізатором Adam продемонстрували загальну точність 83% у обох варіантах, що є нижчим за показники BERT (86%), але вищим за результати Random Forest та XGBoost. Використання оптимізатора Adam дозволило покращити показники для позитивного емоційного стану за метрикою F1-Score. Моделі Random Forest, Extra-tree та XGBoost досягли значно нижчих значень метрик оцінки якості класифікації настроїв користувачів, порівняно з BERT та CNN. Зокрема, точність для позитивного емоційного стану в Random Forest становила лише 57%, а в XGBoost – 54%, що свідчить про обмежену здатність цих

моделей до точного розпізнавання емоцій. Модель Extra-tree досягла загальної точності 68%, але все ж не змогла досягти рівня ефективності CNN та BERT. Таким чином, вибір оптимальної нейронної моделі залежить від специфікацій задачі, що обумовлює її ефективність. Підтверджено, що нейромережеві архітектури є перспективними для аналізу емоційного стану користувачів. Подальші дослідження будуть спрямовані на вивчення впливу різних методів попередньої обробки даних, архітектур нейронних мереж та параметрів навчання на якість аналізу емоційного стану.

Література

1. Tarhini, A., Alalwan, A., Algharabat, R., Masa'deh, R. An analysis of the factors influencing the adoption of online shopping // *International Journal of Technology Diffusion (IJTD)*. – 2018. – Vol. 9, No. 3. – P. 68–87. DOI: 10.4018/IJTD.2018070105.
2. Гнатушенко В., Каштан В., Іванько А., Овчаренко М. Аналіз неструктурованих даних контакт-центру для підтримки прийняття рішень // *Електротехнічні та інформаційні системи*. – 2024. – № 106. – С. 80–86. DOI: 10.32782/EIS/2024-106-14.
3. Salloum, S., Al-Emran, M., Mostafa, M., Azza, S. Using text mining techniques for extracting information from research articles // *Intelligent Natural Language Processing: Trends and Applications*. – Springer, 2018. – P. 373–397. DOI: 10.1007/978-3-319-67056-0_18.
4. Saxena, P., Sharma, S. Systematic Literature Review for Sentiment Analysis Using Big Data social media Streams // *Proceedings of the 2022 5th International Conference on Contemporary Computing and Informatics (IC3I)*. – Uttar Pradesh, India, 2022. – P. 707–714. DOI: 10.1016/j.procs.2019.11.174.
5. Dansana, D., Adhikari, J., Mohapatra, M., Sahoo, S. An approach to analyse and Forecast Social media data using Machine Learning and Data Analysis // *Proceedings of the 2020 International Conference on Computer Science, Engineering and Applications (ICCSEA)*. – Gunupur, India, 2020. – P. 1–5. DOI: 10.1109/ICCSEA49143.2020.9132895.
6. LeCun, Y., Bengio, Y., Hinton, G. Deep learning // *Nature*. – 2015. – Vol. 521, No. 7553. – P. 436. DOI: 10.1038/nature14539.
7. Kharde, V., Sonawane, P. Sentiment analysis of twitter data: a survey of techniques // *arXiv preprint arXiv:1601.06971*. – 2016. – P. 5–15. DOI: 10.5120/ijca2016908625.
8. Araque, O., Corcuera-Platas, S., Ignacio, J. F., Iglesias, C. Enhancing deep learning sentiment analysis with ensemble techniques in social applications // *Expert Systems with Applications*. – 2017. – Vol. 77. – P. 236–246. DOI: 10.1016/j.eswa.2017.02.002.
9. Tul, Q., Ali, M., Riaz, A., Noureen, A., Kamranz, M., Hayat, B., Rehman, A. Sentiment analysis using deep learning techniques: a review // *International Journal of Advanced Computer Science and Applications (IJACSA)*. – 2017. – Vol. 8, No. 6. – P. 424. DOI: 10.14569/IJACSA.2017.080657.
10. Giatsoglou, M., Vozalis, M., Diamantaras, K., Vakali, A., Sarigiannidis, G., Chatzisavvas, K. Sentiment analysis leveraging emotions and word embeddings // *Expert Systems with Applications*. – 2017. – Vol. 69. – P. 214–224. DOI: 10.1016/j.eswa.2016.10.043.
11. Vateekul, P., Koomsubha, T. A study of sentiment analysis using deep learning techniques on Thai Twitter data // *13th International Joint Conference on Computer Science and Software Engineering (JCSSE)*. – 2016. – P. 1–6. DOI: 10.1109/JCSSE.2016.7748849.
12. Coronavirus tweets NLP [Електронний ресурс]. – Режим доступу: <https://www.kaggle.com/datasets/datatattle/covid-19-nlp-text-classification>.
13. Tiwari, D., Nagpal, B., Bhati, B., Kumar, M. A systematic review of social network sentiment analysis with comparative study of ensemble-based techniques // *Artificial Intelligence Review*. – 2023. – Vol. 56. – P. 13407–13461. DOI: 10.1007/s10462-023-10472-w.
14. Pavan, K., Kumar, G. Latent Personality Traits Assessment From Social Network Activity Using Contextual Language Embedding // *IEEE Transactions on Computational Social Systems*. – 2021. – P. 1–12. DOI: 10.1109/TCSS.2021.3108810.

References

1. Tarhini, A., Alalwan, A., Algharabat, R., Masa'deh, R. An analysis of the factors influencing the adoption of online shopping // *International Journal of Technology Diffusion (IJTD)*. – 2018. – Vol. 9, No. 3. – P. 68–87. DOI: 10.4018/IJTD.2018070105.
2. Hnatushenko V., Kashtan V., Ivanko A., Ovcharenko M. Analiz nestrukturevanykh danykh kontakt-tsentru dlia pidtrymky pryiniattia rishen // *Elektrotekhnichni ta informatsiini systemy*. – 2024. – № 106. – С. 80–86. DOI: 10.32782/EIS/2024-106-14.
3. Salloum, S., Al-Emran, M., Mostafa, M., Azza, S. Using text mining techniques for extracting information from research articles // *Intelligent Natural Language Processing: Trends and Applications*. – Springer, 2018. – P. 373–397. DOI: 10.1007/978-3-319-67056-0_18.
4. Saxena, P., Sharma, S. Systematic Literature Review for Sentiment Analysis Using Big Data social media Streams // *Proceedings of the 2022 5th International Conference on Contemporary Computing and Informatics (IC3I)*. – Uttar Pradesh, India, 2022. – P. 707–714. DOI: 10.1016/j.procs.2019.11.174.
5. Dansana, D., Adhikari, J., Mohapatra, M., Sahoo, S. An approach to analyse and Forecast Social media data using Machine Learning and Data Analysis // *Proceedings of the 2020 International Conference on Computer Science, Engineering and Applications (ICCSEA)*. – Gunupur, India, 2020. – P. 1–5. DOI: 10.1109/ICCSEA49143.2020.9132895.
6. LeCun, Y., Bengio, Y., Hinton, G. Deep learning // *Nature*. – 2015. – Vol. 521, No. 7553. – P. 436. DOI: 10.1038/nature14539.

- 7. Kharde, V., Sonawane, P. Sentiment analysis of twitter data: a survey of techniques // arXiv preprint arXiv:1601.06971. – 2016. – P. 5–15. DOI: 10.5120/ijca2016908625.
8. Araque, O., Corcuera-Platas, S., Ignacio, J. F., Iglesias, C. Enhancing deep learning sentiment analysis with ensemble techniques in social applications // Expert Systems with Applications. – 2017. – Vol. 77. – P. 236–246. DOI: 10.1016/j.eswa.2017.02.002.
9. Tul, Q., Ali, M., Riaz, A., Noureen, A., Kamranz, M., Hayat, B., Rehman, A. Sentiment analysis using deep learning techniques: a review // International Journal of Advanced Computer Science and Applications (IJACSA). – 2017. – Vol. 8, No. 6. – P. 424. DOI: 10.14569/IJACSA.2017.080657.
10. Giatsoglou, M., Vozalis, M., Diamantaras, K., Vakali, A., Sarigiannidis, G., Chatzisavvas, K. Sentiment analysis leveraging emotions and word embeddings // Expert Systems with Applications. – 2017. – Vol. 69. – P. 214–224. DOI: 10.1016/j.eswa.2016.10.043.
11. Vateekul, P., Koomsubha, T. A study of sentiment analysis using deep learning techniques on Thai Twitter data // 13th International Joint Conference on Computer Science and Software Engineering (JCSSE). – 2016. – P. 1–6. DOI: 10.1109/JCSSE.2016.7748849.
12. Coronavirus tweets NLP [Elektronnyi resurs]. – Rezhym dostupu: <https://www.kaggle.com/datasets/datatattle/covid-19-nlp-text-classification>.
13. Tiwari, D., Nagpal, B., Bhati, B., Kumar, M. A systematic review of social network sentiment analysis with comparative study of ensemble-based techniques // Artificial Intelligence Review. – 2023. – Vol. 56. – P. 13407–13461. DOI: 10.1007/s10462-023-10472-w.
14. Pavan, K., Kumar, G. Latent Personality Traits Assessment From Social Network Activity Using Contextual Language Embedding // IEEE Transactions on Computational Social Systems. – 2021. – P. 1–12. DOI: 10.1109/TCSS.2021.3108810.