

УДК 004.4'6

DOI: 10.30837/0135-1710.2025.186.005

*К.Е. ПЕТРОВ, І.П. БОКОВ, І.В. КОБЗЕВ***РОЗРОБКА КОМБІНОВАНОГО МЕТОДУ АНАЛІЗУ ЕМОЦІЙНОЇ ЗАБАРВЛЕНOSTІ ТЕКСТІВ**

Розглянуто психологічні й лінгвістичні основи визначення емоційної забарвленості природномовних текстів, різновиди класифікацій та роль емоцій у цифровій комунікації. Проведено порівняльний аналіз існуючих методів аналізу емоційної забарвленості текстів. Запропоновано комбінований метод, який базується на поєднанні двох векторних представлень тексту – статистичного (TF-IDF) та контекстуального (BERT). Таке поєднання дозволяє враховувати як частотні закономірності, так і глибокі семантичні залежності між словами, що в свою чергу, підвищує точність визначення емоційної забарвленості текстів. Наведено результати комп'ютерного моделювання, які демонструють працездатність та ефективність запропонованого методу.

1. Вступ

Визначення емоційної забарвленості тексту є однією з ключових задач обробки природної мови (Natural language processing, NLP) [1] та відіграє важливу роль у численних прикладних сферах, зокрема в маркетингу, соціології, психології, аналізі громадської думки та інформаційній безпеці. Суттєве зростання обсягу текстових даних у соціальних мережах, новинних ресурсах, блогах, коментарях та на форумах значно посилює потребу в автоматичному визначенні забарвленості текстів [1]. Системи аналізу емоційної забарвленості текстів дозволяють оперативно отримувати структуровану інформацію про настрої суспільства, прогнозувати реакцію на певні події, виявляти потенційні загрози або деструктивний контент, а також здійснювати моніторинг громадської думки в реальному масштабі часу.

Попри значні успіхи в цій сфері, існуючі методи аналізу забарвленості мають ряд обмежень, які знижують їхню ефективність. Зокрема, традиційні методи, що базуються на використанні словників чи статистичного аналізу, часто не враховують контекстуального значення слів, що є критично важливим для точного розпізнавання емоційної забарвленості текстів [2]. Крім того, багато методів стикаються з труднощами при аналізі багатозначних слів, сарказму, іронії, неформальних та сленгових виразів, які досить часто використовуються в сучасній комунікації. Ці нюанси роблять аналіз текстів складнішим завданням та потребують глибшого розуміння контексту й лінгвістичних особливостей мови. Крім того, тексти, що містять змішані емоції чи неоднозначні почуття, також є важкими для їх коректної класифікації за допомогою стандартних методів, які часто орієнтовані на чітко виражені емоції [1]. Для подолання цих проблем необхідно розвивати та застосовувати гібридні методи, які б ефективно поєднували переваги існуючих рішень.

Таким чином, існує потреба у подальшому вдосконаленні методів аналізу забарвленості тексту шляхом інтеграції різних підходів, зокрема з використанням сучасних моделей глибокого навчання, таких як нейронні мережі з механізмом уваги (attention mechanism), трансформерні моделі (наприклад, BERT, RoBERTa, GPT), які дозволяють краще враховувати контекстуальні та семантичні особливості текстів [3], [4]. Використання трансформерних моделей забезпечує глибше розуміння контексту через можливість моделювати семантичні взаємозв'язки між словами на різних відстанях у тексті. Водночас, поєднання нейронних мереж із механізмами уваги дозволяє акцентувати увагу моделі на найзначущіших елементах тексту, підвищуючи точність і чутливість аналізу [5]. Тому актуальною є розробка комбінованого методу визначення емоційної

забарвленості природномовних текстів [6], який би забезпечив підвищення точності визначення емоційної забарвленості за рахунок ефективного поєднання різних методів і передових моделей машинного навчання.

2. Аналіз існуючих методів аналізу емоційної забарвленості текстів та визначення проблеми дослідження

Виявлення емоцій у тексті є складним завданням, оскільки воно тісно пов'язане з особливостями контексту, багатозначністю мови, наявністю стилістичних засобів, таких як сарказм, іронія чи метафори. Через це для класифікації емоцій у текстах використовуються різноманітні методи.

Одним з найпоширеніших – є метод класифікації за полярністю емоцій, який передбачає поділ текстів на три основні категорії: позитивні, негативні та нейтральні.

Інший популярний метод базується на теорії базових емоцій Пола Екмана [7], в якій виділяються шість універсальних базових емоцій: радість (joy), смуток (sadness), гнів (anger), страх (fear), здивування (surprise) та огида (disgust),.

Подальшим розвитком концепції [6] стала модель «Колесо емоцій» («Wheel of Emotions») Роберта Плутчика [8]. Вона охоплює вісім основних емоцій: радість, довіру (trust), страх, здивування, сум, відразу, гнів та очікування (expectation). Ці емоції можуть комбінуватись між собою, формуючи складніші та тонші емоційні стани, наприклад, любов (love, поєднання радості та довіри).

Окремо можна виділити метод континуальної класифікації емоцій, в якому використовуються багатовимірні простори для опису емоційних станів. Однією з найпоширеніших є тривимірна модель, що описує емоції через три параметри [9]: валентність (ступінь позитивності чи негативності емоції), активацію (від стану спокою до інтенсивного збудження) і домінантність (ступінь контролю над ситуацією). Наприклад, щастя характеризується високою валентністю, високою активацією та значною домінантністю.

Виявлення емоцій у текстах базується на двох ключових аспектах: лінгвістичному та психологічному. Лінгвістичний аспект [10] передбачає аналіз мовних засобів, які використовуються авторами для вираження емоцій, зокрема лексики, морфології, синтаксису та стилістичних особливостей тексту. Психологічний аспект акцентує увагу на емоціях як когнітивно-емоційних станах [11], що відображають внутрішній досвід автора, проявляючись через специфічні експресивні засоби у тексті. Це передбачає моделювання емоційних станів [7], [8], [9] на основі різних психологічних теорій.

Сучасні методи аналізу емоцій у текстах намагаються інтегрувати лінгвістичні та психологічні аспекти, створюючи комплексні моделі, які здебільшого використовують машинне навчання та нейронні мережі, що дозволяє враховувати як лексичні особливості слів, так і контекстуальні та психологічні аспекти текстів. Такий інтегрований метод суттєво покращує точність автоматичного аналізу емоцій.

У галузі автоматичного аналізу емоцій у текстах використовуються різноманітні методи, які умовно можна поділити на лексиконні, статистичні, методи машинного навчання та гібридні. Кожен з них має свої переваги та обмеження.

Лексиконні методи до аналізу емоційної забарвленості тексту [1], [12] посідають ключове місце серед традиційних методів обробки природної мови. Вони базуються на гіпотезі, що емоційний стан автора можна визначити через використані ним мовні одиниці – переважно слова, що мають певне емоційне навантаження [13]. У центрі цих методів знаходяться так звані лексикони – попередньо сформовані спеціальні словники (наприклад, SentiWordNet, AFINN, NRC Emotion Lexicon), що містять слова, кожне з яких має заздалегідь визначену емоційну категорію або числову оцінку полярності (від сильно негативної до сильно позитивної) [1, 12]. Залежно від типу лексикону, емоційні характеристики можуть

бути представлені у вигляді бінарної шкали (позитивне або негативне), багаторівневої числової шкали (наприклад, від -5 до +5), або через категоріальну приналежність до базових емоцій за теоріями Плутчика [8], Екмана [7] та інших.

Серед лексиконних методів можна виділити:

- методи, що використовують сентиментні лексикони – спеціально сформовані переліки слів із визначеною емоційною характеристикою (позитивна, негативна чи нейтральна). Наприклад, до таких ресурсів належать WordNet-Affect, SentiWordNet, LIWC (Linguistic Inquiry and Word Count);

- методи, що базуються на використанні розширених лексиконів – баз даних, які окрім емоційних слів, включають також їхні синоніми, антоніми та оцінку інтенсивності емоцій. Одним із прикладів є VADER (Valence Aware Dictionary and sEntiment Reasoner), який добре підходить для аналізу текстів із соціальних мереж;

- методи на основі правил (Rule-based methods), які ґрунтуються на використанні наборів правил, що дозволяють враховувати контекст для посилення чи зменшення емоційного забарвлення фраз (наприклад, різниця між «дуже щасливий» і «трохи щасливий»).

Лексиконні методи аналізу емоційної забарвленості текстів відіграють важливу роль у задачах класифікації настроїв завдяки своїй простоті, обчислювальній ефективності та високій інтерпретованості результатів. Використання попередньо сформованих лексиконів дозволяє швидко оцінити емоційний зміст тексту без потреби в навчанні моделей на великих корпусах. Такі методи є особливо ефективними у випадках обробки коротких текстів, а також у системах, де необхідна прозорість прийнятих рішень.

Водночас обмеження лексиконних методів, зокрема нечутливість до контексту, труднощі з інтерпретацією заперечень, іронії та багатозначних слів, знижують їх точність у складних мовних конструкціях. Через це сучасні дослідницькі підходи дедалі частіше комбінують лексичний аналіз із контекстно-залежними моделями, використовуючи лексикони як модулі попередньої обробки або джерела інтерпретованих ознак. Таким чином, лексиконні методи зберігають свою актуальність у сучасному ландшафті обробки природної мови і використовуються як основа для гібридних систем емоційної класифікації.

На відміну від лексиконних методів, які ґрунтуються на заздалегідь сформованих словниках з емоційними маркерами, статистичні методи класифікації текстів використовують формальні математичні моделі для аналізу текстових даних, визначення статистичних закономірностей і подальшого віднесення тексту до певної емоційної категорії [14]. Ці методи ґрунтуються на ідеї, що текст можна представити у вигляді числових ознак, які можна обробити так само, як будь-які інші дані у задачах класифікації. Завдяки цьому емоційна забарвленість повідомлень виявляється шляхом виявлення кореляцій між статистичними характеристиками та цільовими етикетками (позитивний, негативний, нейтральний тощо).

Статистичний метод є особливо цінним у контексті автоматизації обробки великих обсягів інформації, де ручне анотоване кодування або використання фіксованих словників є малоефективним або взагалі неможливим [3]. Замість цього моделі навчаються розпізнавати шаблони у тексті на основі аналізу розподілу слів, частоти їх вживання та інших числових метрик. Це дозволяє виявити латентні (приховані) закономірності, які могли бути непомітними для людини або непридатними для лексичного аналізу. Наприклад, статистичні моделі здатні враховувати характерні поєднання слів або тенденції до використання певної лексики у різних емоційних контекстах, що особливо важливо в умовах неоднозначності та варіативності мови.

Статистичні методи ґрунтуються на аналізі частоти вживання слів, виявленні співвідношення емоційно забарвленої лексики та дослідженні словосполучень (колокацій),

які часто зустрічаються разом у тексті [14].

Найпопулярніші статистичні методи [2], [3] – Term Frequency (TF), Term Frequency–Inverse Document Frequency (TF-IDF), аналіз N-грам, Latent Semantic Analysis (LSA) – дозволяють ефективно обробляти великі корпуси текстів.

Кожен з цих методів має свої переваги й обмеження використання. TF-IDF дозволяє виокремити ключові слова, але не враховує порядок і контекст. Аналіз N-грам краще відображає локальний контекст, проте зі збільшенням N зростає обчислювальна складність. LSA виявляє приховані семантичні зв'язки, однак потребує великого обсягу даних.

Застосування статистичних методів є ефективним для аналізу великих корпусів тексту, але вони також не враховують глибокий контекст слів та їхній емоційний зміст.

Для подолання вказаних вище проблем аналізу емоційної забарвленості текстів активно використовуються методи машинного навчання та глибокі нейронні мережі.

Основні методи включають використання: класичних методів машинного навчання, векторного представлення слів (word embeddings), а також нейронних мереж, зокрема трансформерних моделей.

Методи класичного машинного навчання широко використовуються при розробці моделей, що здатні автоматично навчатися на основі розмічених текстових корпусів і прогнозувати емоційні мітки для нових текстів. Завдяки своїй обчислювальній ефективності, відносній простоті реалізації та високій швидкодії, ці методи широко застосовуються в практиці аналізу текстових даних. Найпоширенішими з них є моделі наївного байєсівського класифікатора (Naive Bayes), логістичної регресії (Logistic Regression), метод опорних векторів (SVM), ансамблевий метод випадкових лісів (Random Forest) та метод k найближчих сусідів (k-NN) [13], [15]. Кожен з цих методів має унікальні характеристики, переваги та сфери оптимального застосування, що дозволяє гнучко адаптувати їх до різних сценаріїв класифікації емоцій.

Останнім часом для інтелектуального аналізу текстів все частіше застосовують глибоке навчання, що базується на використанні, зокрема, рекурентних нейронних мереж (RNN), мереж довгої короткострокової пам'яті (LSTM), згорткових мереж (CNN) та трансформерів (BERT, GPT, RoBERTa тощо) [16]-[19]. Вони враховують контекст і семантичні зв'язки між словами, а також можуть інтерпретувати складні стилістичні прийоми, такі як іронія та сарказм.

Переваги використання глибокого навчання полягають у високій точності навіть у складних мовних випадках, здатності обробляти великі корпуси неструктурованих текстових даних та можливості донавчання на вузькоспеціалізованих доменах [12].

Здатність архітектур глибокого навчання адаптуватися до особливостей природної мови, враховувати контекст, розпізнавати приховані патерни та забезпечувати конкурентну точність дозволяє створювати системи, які перевершують можливості систем, що базуються на використанні класичних методів за результатами.

Разом з тим існують і недоліки, притаманні моделям глибокого навчання. Зокрема, їхня висока обчислювальна складність потребує наявності потужного апаратного забезпечення, що обмежує використання в реальному часі на малопотужних пристроях [12], [17]. Крім того, складність їхньої архітектури утруднює інтерпретацію результатів, що є критичним чинником у сферах, де потрібна прозорість прийняття рішень. Ще однією проблемою є потреба у великій кількості якісно анотованих текстових даних, без яких повноцінне навчання моделі стає неможливим.

Для підвищення точності аналізу в сучасних системах визначення емоційної забарвленості текстів доцільно використовувати комбіновані методи, що поєднують лексиконні методи, статистичний аналіз та глибокі нейронні мережі. Наприклад, спочатку може виконуватися

попередній аналіз тексту за допомогою лексиконних методів, а потім отримані дані подаються на вхід глибокої нейронної мережі для уточнення емоційної забарвленості.

Такі комбіновані методи дозволять враховувати як поверхневі лексичні ознаки, так і глибокі семантичні зв'язки у тексті. Застосування декількох методів у межах однієї системи сприятиме підвищенню загальної точності класифікації текстів. Тому вирішення проблеми адекватної класифікації природномовних текстів за їхньою емоційною забарвленістю на основі комбінації кількох методів є актуальним з теоретичної та прикладної точок зору.

3. Мета і задачі дослідження

Метою дослідження є підвищення точності класифікації емоційної забарвленості природномовних текстів за рахунок використання лексиконних, статистичних і контекстуальних методів, які дозволять врахувати як поверхневі лексичні ознаки, так і глибокі семантичні зв'язки у тексті.

Для досягнення поставленої мети необхідно вирішити такі основні задачі:

- розробити комбінований метод для вирішення задачі визначення емоційної забарвленості тексту;
- провести експериментальну перевірку працездатності та ефективності запропонованого методу.

4. Розробка комбінованого методу визначення емоційної забарвленості текстів

Комбінований метод визначення емоційної забарвленості текстів, що пропонується, ґрунтується на поєднанні лексиконних, статистичних і контекстуальних методів для досягнення максимальної точності при класифікації емоцій.

Основна ідея методу полягає в тому, щоб інтегрувати переваги класичних методів машинного навчання, які працюють з лексичними ознаками, із сучасними можливостями глибоких нейронних мереж, зокрема трансформерів, що здатні враховувати контекст і семантичні зв'язки між словами [4]. Архітектура методу передбачає чітко визначену послідовність етапів (рис. 1), кожен з яких відіграє ключову роль у забезпеченні високої якості результатів.



Рис. 1. Основні етапи комбінованого методу класифікації емоцій

Розглянемо основні етапи запропонованого методу детальніше.

На першому етапі здійснюється попередня обробка тексту. На цьому етапі Вхідний текст проходить стандартні процедури очищення: видалення зайвих символів, пунктуації, HTML-тегів, зниження регістру, а також лематизацію – приведення слів до їхньої базової форми. Здійснюється також видалення стоп-слів, які не несуть значущого семантичного навантаження. Цей крок сприяє зменшенню шуму в даних та покращенню якості подальших перетворень.

Другий етап – векторизація тексту, яка реалізується паралельно за двома напрямками. Перший напрямок – лексичне векторне представлення тексту, зокрема через TF-IDF.

Цей метод дозволяє відобразити статистичну значущість кожного терміну у тексті з урахуванням його частоти у документі та у всьому корпусі. У запропонованому методі TF-IDF використовується для побудови векторів фіксованої довжини (5000 найрелевантніших ознак), які передають частотну і тематичну інформацію про текст.

Другий напрямок – контекстуальне представлення тексту за допомогою трансформерної моделі BERT (Bidirectional Encoder Representations from Transformers). Замість використання BERT у режимі fine-tuning, обрано режим feature extraction: текст пропускається через попередньо натреновану модель bert-base-uncased і потім з нього видобувається векторне представлення, отримане шляхом усереднення CLS-токенів з останнього шару моделі. Це дозволяє отримати глибоко семантизоване уявлення про зміст тексту, що враховує контекстне вживання слів у реченні.

Третій етап передбачає інтеграцію ознак з обох представлень: TF-IDF та BERT-векторів. Отримані два векторні простори з'єднуються в єдиний гібридний вектор шляхом конкатенації. Це дозволяє комбінувати статистичні закономірності з семантичними зв'язками, що значно розширює інформаційне поле для класифікації текстів та знижує ймовірність втрати важливих ознак.

Четвертим етапом є класифікація текстів за їхньою емоційною забарвленістю з використанням моделі машинного навчання. У запропонованому методі використовується Random Forest Classifier – ансамблевий метод, що базується на ухваленні рішення за допомогою колективу дерев. Цей метод було обрано через його стійкість до перенавчання, здатність працювати з великими масивами ознак і високу інтерпретованість результатів. Для оптимізації продуктивності та зменшення ризику перенавчання, перед навчанням класифікатора здійснювалася стандартизація ознак [2]. Усі вектори приводилися до єдиного діапазону значень з використанням Z-score normalization. Далі здійснюється передбачення емоційної забарвленості текстів, що не входили до навчального набору. Вихідним результатом є прогноз емоційної категорії (позитивна, негативна або нейтральна) разом із ймовірнісною оцінкою, яка дає змогу враховувати рівень упевненості моделі у власному рішенні.

У структурному плані запропонований метод реалізовано як послідовність етапів, завдяки чому його можна легко адаптувати до вирішення нових завдань або замінити окремі компоненти, що використовуються на кожному окремому етапі, на досконаліші – наприклад, використати замість TF-IDF сучасніші методи векторизації, як-от BM25, або змінити базову трансформерну модель BERT на RoBERTa чи DistilBERT [3], [4]. Це забезпечує аналітичне підґрунтя для подальшої оптимізації або модифікації окремих складових методу, що є особливо важливим при масштабуванні системи або адаптації до нових наборів даних та доменно-специфічного контенту.

Для реалізації методу було обрано мову програмування Python, яка є найпоширенішою у сферах обробки природної мови і машинного навчання [17] та має розвинену екосистему бібліотек для обробки текстів, векторизації, роботи з моделями глибокого навчання та візуалізації результатів [20]. Зокрема, при розробці методу було використано такі бібліотеки [20]: Pandas і NumPy – для забезпечення попередньої обробки й агрегації даних; Scikit-learn – для реалізації ключових етапів векторизації, навчання моделі, масштабування ознак і побудови метрик оцінки ефективності; Transformers у поєднанні з PyTorch – для отримання та інтеграції трансформерного представлення тексту, зокрема BERT, що значно підвищило якість ознак за рахунок урахування контексту; Matplotlib та Seaborn – для візуалізації результатів та аналізу точності класифікації емоцій.

5. Експериментальна перевірка ефективності використання комбінованого методу

Для експериментальної перевірки працездатності та ефективності запропонованого методу було використано набір даних з емоційної класифікації англomовних твітів, що доступний на платформі Kaggle – Emotion Dataset for NLP [21]. Цей датасет містить короткі повідомлення з мітками емоцій, які відповідають певним емоційним станам автора. Набір даних включає понад 20000 рядків тексту, які містять емоційні мітки з таких категорій, як joy, anger, sadness, fear, love, surprise тощо. Для спрощення задачі та приведення її до трикласової моделі (позитивна, негативна, нейтральна), було здійснено агрегування класів: позитивні емоції – joy, love, surprise; негативні – anger, sadness, fear; нейтральні – невиражені або конфліктні випадки.

Балансування класів було здійснено шляхом підсумовування підвибірок, щоб уникнути домінування одного емоційного класу над іншими. Перед векторизацією всі тексти пройшли однакову обробку: зниження регістру, видалення пунктуації та спеціальних символів, токенизацію та видалення стоп-слів. Ці процедури дозволили уніфікувати структуру тексту, покращити якість векторизації та зменшити розмірність вхідного простору.

Як зазначено вище, комбінований метод реалізує концепцію двонаправленого опрацювання текстових даних, що передбачає паралельне використання двох незалежних, але комплементарних методів векторного представлення:

- TF-IDF – для обчислення було обрано 5000 найпоширеніших термінів. Використано TfidfVectorizer зі стандартними параметрами, включно з ngram_range=(1,2) та max_df=0.9, щоб відсіяти занадто часті слова. Результатом є розріджена матриця розміром $n \times 5000$;

- BERT (base-uncased): використано модель bert-base-uncased із Hugging Face, без fine-tuning. Отримання векторів здійснювалося через обчислення CLS-токену або середнього по всіх токенах на останньому прихованому шарі. Розмірність отриманого вектору – 768 ознак.

Обидва представлення об'єднувались у єдиний вектор ознак розмірності 5768, після чого здійснювалась стандартизація за допомогою StandardScaler.

Як класифікатор обрано Random Forest, оскільки він добре справляється з великою кількістю ознак, не потребує масштабування і стійкий до перенавчання. Було використано такі параметри моделі: n_estimators = 300; max_depth = 20; min_samples_leaf = 2; random_state = 42; n_jobs = -1 (для використання всіх CPU-ядер).

Модель навчалась на 80 % даних, інші 20 % використовувалися для тестування. Було також реалізовано п'ятикратну крос-валідацію для перевірки стабільності результатів.

Для оцінки ефективності моделі було використано стандартні метрики: accuracy – загальна точність класифікації (частка документів, за якими класифікатор прийняв правильне рішення, відносно до загальної кількості класифікованих документів); precision – точність в межах класу (частка документів, що істинно належать даному класу, відносно до загальної кількості документів які система віднесла до цього класу); recall – повнота (частка знайдених класифікатором документів, що належать класу, відносно до загальної кількості документів цього класу в тестовій вибірці); F1-score – гармонійне середнє між precision і recall [22]. Крім того, для візуалізації правильних і неправильних передбачень щодо віднесення тексту до відповідного класу в дослідженні використано Confusion Matrix (матриця похибок класифікації).

На рис. 2 представлено матрицю похибок, яка демонструє достатньо високу точність класифікації, з невеликим рівнем помилок між класами, що характерно для комбінованих моделей із сильними векторними поданнями представленнями.

Для оцінювання ефективності використання комбінованого методу було здійснено його порівняння з іншими поширеними методами, а саме TF-IDF + Random Forest (RF) та BERT + Random Forest (RF). Метою такого аналізу було визначення, наскільки запропонований

комбінований метод випереджає базові методи як за точністю, так і за стабільністю результатів.

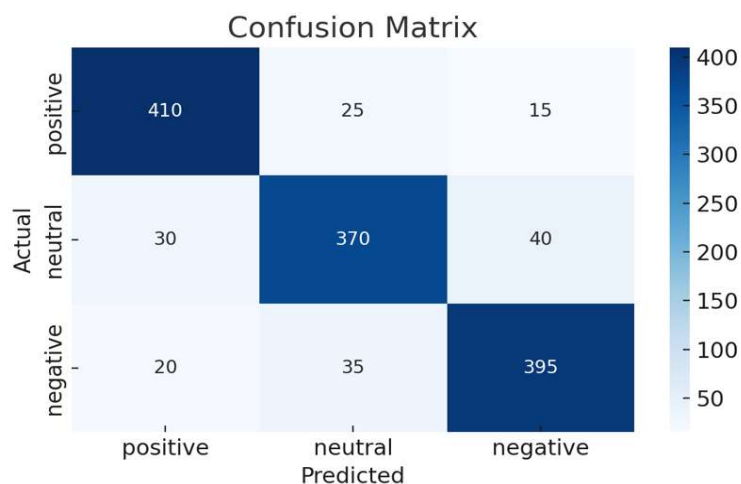


Рис. 2. Матриця похибок класифікації емоційної забарвленості текстів

На рис. 3 представлено діаграму значень точності різних моделей класифікації.

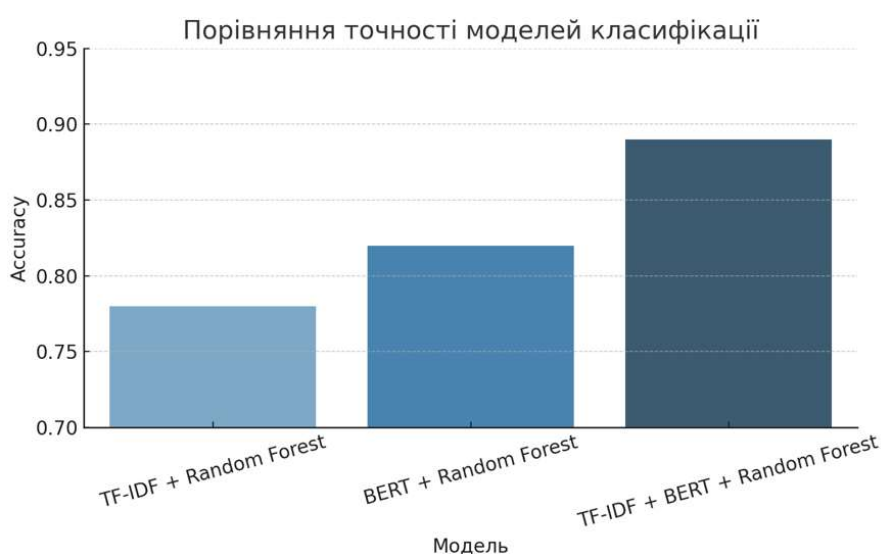


Рис. 3. Діаграма значень точності моделей класифікації

Значення F1-score представлено на рис. 4, який наочно ілюструє різницю в збалансованості точності та повноти між моделями, що базуються на різних комбінаціях ознак.

Оцінки часових витрат на тренування різних моделей, представлено на рис. 5.

З рис. 5 видно, що комбінований метод, який базується на використанні моделі TF-IDF + BERT + RF тренується найдовше, оскільки включає обидва типи векторизації та обробку великої кількості ознак.

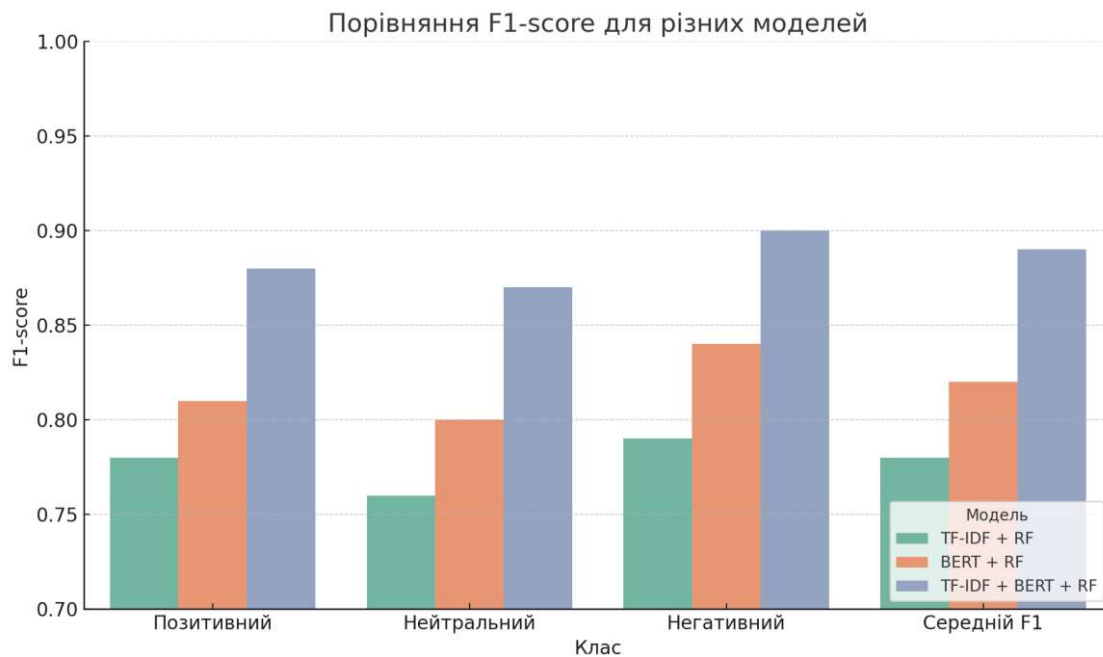


Рис. 4. Діаграма значень F1-score для різних моделей

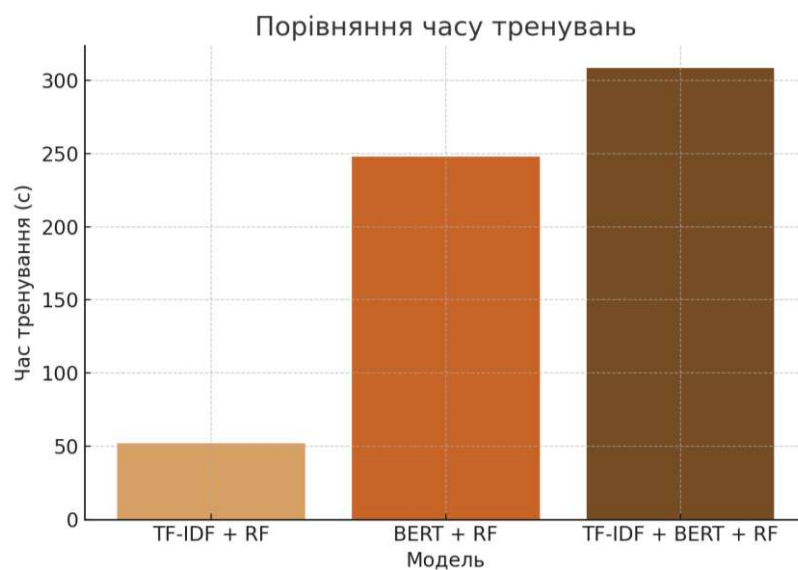


Рис. 5. Діаграма оцінок витрат часу тренування різних моделей

6. Обговорення результатів дослідження

Результати експериментальних досліджень підтвердили ефективність запропонованого комбінованого методу до визначення емоційної забарвленості текстів, що поєднує ознаки TF-IDF і контекстні вектори BERT.

З рис. 3 видно, що більшість прикладів правильно класифіковані, зокрема позитивні й негативні емоції мають найвищий рівень точності. Найбільше похибок спостерігається між класами «нейтральний» і «позитивний», що є типовим для коротких текстів, у яких емоційна забарвленість виражена неявно. Незначна кількість хибних спрацьовувань для негативного класу вказує на високу чутливість моделі до емоцій негативного спектра.

Загалом отримана матриця підтверджує ефективність комбінованого методу:

поєднання статистичних і контекстуальних ознак дозволило краще відрізнити тонко виражені емоційні стани, порівняно з класичними методами.

Порівняння з базовими методами показав, що саме запропонований метод демонструє найвищі значення асигасу та F1-score, перевершуючи як моделі на основі лише лексичних ознаках, так і моделі, що використовують лише BERT.

Комбінована модель TF-IDF + BERT + Random Forest забезпечує найвищий рівень точності (рис. 3), яка досягає значення близько 89 %. Це підтверджує ефективність інтеграції двох типів ознак – статистичних (TF-IDF) і контекстуальних (BERT) – у рамках одного векторного простору. Модель BERT + Random Forest показала середню точність – приблизно 82 %, що вказує на добру здатність BERT моделювати контекст, однак без доповнення частотними ознаками вона поступається гібридному методу. Найнижчу точність продемонструвала модель TF-IDF + Random Forest – близько 78 %, що свідчить про її обмежену здатність враховувати глибокі семантичні зв'язки в тексті.

Отримані значення F1-score (рис. 4) свідчать, що гібридна модель TF-IDF + BERT + RF демонструє найкращі результати для всіх трьох класів – позитивного, нейтрального та негативного. Середній F1 для цієї моделі становить близько 88 %, що помітно перевищує показники інших моделей. Модель BERT + RF посіла друге місце за ефективністю, особливо добре впоравшись із класифікацією негативних повідомлень. Її результати свідчать про високу цінність контекстуального представлення, хоча в сукупності ця модель поступається комбінованій моделі. Найгірші результати показала модель TF-IDF + RF, яка не змогла досягти високої точності в жодному з класів. Це ще раз підкреслює обмеженість використання виключно частотного методу в задачах, де важливе розуміння семантики.

Водночас, комбінована модель потребує більше часу на тренування, що було продемонстровано на відповідній діаграмі (рис. 5). Однак ці витрати компенсуються суттєвим зростанням якості результатів.

На основі аналізу результатів, отриманих в ході проведення експериментальних досліджень можна сформулювати низку рекомендацій щодо підвищення точності аналізу емоційної забарвленості природномовних текстів. По-перше, доцільно комбінувати різні типи ознак – частотні, контекстуальні та синтаксичні – для отримання повнішого представлення тексту. По-друге, варто проводити розширену попередню обробку, включно з врахуванням емодзі, синонімів, синонімічних перетворень і фразеологізмів, що є носіями прихованих емоцій. Крім того, покращення може бути досягнуто за рахунок тонкого налаштування трансформерної моделі під конкретний домен або використання ансамблевих методів, що поєднують прогнози декількох моделей. У сукупності ці заходи здатні забезпечити вищу чутливість моделі до нюансів емоційної виразності, особливо в коротких повідомленнях.

Подальші дослідження можуть бути зосереджені на оптимізації швидкодії моделі, зокрема шляхом використання полегшених варіантів BERT, таких як DistilBERT або TinyBERT, а також застосування методів квантування або дистиляції знань. Перспективним напрямом є також дослідження впливу різних схем агрегування ознак, включно з механізмами уваги (attention mechanism) [23] або autoencoder-об'єднанням векторів. У майбутньому доцільним є адаптація методу для багатомовного середовища та доменно-специфічного аналізу, що дозволить ширше застосовувати його для вирішення реальних завдань.

Крім того, доцільним є дослідження інтеграції зовнішніх знань, наприклад, у вигляді семантичних мереж або емоційних онтологій, які можуть підсилити якість інтерпретації тексту. Застосування багаторівневого ансамблю моделей може забезпечити гнучкість та адаптивність системи до різних джерел текстових даних. Перспективним є також включення модулів самооцінювання моделі, які дозволять оцінювати рівень впевненості у класифікаційних рішеннях. Це особливо важливо для критичних застосувань, де недостовірні

емоційна інтерпретація може призвести до хибних висновків. У напрямі адаптації метода до використання в системах реального часу доцільно розглянути гібридні архітектури, що поєднують глибоке навчання з класичними методами для досягнення компромісу між точністю та швидкістю. Важливо також досліджувати вплив різних стратегій попередньої обробки, зокрема нормалізації тексту, фільтрації шуму та виявлення іронії, які суттєво впливають на результати класифікації. Нарешті, інтеграція з інтерфейсами користувача (наприклад, візуалізація емоційної динаміки) може розширити функціональність систем та зробити їх зручнішими й інформативнішими для кінцевих користувачів.

7. Висновки

В рамках проведеного дослідження було розроблено комбінований метод визначення емоційної забарвленості текстів, що поєднує його статистичні й семантичні представлення для досягнення підвищеної точності класифікації емоцій.

Запропонований у роботі комбінований метод є прикладом гібридного методу до аналізу емоційної забарвленості тексту, в якому поєднуються два паралельні векторні представлення тексту: статистичне (TF-IDF) та контекстуальне (BERT). Таке поєднання дозволило враховувати як частотні закономірності, так і глибокі семантичні залежності між словами. У поєднанні з ансамблевим класифікатором Random Forest вдалося побудувати стійку модель, яка здатна ефективно класифікувати короткі англійські тексти з високим рівнем точності.

Для оцінки ефективності запропонованого методу було проведено його порівняння з базовими методами, що базуються на використанні моделей TF-IDF + RF та BERT + RF. Експерименти показали, що комбінована модель досягає точності класифікації до 89%, що перевищує результати альтернативних моделей (78 % і 82 % відповідно). Середній F1-score також виявився найвищим у комбінованій моделі – близько 88 %. Такий результат підтверджує доцільність інтеграції різних типів ознак. Було також розглянуто часові витрати на тренування моделей. Комбінована модель очікувано виявилася найресурсовитратнішою, що пов'язано з необхідністю одночасного обчислення двох типів векторів. Однак ці витрати повністю компенсуються зростанням точності та стабільності результатів.

Подальші дослідження можуть бути зосереджені на оптимізації архітектури моделі для скорочення часу її тренування, зокрема шляхом використання квантованих або дистильованих трансформерів. Крім того, гібридна архітектура може бути розширена на багатокласову класифікацію емоцій, що дозволить досягти детальнішого аналізу психологічного стану автора повідомлення. Не менш важливою є розробка explainable AI-механізмів для пояснення прийнятих рішень, що особливо актуально у таких чутливих сферах, як медицина, освіта та юриспруденція. Запропонований метод доцільно інтегрувати у веб-сервіси, CRM-платформи, чат-боти і аналітичні системи соціальних мереж та медіа, що забезпечить його практичну реалізацію.

Перелік посилань:

1. Goldberg Y. Neural Network Methods in Natural Language Processing. San Rafael: Morgan & Claypool, 2017. 287 p. URL: <https://doi.org/10.1007/978-3-031-02165-7>
2. Joshi D. Emotion Detection using Transformer Model with Deep Learning. *International Journal of Engineering Research and Technology*. 2025. Vol. 14 (3). P. 1–7 p. URL: <https://doi.org/10.17577/IJERTV14IS030116>
3. Acheampong F., Wenyu C., Nunoo-Mensah H. Text-based emotion detection: Advances, challenges, and opportunities. *Engineering Reports*. 2020. Vol. 2 (7). P. 1–24. URL: <https://doi.org/10.1002/eng2.12189>
4. Bing L. Sentiment Analysis: Mining Opinions, Sentiments, and Emotions. Cambridge : Cambridge University Press, 2020, 448 p.
5. Vanshika, Rani N., Walia R. A Comprehensive Review of Sentiment Analysis: Techniques, Datasets, Limitations, and Future Scope. *2024 Sixth International Conference on Computational Intelligence and*

Communication Technologies (CCICT), Sonapat, India, April 19-20, 2024. P. 403-409, URL: <https://doi.org/10.1109/CCICT62777.2024.00072>.

6. Боков І. П., Петров К. Е. Дослідження методів аналізу емоційного забарвлення текстів. *Радіоелектроніка та молодь у XXI столітті: матеріали XXVIII Міжнар. молодіж. форуму*, 16–18 квіт. 2025 р. Харків, 2025, Т. 6. С. 10–11.

7. Ekman P. Telling lies: Clues to deceit in the marketplace, politics, and marriage. W. W. Norton & Company, 2009. 400 p.

8. Emotion: Theory, Research and Experience. Volume 1. Theories of Emotion. Edited by R. Plutchik and H. Kellerman. Academic Press: London, 1980. 399 p. URL: <https://doi.org/10.1017/S0033291700053769>

9. Russell S., Norvig P. Artificial Intelligence: A Modern Approach. 4th ed. Boston : Pearson, 2020. 1136 p.

10. Jurafsky D., Martin J. Speech and Language Processing. 3rd ed. New Jersey : Pearson, 2008. 1024 p.

11. Sailunaz K., Alhajj R. Emotion and Sentiment Analysis from Twitter Text. *Journal of Computational Science*. 2019. Vol. 36. P. 437–448. URL: <https://doi.org/10.1016/j.jocs.2019.05.009>.

12. Dang N., Moreno-García M., De la Prieta F. Sentiment analysis based on deep learning: a comparative study. *Electronics*. 2020. Vol. 9(3). P. 1–29. URL: <https://doi.org/10.3390/electronics9030483>

13. Ghafoor Y., Jinping S., Calderon F., Huang Y., Chen K., Chen Y. TERMS: textual emotion recognition in multidimensional space. *Applied Intelligence*. 2023. Vol. 53(3). P. 2673–2693. URL: <https://doi.org/10.1007/s10489-022-03567-4>

14. Manning C., Schütze H. Foundations of Statistical Natural Language Processing. Cambridge: MIT Press, 1999. 718 p.

15. Minaee S., Kalchbrenner N., Cambria E., Nikzad N., Chenaghlu M., Gao J. Deep learning based text classification: A comprehensive review. *ACM Computing Surveys*. 2022. Vol. 54(3). P. 1–40. URL: <https://doi.org/10.1145/3439726>

16. Otter D., Medina J., Kalita J. A Survey of the Usages of Deep Learning for Natural Language Processing. *IEEE Transactions on Neural Networks and Learning Systems*. 2020. Vol. 32(2). P. 604–624. URL: <https://doi.org/10.1109/TNNLS.2020.2979670>

17. Chen S., Zhang Y., Yang Q. Multi-Task Learning in Natural Language Processing: An Overview. *ACM Computing Surveys*. 2024. Vol. 56(12). P. 1–32. URL: <https://doi.org/10.1145/3663363>

18. Pimpalkar P., Ingle G., Sonewane R., Lad V., Bangre R. Deep Learning for Sentiment Analysis: A Survey. *International Journal on Advanced Computer Theory and Engineering*. 2025. Vol. 14(1). P. 347–351. URL: <https://journals.mriindia.com/index.php/ijacte/article/view/479>

19. Петров К. Е., Воробйов Є. К., Кобзев І. В. Синтез моделі класифікації діалогових актів на основі використання рекурентних нейронних мереж. *АСУ та прилади автоматики*. 2022. Вип. 178. С. 4–12. URL: <https://doi.org/10.30837/0135-1710.2022.178.004>

20. Howard J., Gugger S. Deep Learning for Coders with Fastai and PyTorch. Sebastopol: O'Reilly Media, 2020. 624 p.

21. Emotions dataset for NLP classification tasks. <https://www.kaggle.com/datasets/praveengovi/emotions-dataset-for-nlp> (дата звернення: 05.05.2025).

22. Understanding Precision, Recall, and F1 Score Metrics. <https://medium.com/@piyushkashyap045/understanding-precision-recall-and-f1-score-metrics-ea219b908093> (дата звернення: 05.06.2025).

23. Yang C., Cao J. Interpretable Sentiment Analysis Using the Attention-Based Multiple Instance Classification Model: An Application to Wine Reviews. *Harvard Data Science Review*. 2025. Vol. 7(2). P. 1–11. URL: <https://doi.org/10.1162/99608f92.caab9466>

Надійшла до редколегії 13.07.2025 р.

Петров Костянтин Едуардович, доктор технічних наук, професор, завідувач кафедри ІУС ХНУРЕ, м. Харків, Україна, e-mail: kostiantyn.petrov@nure.ua, ORCID: <https://orcid.org/0000-0003-1973-711X>.

Боков Ігор Петрович, здобувач вищої освіти, група СШМ-23-2, факультет комп'ютерних наук ХНУРЕ, м. Харків, Україна, e-mail: igor.bokov@nure.ua

Кобзев Ігор Володимирович, кандидат технічних наук, доцент, доцент кафедри мультимедійних систем і технологій ХНЕУ ім. Семе́на Кузне́ця, м. Харків, Україна, e-mail: ikobzev12@gmail.com, ORCID: <https://orcid.org/0000-0002-7182-5814>