

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ

Харківський національний університет радіоелектроніки
Факультет Комп'ютерних наук
Кафедра Програмної інженерії

КВАЛІФІКАЦІЙНА РОБОТА

Пояснювальна записка

другий (магістерський)
(рівень вищої освіти)

Дослідження алгоритмів нейронних мереж для розпізнавання ментальних розладів людини за текстом

Виконав:
студент 2 курсу групи ІПЗМ-20-4

Головачова О.А.
(прізвище, ініціали)

Спеціальність 121 – Інженерія програмного забезпечення

Тип програми	Освітньо-наукова
--------------	------------------

	<u>доц. Афанасьєва І.В.</u>
Керівник	(посада, прізвище, ініціали)

Допускається до захисту

Зав. Кафедри

З.В. Дудар

2022 p.

Харківський національний університет радіоелектроніки

Факультет _____ Комп'ютерних наук _____
Кафедра _____ Програмної інженерії _____
Рівень вищої освіти _____ другий (магістерський) _____
Спеціальність _____ 121– Інженерія програмного забезпечення _____
(код і повна назва)
Тип програми _____ освітньо-наукова програма _____
Освітня програма _____ Інженерія програмного забезпечення _____

ЗАТВЕРДЖУЮ:

Зав. кафедри _____

(підпис)

«__» _____ 202_ р.

ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ

студента _____ Головачової Оксани Анатоліївни _____

(прізвище, ім'я, по батькові)

1. Тема роботи «Дослідження алгоритмів нейронних мереж для розпізнання ментальних розладів за текстом»

затверджена наказом університету від « 24 » березня 2022р. № 412 Ст

2. Термін подання студентом роботи до екзаменаційної комісії « 21 » травня 2022р.

3. Вихідні дані до роботи алгоритми та порівняння нейронних мереж, розробка гібридного підходу для розпізнання ментальних розладів за текстом.

4. Перелік питань, що потрібно опрацювати в роботі мета роботи, аналіз предметної галузі та постановка задачі, дослідження алгоритмів глибокого навчання, розробка гібридного підходу, алгоритми нейронних мереж для розпізнавання ментальних розладів за текстом..

КАЛЕНДАРНИЙ ПЛАН

№	Назва етапів роботи	Терміни виконання етапів роботи	Примітка
1	Аналіз предметної галузі та постановка задачі	17.01.2022	виконано
2	Аналіз існуючих методів та алгоритмів	31.01.2022	виконано
3	Виконання експериментальної частини	21.02.2022	виконано
4	Підготовка пояснювальної записки	03.04.2022	виконано
5	Підготовка презентації та доповіді	03.04.2022	виконано
6	Перевірка на академічний плагіат	11.05.2022	виконано
7	Нормоконтроль	14.05.2022	виконано
8	Рецензування	15.05.2022	виконано
9	Занесення диплома в електронний архів	16.05.2022	виконано
10	Попередній захист	17.05.2022	виконано
11	Допуск до захисту у зав. кафедри	18.05.2022	виконано
12	Захист кваліфікаційної роботи	21.05.2022	

Дата видачі завдання 17 січня 2022 р.

Студент _____
(підпис)

Керівник роботи _____ доц. Афанасьєва І.В.
(підпис) (посада, прізвище, ініціали)

РЕФЕРАТ / ABSTRACT

Кваліфікаційна робота магістра містить: 74 с., 19 рис., 4 табл., 33 джерел, 6 додатків.

ГЛИБОКЕ НАВЧАННЯ, ДЕПРЕСІЯ, НЕЙРОННІ МЕРЕЖІ, ТЕКСТ, РОЗПІЗНАННЯ, ЕМОЦІЇ, BERT, ДАНІ.

Мета роботи – дослідити алгоритми нейронних мереж для розпізнавання ментальних розладів людини за тестом.

Об’єкт дослідження – розпізнавання ментальних розладів людини за текстом.

Предмет дослідження – алгоритми нейронних мереж для розпізнавання ментальних розладів людини за текстом.

У роботі проведено порівняльний аналіз найсучасніших архітектур нейронних мереж, в результаті якого визначено ефективність роботи з текстовими даними для виявлення ментальних розладів людини.

У роботі було розроблено гібридний підхід та проаналізовано його роботу на тестових наборах даних.

Було проведено аналіз між існуючими роботами та створеним підходом.

У якості методів розробки та аналізу була використана мова програмування Python, фреймворк для розробки нейронних мереж Keras, середа розробки Jupyter Notebook.

DEEP LEARNING, DEPRESSION, NEURAL NETWORKS, TEXT, RECOGNITION, EMOTIONS, BERT, DATA.

The aim of the work is to investigate the algorithms of neural networks for the recognition of human mental disorders by the test.

The object of research is the recognition of human mental disorders by the text.

The subject of research is neural network algorithms for recognizing human mental disorders by text.

The comparative analysis of the most modern architectures of neural networks is carried out in the work, as a result of which the efficiency of work with textual data for detection of mental disorders of the person is defined.

A hybrid approach was developed and its work on test data sets was analyzed.

An analysis was conducted between the existing works and the established approach.

Python programming language, Keras neural network development framework, Jupyter Notebook development environment were used as development and analysis methods.

Я, Головачова Оксана Анатоліївна, студент гр. ІПЗм-20-4, здобувач вищої освіти на другому (магістерському) рівні кафедри «Програмна інженерія», заявляю: моя кваліфікаційна робота на тему «Дослідження алгоритмів нейронних мереж для розпізнавання ментальних розладів за текстом», що буде представлена в екзаменаційну комісію для публічного захисту, виконана самостійно, в ній не містяться елементи плагіату і вона може бути опублікована в електронному архіві відкритого доступу ElAr KhNURE. Всі запозичення з друкованих та електронних джерел мають відповідні посилання. Я ознайомлений (а) з діючим положенням «Про протидію академічному плагіату в ХНУРЕ», згідно з яким виявлення плагіату є підставою для відмови в допуску кваліфікаційної роботи до захисту та застосування дисциплінарних заходів.

ЗМІСТ

Вступ.....	9
1 Аналіз предметної галузі	12
1.1 Аналіз літератури	12
1.2 Алгоритми нейронних мереж	17
1.3 Методи нейронних мереж.....	21
1.4 Постановка задачі.....	25
2 Опис методів досліджень.....	26
2.1 Попередня обробка даних	26
2.2 Метрики	30
2.3 Гібридний підхід	31
3 Експериментальні дослідження.....	36
3.1 Порівняння підходу з іншими алгоритмами	36
3.2 Дата-сети.....	40
3.3 Порівняння алгоритму з вже існуючими системами	45
Висновки.....	51
Перелік джерел посилання	54
Перелік джерел посилання за науковими напрямками керівника та науковців кафедри програмної інженерії.....	58
ДОДАТОК А Звіт результатів перевірки кваліфікаційної роботи на унікальність тексту	59
ДОДАТОК Б Слайди презентації.....	60
ДОДАТОК В Апробація результатів роботи.....	69
ДОДАТОК Г Експертний висновок результатів перевірки кваліфікаційної роботи на відповідність оформлення вимогам ДСТУ	70
ДОДАТОК Д Фрагмент коду програми.....	71

ДОДАТОК Е Впровадження системи 74

СКОРОЧЕННЯ ТА УМОВНІ ПОЗНАКИ

LSTM – Long Short-Term Memory

BI-LSTM – Bi-Directional Long Short-Term Memory

RNN – Recurrent neural network

CNN – Convolutional neural network

NLP – Natural Language Processing

CBoW – The Continuous Bag of Words

ML – Machine Learning

BERT – Bidirectional Encoder Representations from Transformers

NB – Naïve Bayes

NLTK – Natural Language Toolkit

RF – Random Forest

ВСТУП

Здоров'я людини – є одним з найважливіших складових життя людини, бо без нього неможливо повноцінно існувати у нашому суспільстві. У наш час створено багато медичних закладів, приватних місць, таких як, спортивні зали, реабілітаційні зали, різні корти, а також програмного забезпечення для регулювання та відновлення фізичного здоров'я людини. Так, наприклад, є багато мобільних застосунків для контролю: кількості випитої води, щоденних фізичних вправ, метрик здоров'я та отримання статистики з рекомендаціями. Але багато людей забувають, що треба піклуватись не тільки про фізичне здоров'я, а і про ментальне або психологічне здоров'я.

Психічне здоров'я впливає на поведінку, мислення та настрій людей, які взаємодіють із навколишнім світом. Загалом кількість людей з ознаками депресії в 2015 році прогнозувалась на рівні 4,4% (понад 332 млн осіб) [1].

Всесвітня організація охорони здоров'я (ВООЗ) визначає психічне здоров'я як стан благополуччя, в якому людина може реалізувати свій власний потенціал, справлятися із звичайними життєвими стресами, може плідно та ефективно працювати і мати можливість робити внесок у життя своєї спільноти [2]. За даними ВООЗ, кожна четверта людина в світі буде вражена психічним або неврологічним розладом в якийсь момент її життя.

Молоді люди особливо схильні до ризику розвитку психічних порушень у процесі переходу від залежного до самостійного або дорослого життя. У людей з'являються нові потреби або зміни до яких потрібно адаптуватись, наприклад, нові відносини, пандемія, народження дитини, смерть близької людини, звільнення з роботи тощо.

Депресія є одним з найпоширеніших розладів психічного здоров'я, і велика кількість людей з депресією щороку кінчають життя самогубством. Люди, які

страждають на депресію, не звертаються до психологів, тому що відчують сором або не знають про депресію, не усвідомлюють, що це може призвести до серйозної затримки діагностики та лікування.

Тим часом результати аналізу показують, що дані соціальних мереж дають цінні підказки про стан фізичного та психічного здоров'я. У наш час стали дуже розвинуті соціальні мережі, наприклад, Instagram, Facebook або Twitter, де люди можуть обмінюватись фотографіями, різними думками або вести свій блог, де описують кожний крок свого життя. У наш час для людей стали більш близькі твіти, ніж розмови із спеціалістами про їх ментальні проблеми.

Однак використовувані методи представлення тексту та глибокого навчання надають лише обмежену інформацію та знання про різні тексти, опубліковані користувачами. У роботі була запропонована структура для ефективної ідентифікації повідомлень, пов'язаних із депресією та тривогою, зберігаючи контекстне та семантичне значення слів, які використовуються в усьому корпусі при застосуванні BERT. Була розроблена власна система збору даних з Reddit [3] і Twitter [4], які є найпоширенішими сайтами соціальних мереж.

Створений гібридний підхід було порівняно з іншими алгоритмами, такими як RF, SVM, CNN, MLP та Bi-LSTM. Були використані word2vec і BERT з Bi-LSTM, щоб ефективно аналізувати та виявляти ознаки депресії та тривоги з публікацій у соціальних мережах. Створена система перевершує інші найсучасніші методи та досягає точності 98%.

У роботі проведено порівняльний аналіз роботи найсучасніших архітектур нейронних мереж, в результаті якого визначено ефективність роботи з текстовими даними для вилучення ознак депресії, щоб зробити висновок про ментальний стан людей.

Практичні результати кваліфікаційної роботи можуть бути використані для створення програмного забезпечення для моніторингу ментального стану суспільства у соціальних мережах, як на державному рівні, щоб використовувати отримані

результати для оцінки думки суспільства щодо важливих змін чи питань, або, як приклад, використовувати у різних компаніях для вистежування ментального стану людини, та у разі необхідності допомогти працівнику прийти у норму. Зараз це є актуальною темою, та у Японії вже починають вводити аналіз за ментальним станом працівників, щоб стежити, що людина у нормі, там не тільки проводять тести, а й створюють різні програми для моніторингу.

Результати експерименту та опис створеного гібридного підходу були опубліковані у дванадцятій міжнародній науково-технічній конференції «Сучасні напрями розвитку інформаційно-комунікаційних технологій та засобів управління» у другому томі, секції п'ятій.

1 АНАЛІЗ ПРЕДМЕТНОЇ ГАЛУЗІ

1.1 Аналіз літератури

Глибоке навчання – це сукупність алгоритмів, що використовуються в машинному навчанні для моделювання абстракцій високого рівня даних за допомогою модельних архітектур, які складаються з безлічі нелінійних перетворень. Це частина широкого сімейства методів, що застосовуються для машинного навчання, які ґрунтуються на навчальних уявленнях даних. Більшість досліджень щодо виявлення депресії базуються на текстових даних або методах опису особистості з використанням публікацій у соціальних мережах для вибору ознак.

Підходи штучного інтелекту, такі як машинне навчання, були використані в кількох дослідженнях ранньої класифікації депресії. Сімс і Дж. Радж аналізували глибоке навчання для виявлення ранньої стадії депресії у соціальних мережах. Ця дослідницька стаття зосереджується на емоційній процедурі, мовній основі та прогнозуванні аналізу даних, таких як публікації в Інтернеті. Окремі класифікатори можуть працювати незалежно, наприклад дерево рішень, машина опорних векторів, випадковий ліс, наївний Метод Байєса та запропоновані ними гібридні методи в поєднанні зі статистикою ймовірностей. Були протестовані 2500 твітів з Твіттер датасету, та найкращий результат був показаний з використанням їх гібридного методу: злиття SVM та алгоритму Баєса [5].

Орабі запропонував використання техніки безперервного мішка слів (CBOW) для виявлення депресії з набору даних Twitter. У своєму аналізі він використовував 1145 Твіттер користувачів, де було більше 2000000 твітів. Кожний користувач був помічений, як «контроль» та «депресія», також кожен користувач мав мітку статі та віку. Головна ціль підходу: виявлення користувачів соціальної мережі Твітер, які знаходяться у зоні ризику та можуть у майбутньому мати депресію. У експерименті найкращий результат був зроблений з допомогою CNN моделі [6].

СВoW – архітектура, яка передбачає поточне слово, виходячи з навколишнього контексту. Ця архітектура описує як саме нейронна мережа вчиться на різних датасетах та запам'ятовує представлення слів. Це звичайна модель мішка слів з використанням чотирьох ближніх сусідів терміну без рахунку порядку. Процес тренування влаштований таким чином: беруться послідовно $(2k + 1)$ слів, слово в центрі є тим словом, яке має обчислюватися. А сусідні слова є контекстом довжини по k з кожного боку. Кожному слову в моделі ставиться у відповідність унікальний вектор, який змінюється в процесі навчання моделі [7].

Також Міколовим відразу був запропонований й інший підхід – протилежний до SBOW, який він назвав skip-gram, тобто «словосполучення з пропуском». Він відрізняється тим, що ми намагаємося з даного нам слова вгадати його контекст (точніше вектор контексту). У всьому іншому модель має відмінностей [8].

CNN – це згорткова нейронна мережа, яка використовується для класифікації та розпізнання різних об'єктів: обличч на фотокартках, розпізнання мови. Кожний шар нейронної мережі використовує своє перетворення. За допомогою такої обробки можна не тільки класифікувати об'єкт, а й знайти його на зображення. Ця архітектура нейронних мереж, запропонована Яном Лекуном, від самого початку націлена на ефективне розпізнавання зображень. Згорткова нейронна мережа зазвичай являє собою чергування шарів згортки (convolution layers, C-layers), шарів субдискретизації (subsampling layers, S-layers) і наявності повнозв'язних шарів (fully-connected layer, F-layers) на виході. Всі три види шарів можуть чергуватися в довільному порядку [9].

Кім Джін досліджував машинне навчання для ментального здоров'я у соціальних мережах. Це дослідження включало бібліометричний аналіз публікацій, пов'язаних із психічним здоров'ям у соціальних мережах з 2015 по 2020 рік із двома базами даних цитування, WoS та Scopus, а також аналізом тенденцій [10].

Wonkoblar.A запропонував метод глибокої нейронної мережі для аналізу депресії в соціальних мережах. Вони використали для аналізу набір даних Twitter: 3682 користувачів, 1983 з яких були визначені як користувачі з депресією, та 1699 не

знаходили у собі ознак депресії. Було розроблено дві моделі навчання з кількома екземплярами (MIL) з анафоричним кодером роздільної здатності та без нього, щоб ідентифікувати користувачів із депресією та виділити повідомлення, пов'язані з психічним здоров'ям автора [11].

Qun Nisa використовував згорткові нейронні мережі для порівняння різноманітних моделей з використанням Bert. У якості датасету були використані колекції текстів користувачів, коментарів користувачів, назви постів з форуму Reddit [12].

BERT – це мовна модель від Гугл. Мовна модель приймає на вхід ембедінг-токенів і видає результат залежно від завдання. Існує стандартний набір завдань, який потрібно виконати на стандартному наборі даних, щоб довести, що нейромережа справляється з розумінням тексту. Приклад завдання – видати 1, якщо у двох різних питаннях на Quora запитують те саме, видати 0 – якщо ні. Стандартні завдання називаються в NLP бенчмарками. BERT тестували на наборах бенчмарків GLUE – «Оцінка загального розуміння мови», SQuAD та SWAG.

Ibitoye A.O розглянув два дослідження, які використовували передбачувану здатність класифікаторів машинного навчання для аналізу взаємодії емоцій. Вони класифікували дописи про депресію в соціальних мережах із застосуванням методів класифікації [13].

LSTM (long short-term memory) — тип рекурентної нейронної мережі, здатний навчатися довгостроковим залежностям. LSTM спеціально розроблені для вирішення проблеми довгострокової залежності [14]. Їхня спеціалізація – запам'ятовування інформації протягом тривалих періодів часу, тому їх практично не потрібно навчати. Всі рекурентні нейронні мережі мають форму ланцюжка модулів нейронної мережі, що повторюються.

Дослідники некомерційної групи OpenAI із Сан-Франциско навчили нейронну мережу LSTM точніше розпізнавати емоційну складову тексту. Тепер машина майже безпомилково впізнає настрій у відгуках покупців на Amazon та кінорецензіях на

Rotten Tomatoes. Запропонована модель набагато перевершує існуючі методи розпізнавання настрою тексту. Це було доведено під час аналізу одноманітних відгуків та рецензій з інтернету. Програма точно визначає позитивний та негативний зміст тексту. Дослідники OpenAI застосували у цій роботі вид рекурентної нейронної мережі – LSTM (Long short-term memory – «Довга короткострокова пам'ять»). Цей вид нейронної мережі добре зарекомендував себе у розпізнаванні написаного тексту та людської мови. Їй належить рекорд мінімуму помилок при розпізнаванні мови – 17,7%. Застосування методики зазвичай полягали у пошуку ключових слів. І це завжди ефективний метод. Наприклад, вираз «Сподіваюся, ти щасливий» можна віднести до текстів із позитивною конотацією лише тому, що в ньому є слово «щасливий». Але якщо цей вислів поставити в контекст – Ти вбив його. Сподіваюся, ти щасливий» - висновок виявиться помилковим [15].

Таблиця 1.1 – Порівняння досліджень стосовно виявлення ментальних розладів

Автор	Рік	Методи	Висновки	Обмеження
Aldarwish	2017	SVM, NB	NB дає високу точність	Дата-сет є арабським
Shen G.	2017	MDL, MSNL, NB, WDL	MDL точність є високою	Фокус на зізнання юзерів
Islam M.R.	2018	DT, SVM, k-nearest	DT дає найкращі результати	Фокус на коментарі
Nguyen T.L.	2018	DT, SVM, Ensemble	DT дає найкращу точність	Фокус на коментарі
Gaikar M.	2019	NB, SVM hybrid model	85 % точність	Витрачається багато часу

Кінець таблиці 1.1

Автор	Рік	Методи	Висновки	Обмеження
Burdisso S.G.	2019	SS3, NB, SVM, KNN	SS3 – найкращі результати	Витрачається багато часу
Farima	2019	MLP, SVM, NB	MLP – найкращі результати	Обмежений дата-сет
Alsagri	2020	SVM, DT, NB	SVN – дає найкращу точність	Неможливо уникнути переповненості даних
Kim J.	2020	CNN, XGBoost	CNN – дає найкращу точність	Обмежений обсяг
Filho De Souza	2021	DT, SVM, RF, GB, LR	RF – дає найкращу точність	Фокус на зізнання юзерів

Попередні фреймворки були зосереджені на виявленні депресії за допомогою обмежених фраз, речень. Для діагностики депресії більшість дослідників використовували єдиний метод машинного навчання. Крім того, дослідники зосередилися на значення точності, щоб визначити, який алгоритм найбільш підходить для ситуації.

1.2 Алгоритми нейронних мереж

Розглянемо CNN. CNN – клас нейронних мереж, який спеціалізується на обробці даних, які мають топологію у вигляді сітки, наприклад зображення. Цифрове зображення – це бінарне представлення візуальних даних. CNN зазвичай має три рівня: згортковий шар, шар об'єднання та повністю пов'язаний шар [16].

CNN використовують різновид багат шарових перцептронів, розроблений так, щоб вимагати використання мінімальної обробки. Виходячи з їхньої архітектури спільних ваг та характеристик інваріантності відносно паралельного перенесення. ЗНМ використовують порівняно мало попередньої обробки, в порівнянні з іншими алгоритмами класифікації зображень. Це означає, що мережа навчається за допомогою фільтрів, що в традиційних алгоритмах приходиться розробляти вручну. Ця незалежність у конструюванні ознак від апріорних знань та людських зусиль є великою перевагою. CNN складається з шарів входу та виходу, а також із декількох прихованих шарів. Приховані шари CNN зазвичай складаються зі згорткових шарів, агрегувальних шарів, повноз'єднаних шарів та шарів нормалізації. Згорткові шари застосовують до входу операцію згортки, передаючи результат до наступного шару. Згортка імітує реакцію окремого нейрону на зоровий стимул. Згортковий шар є основним будівельним блоком CNN. Він несе основну частину обчислювального навантаження мережі. Цей рівень виконує скалярний добуток між двома матрицями, де одна матриця – це набір навчальних параметрів, інакше званих ядром, а інша матриця є обмеженою частиною сприймаючого поля. Ядро просторово менше зображення, але має більшу глибину. Це означає, що якщо зображення складається з трьох RGB-каналів, висота та ширина ядра будуть просторово малі, але глибина поширюється на всі три канали.

Шар пулінгу (об'єднання) замінює вихідні дані мережі у певних місцях, отримуючи зведену статистику найближчих виходів. Це допомагає зменшити просторовий розмір уявлення, що зменшує необхідну кількість обчислень та ваг.

Нейрони у повністю пов'язаному шарі мають повний зв'язок з усіма нейронами попереднього та наступного шару, як це видно у звичайному FCNN (Fully-Connected Convolutional Neural Network). Ось чому його можна обчислити як завжди, множенням матриці з подальшим ефектом усунення.

Повністю пов'язаний шар допомагає зіставити уявлення між даними на вході та виході.

Розглянемо MLP. Багатошаровий персептрон – це нейронна мережа, область науки, в якій досліджується, як прості моделі біологічного мозку, можуть використовуватися для вирішення складних вирішувальних задач, таких як прогнозування в машинному навчанні (ML) [17].

Мета полягає не у створенні реалістичних моделей мозку, а в розробці надійних алгоритмів, які використовуються для моделювання складних проблем.

Сила нейронних мереж полягає в їх здатності вивчати патерни в Тренувальні дані і як найкраще пов'язати його з Цільовою змінною яку ми хочемо передбачити. Математично вони здатні створювати будь-яку функцію та зарекомендували себе як універсальний алгоритм апроксимації, тобто заміні одних об'єктів на інші, більш спрощені.

Прогностична здатність нейронних мереж обумовлена ієрархічною або багаторівневою структурою мереж. Структура даних може виділяти ознаки у різних масштабах і об'єднувати в ознаки вищого порядку.

Розглянемо LSTM. LSTM-мережі складаються з повторюваних елементів. Кожен такий елемент містить чотири шари і відрізняється тим, що має комірку довгої короткочасної пам'яті. LSTM розроблено спеціально, щоб уникнути проблеми довготривалої залежності. Запам'ятовування інформації на довгі періоди часу – це їхня звичайна поведінка, а не щось, чому вони намагаються навчитися [18].

Будь-яка рекурентна нейронна мережа має форму ланцюжка модулів нейронної мережі, що повторюються. У звичайній RNN структура одного такого модуля дуже проста, наприклад, він може бути одним шаром з функцією активації $\tanh$ (гіперболічний тангенс). Структура LSTM також нагадує ланцюжок, але модулі виглядають інакше. Замість одного шару нейронної мережі вони містять чотири, і ці шари взаємодіють особливим чином.

Ключовий компонент LSTM – це стан осередку (cell state) – горизонтальна лінія, що проходить у верхній частині схеми.

Стан осередку нагадує конвеєрну стрічку. Вона проходить безпосередньо через весь ланцюжок, беручи участь лише у кількох лінійних перетвореннях. Інформація може легко йти по ній, не змінюючись.

Тим не менш, LSTM може видаляти інформацію зі стану комірки; цей процес регулюється структурами, які називаються фільтрами. Фільтри дозволяють пропускати інформацію на підставі певних умов.

Розглянемо BI-LSTM. Двонаправлені рекурентні нейронні мережі (RNN) насправді просто об'єднують дві незалежні RNN. Ця структура дозволяє мережам мати як зворотну, так і пряму інформацію про послідовність на кожному кроці часу. Використання двонаправленого буде виконувати ваші введення двома способами: один з минулого в майбутнє, а другий з майбутнього в минуле, і те, що цей підхід відрізняється від односпрямованого, полягає в тому, що в LSTM, який працює назад, ви зберігаєте інформацію з майбутнього, а за допомогою двох прихованих станів у поєднанні ви здатні в будь-який момент часу зберегти інформацію як з минулого, так і з майбутнього.

Розглянемо SVM. Метод Опорних Векторів або SVM – це лінійний алгоритм, що використовується в задачах класифікації та регресії. Цей алгоритм має широке застосування практично і може вирішувати як лінійні і нелінійні завдання. Суть роботи: алгоритм створює лінію або гіперплощину, яка поділяє дані на класи [19].

Основна ідея методу полягає у відображенні векторів простору ознак, які представляють об'єкти, що класифікуються, в простір більш високої розмірності. Це пов'язано з тим, що у просторі більшої розмірності лінійна роздільність множини виявляється вище, ніж у просторі меншої розмірності. Причини цього інтуїтивно зрозумілі: що більше ознак використовується для розпізнавання об'єктів, то вище очікувана якість розпізнавання.

Після переведення в простір більшої розмірності, в ньому будується роздільна гіперплощина. При цьому всі вектори, розташовані з однієї сторони гіперплощини, відносяться до одного класу, а розташовані з іншого – до другого. Також, по обидва боки основної розділяє гіперплощини, паралельно їй і на рівній відстані від неї будуються дві допоміжні гіперплощини, відстань між якими називають зазор.

Передбачається, що розділяюча гіперплощина, побудована за цим правилом, забезпечить найбільш впевнений поділ класів і мінімізує середню помилку розпізнавання.

Розглянемо RF. Випадковий ліс – алгоритм машинного навчання, створений Лео Брейманом та Адель Катлер.

RF (random forest) – це безліч вирішальних дерев. У задачі регресії їх відповіді усереднюються, у задачі класифікації приймається рішення голосуванням у більшості [20]. Усі дерева будуються незалежно за такою схемою:

Вибирається підвибірка навчальної вибірки розміру `samplesize` (м.б. з поверненням) – по ньому будується дерево (для кожного дерева – своя підвибірка).

Для побудови кожного розщеплення у дереві переглядаємо `max_features` випадкових ознак (для кожного нового розщеплення – свої випадкові ознаки). Вибираємо найкращі ознаки та розщеплення по ньому (за заздалегідь заданим критерієм). Дерево будується, як правило, до вичерпання вибірки (поки в листі не залишаться представники лише одного класу), але в сучасних реалізаціях є параметри, які обмежують висоту дерева, кількість об'єктів у листі та кількість об'єктів у підвибірці, при якому проводиться розщеплення.

Така схема побудови відповідає головному принципу ансамблювання (побудови алгоритму машинного навчання на базі кількох, у даному випадку вирішальних дерев): базові алгоритми повинні бути хорошими та різноманітними (тому кожне дерево будується на своїй навчальній вибірці та при виборі розщеплень є елемент випадковості) .

1.3 Методи нейронних мереж

Нейромережеві моделі зазнають складнощів при необхідності роботи з розрідженими категоріальними ознаками. Ембеддинги – це можливість зменшення розмірності таких ознак для підвищення продуктивності моделі. Перш ніж говорити про структуровані набори даних, корисно розібратися з тим, як зазвичай використовуються ембеддинги. У сфері обробки природних мов (Natural Language Processing, NLP) зазвичай працюють зі словниками, що з тисяч слів. Ці словники зазвичай вводять у модель з використанням методики швидкого кодування (One-Hot Encoding), що, в математичному сенсі, рівносильне наявності окремого стовпця для кожного зі слів. Коли слово передається в модель, у відповідному стовпці виявляється одиниця, тоді як у всіх інших виводяться нулі. Це веде до появи дуже сильно розріджених наборів даних. Вирішення цієї проблеми полягає у створенні ембеддингу. Наприклад, «кокос» і «білий ведмідь» – це слова, які семантично досить різні, і ембеддінг представлятиме їх як вектори, які будуть розташовані на великій відстані один від одного; але слова «кухня» та «вечеря» є спорідненими словами, тому вони будуть розташовані ближче одне до одного.

GloVe (глобальні вектори для представлення слів) – це альтернативний метод створення вбудованих слів. Він базується на методах розкладання матриць на матрицю слова-контексту. Основою алгоритму є побудова великої матриці інформації

про спільну появу, і підрахунок кожного слова (рядка), і те, як часто це слово зустрічається в якомусь контексті (стовпці) у великому корпусі.

Розглянемо алгоритм GloVe:

- збираємо статистичні дані про спільне використання слів у формі матриці спільного використання слів X . Кожен елемент X_{ij} такої матриці відображає, як часто слово i з'являється в контексті слова j . Зазвичай сканується корпус наступним чином: для кожного терміна шукається терміни контексту в межах певної області, визначеної розміром вікна перед терміном та розміром вікна після терміна. Крім того, надається менше ваги для більш віддалених слів;

- визначаємо м'які обмеження для кожної пари слів;

- визначаємо функцію витрат.

Skip-gram ігнорує структуру слова, але деякі мови мають складові слова, як, наприклад, німецькою. Тому до основної моделі було додано subword-модель. Subword-модель - це уявлення слова через ланцюжки символів (n -грами) з n від 3 до 6 символів від початку до кінця слова плюс слово повністю. Наприклад, слово замок з $n = 3$ буде представлено n -грамами <за, зам, амо, мок, ок > і послідовністю <замок>. Таким чином, для моделі є різниця між послідовністю <зам> у слові зам - і n -грамою зам із слова замок. Такий підхід дозволяє працювати з тими словами, які модель раніше не зустрічала.

Алгоритм word2vec приймає текстовий корпус і генерує еквівалентний вектор, пов'язаний із заданим введенням. Згенерований вектор матиме великі розміри, і кожному слову, включеному в текст, буде надано відповідний вектор у векторному просторі. Вектори розміщені в просторі таким чином, щоб слова зі схожим контекстним значенням залишалися ближче один до одного. Однак алгоритм word2vec не може представити нове слово за допомогою вектора, якщо воно не включено в навчальний набір даних. CBOW – це модель, яка знаходить цільове слово за допомогою сусідніх контекстних слів, тоді як модель пропуску грам знаходить контекстні слова, надані центральним словом як вхідні дані. У запропонованій нами

системі модель пропуску грам була навчена використовувати її покращену продуктивність із узгодженими словами.

Розглянемо BERT. Це нещодавній підхід, запропонований дослідниками, які використовують мову Google AI. BERT використовує архітектуру трансформатора, механізм уваги, який виявляє контекстні контакти між словами [22]. Трансформатор містить два окремих компонента: кодер, який зчитує введений текст, і декодер, який генерує передбачення для завдання. Оскільки головною метою BERT є створення мовної моделі, необхідний лише механізм кодування. Під час навчання BERT використовує дві схеми для захоплення контекстного значення слів як у правому, так і в лівому напрямках.

Модель замаскованої мови (MLM). Перш ніж передавати вхідні послідовності в модель BERT, 15% слів у кожній послідовності були замінені маркером [MASK]. Потім BERT намагається визначити початкове значення замаскованих слів на основі контексту, наданого іншими словами в послідовності. Прогноз здійснюється на основі трьох основних кроків, як показано на рисунку 1.1. Це застосування шару класифікації поверх вихідних даних кодера, помноження вихідних векторів на матрицю вбудовування та перетворення їх у розміри словника, обчислення розподілу ймовірностей для всіх слів у словнику за допомогою softmax.

Наступний етап – це прогноз наступного речення (NSP). Під час навчального процесу модель BERT бере два речення та знаходить певну кореляцію, щоб передбачити, чи є друге речення продовженням (наступним реченням) першого речення в оригінальному документі. Під час фази навчання половина введених даних – це пара, в якій друге речення є наступним реченням у вихідному документі, а інша половина введеного тексту – це друге речення, випадково вибране з корпусу тексту.

BERT використовує кодер з архітектури трансформатора, який являє собою мережу кодер-декодер. BERT буває двох розмірів: базовий BERT (BERTBASE) і великий BERT (BERTLARGE). BERTBASE має 12 шарів у стеку кодерів, тоді як BERTLARGE має 24 шари в стеку кодерів. Архітектури BERT (BASE і LARGE) також

мають більші мережі прямої зв'язку (768 і 1024 прихованих блоків відповідно) і більше уваги (12 і 16 відповідно), ніж архітектура трансформатора, запропонована в оригінальному дослідженні [23]. Він містить 512 прихованих одиниць і 8 головок уваги. BERTBASE містить 110M параметрів, тоді як BERTLARGE має 340M параметрів. Ця модель приймає маркер класифікації (CLS) як перший вхід, за яким слідує послідовність слів, які пересилаються на наступний рівень.

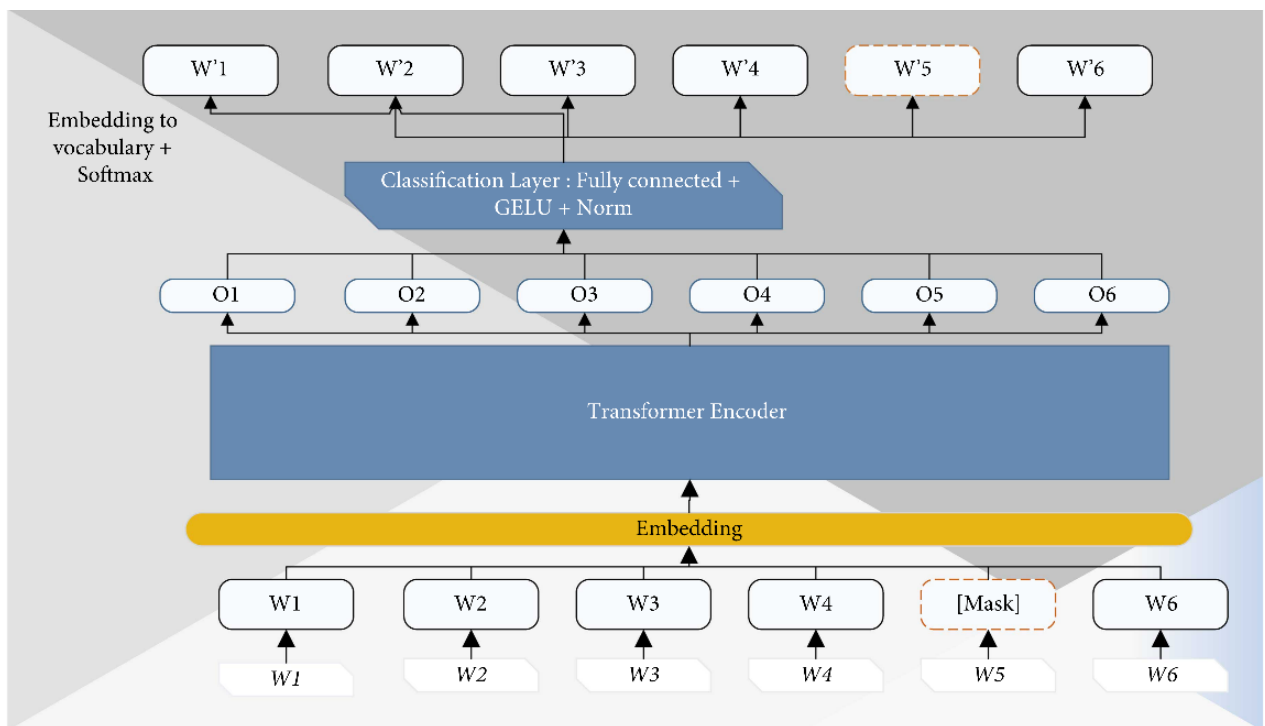


Рисунок 1.1 – Кроки прогнозу BERT

У роботі запропоновано вбудовування слів на основі BERT і точне налаштування для конкретної задачі, пов'язаної зі здоров'ям, на основі гіперпараметрів.

1.4 Постановка задачі

Під час аналізу предметної галузі, а саме літератури та існуючих рішень та постановці проблеми були сформовані вимоги для подальшої роботи, також були виявлені основні напрями та задачі для проведення порівняння.

Основним завдання було поставлено пошук підходу, який дає найкращий показник у якості, для роботи з різними тестовими даними.

Також у даній роботі потрібно реалізувати та проаналізувати різні варіації нейромережових підходів та алгоритмів для роботи з текстом та виявлення у ньому тональності та розпізнання емоцій та, як результат, виявлення ментальних розладів у людей, які створювали ці тексти.

Щоб об'єктивно оцінити та порівняти результат реалізованих варіантів, потрібно розглянути їх роботу на різних вхідних даних. Для цього потрібно зібрати різні тестові дані, які максимально наближені до людських текстів у реальному житті.

При реалізації нейромережових підходів слід виконати наступні завдання:

- для навчання нейронної мережі треба провести збір даних;
- в якості нейронної мережі слід розробити та реалізувати мережу з архітектурою, що підходить для роботи з текстовими даними, яка буде здійснювати класифікацію за емоційною складовою в тексті;
- при реалізації архітектури в програмному коді слід використовувати фреймворк Keras;
- при побудові архітектури слід використовувати технологію обробки текстових даних word embedding.

У результаті потрібно провести комплексну оцінку ефективності отриманих моделей за допомогою обраних метрик.

2 ОПИС МЕТОДІВ ДОСЛІДЖЕНЬ

2.1 Попередня обробка даних

Попередня обробка даних – це метод очищення та фільтрації неякісних даних, щоб їх можна було легко використовувати для вилучення ознак, як показано на рисунку 2.1.

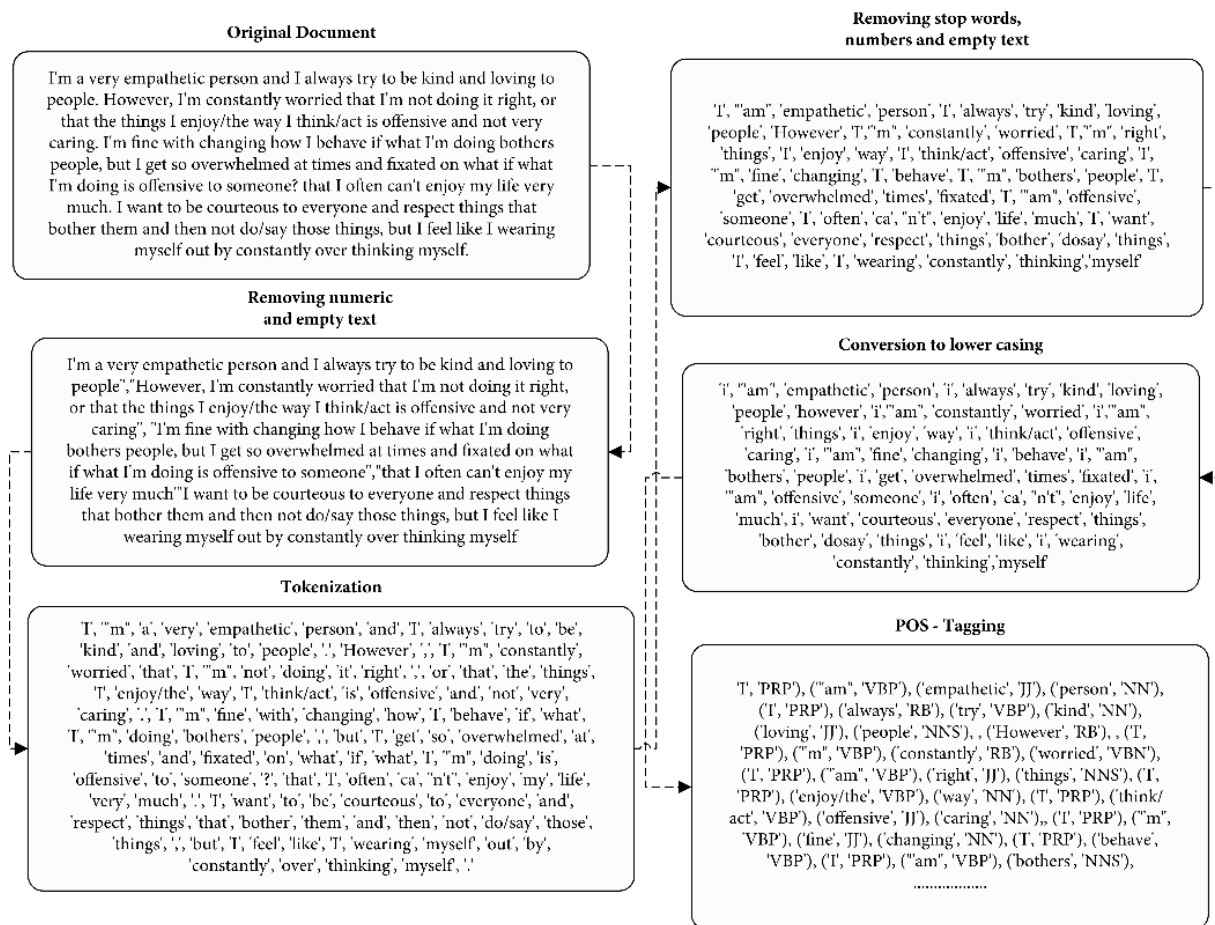


Рисунок 2.1 – Приклад попередньої обробки даних

У реальному світі люди обмінюються інформацією в соціальних мережах неформально, використовуючи тексти, які містять хештеги, спеціальні символи та зайві слова. Потрібно застосувати концепцію машинного навчання та інтелекту даних,

щоб витягти якийсь сенс із такого текстового корпусу, перш ніж подати їх до будь-якої моделі класифікації. Крім того, основним напрямком нашої уваги є заміна жаргону та емодзі відповідним інформативним текстом за допомогою Emojipedia. Відомо, що Інтернет стає комунікаційним засобом, де люди часто використовують розмовну мову та емодзі, щоб передати свої думки та думки. Вилучення змістовного тексту з публікацій у соціальних мережах допомагає зрозуміти контекст і посилити емоції, пов'язані з ним. Були застосовані різні методи попередньої обробки до зібраного тексту, який містить багато жаргонізмів і неформальних слів таким чином, що це допоможе нам визначити депресію та тривожність.

Метод токенизації – процес сегментації тексту на слова або пропозиції. Електронний текст являє собою лінійну послідовність символів (символів, слів або фраз). Природно, перш ніж приступити до реальної обробки тексту, текст повинен бути розділений на лінгвістичні одиниці, такі як слова, числа, цифри і так далі. Цей процес називається токенизацією. В англійській мові слова часто відділені один від одного пробілами, але не всі пробіли рівні. Токенизація – це предобробка, ідентифікація базових одиниць, які належать обробці. Без цих чітко розділених базових одиниць неможливо провести який-небудь аналіз або генерацію. Помилки, зроблені на цьому етапі, викликають більше помилок на більш пізніх етапах обробки тексту. Токенизація – це метод поділу значної кількості тексту на менші частини, відомі як маркери. Ці токени використовуються для виявлення деяких закономірностей і використовуються як вхідні дані для наступних загальних кроків у конвеєрі НЛП, таких як стеммінг і лемматизація. Загалом, великий текст складається з хеш-знаків, розділові знаки та символи, які навіть не є текстом. У запропонованій нами системі цей процес виконується за допомогою TreebankWordTokenizer, який міститься в наборі інструментів природної мови (NLTK), щоб очистити слова, які називаються токенами. Токенизація використовується для скорочення небуквених-цифрових символів і розбиття речень на слова. Нарешті, весь текст у всьому даному файлі представлений пакетом слів для подальшого аналізу.

Лематизація. Лематизація схожа на стеммізацію в тому, що вона наводить слово до його початкової форми, але з однією відмінністю: у разі корінь слова буде існуючим у мові словом. Наприклад, слово "caring" припиниться в "care", а не "car", як у стеммізації. Приклад лематизації зображено на рисунку 2.2.

```
# nltk.download('wordnet')
from nltk.stem import WordNetLemmatizer

lemmatizer = WordNetLemmatizer()

lemm_review = [lemmatizer.lemmatize(word) for word in filtered_review]

print(lemm_review)

['touching', 'movie', 'full', 'emotion', 'wonderful', 'acting', 'could', 'sat', 'second', 'time']
```

Рисунок 2.2 – Приклад лематизації за допомогою Python

Також є стемінг. Алгоритм Stemming працює, вирізаючи суфікс із слова. У більш широкому розумінні відсікає початок або кінець слова.

Навпаки, лематизація є потужнішою операцією і враховує морфологічний аналіз слів. Він повертає лему, яка є базовою формою її флексивних форм. Для створення словників та пошуку правильної форми слова потрібні глибокі лінгвістичні знання. Стеммінг - це спільна операція, а лематизація - інтелектуальна операція, в якій правильна форма виглядатиме у словнику. Отже, лематизація допомагає у формуванні найкращих можливостей машинного навчання.

Також потрібне видалення слів, які не мають суттєвої інформації про проблеми, пов'язані з психічним здоров'ям загалом. Найпоширенішими словами, які вважаються нерелевантними, є займенники, прийменники, символи (наприклад, дати, #), сполучники та артиклі (a, an і the). Крім того, універсальні локатори ресурсів (URL) у будь-яких даних текстових даних слід фільтрувати, оскільки вони не мають корисної інформації для обробки тексту. Щоб видалити стоп-слова з даного речення, ми спочатку розбиваємо наш текст на слова, а потім видаляємо слово, якщо воно

зустрічається в списку стоп-слів, наданому NLTK, як описано вище. Заміна жаргону та емодзі їхнім фактичним текстом за допомогою Emojipedia є надзвичайно важливим кроком, оскільки вони містять корисну інформацію щодо психічного здоров'я. Приклад видалення слів з використанням мови Python зображено на рисунку 2.3.

```
# nltk.download('stopwords') # you will have to download the set of stop words the first time
from nltk.corpus import stopwords

stop_words = stopwords.words('english')

filtered_review = [word for word in tokens if word not in stop_words] #removing stop words.
print(filtered_review)

['touching', 'movie', 'full', 'emotions', 'wonderful', 'acting', 'could', 'sat', 'second', 'time']
```

Рисунок 2.3 – Видалення стоп-слів за допомогою Python

Нижній регістр. Люди зазвичай висловлюють свою думку про те, як вони себе почувають, використовуючи емодзі. Деякі слова можуть бути написані по-різному, наприклад «b4», тобто «до»; «>»), тобто «щасливий»; і «2мого», тобто «завтра». Кожне слово в нашому сценарії перетворюється в його оригінальну і загальну форму, а потім перетворюється в нижній регістр, що зберігає послідовність і дозволяє уникнути плутанини під час аналізу тексту.

Практика класифікації слів за частинами мови та відповідного позначення їх називається позначенням частини мови (POS). Такими мітками можуть бути іменники, дієслова, прикметники, прислівники, визначники та сполучники тощо. Основна причина цього кроку полягає в тому, щоб відкинути будь-який POS, який не має внеску в ідентифікацію депресії та тривоги. Крім того, POS-теги розпізнають іменники та прикметники, які вважаються безцінними індикаторами для аналізу настроїв. Ми використали POS-тегер NLTK для запропонованої системи.

2.2 Метрики

Щоб оцінити модель, були використані показники, такі як precision, recall, accuracy. Матриця помилок – це матриця, яка використовується для оцінки ефективності класифікації, вона показує кількість неправильних прогнозів у порівнянні з кількістю правильних передбачень у табличному порядку. На основі матриці помилок можна обчислити precision, recall, accuracy таким чином :

$$\text{precision} = \frac{TP}{TP+FP} ,$$

$$\text{recall} = \frac{TP}{TP+FN} ,$$

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

Найпоширеніші терміни, які використовуються для обчислення матриці помилок:

- P: фактичний позитивний випадок, який є класом, пов'язаним з депресією в моделі;
- N: фактичний негативний випадок, який у нашій моделі не пов'язаний з депресією;
- істинно-позитивний (TP): випадок, у якому фактичний клас точки даних (зібраних текстових даних) істинний, а клас, передбачений нашою моделлю, також істинний;
- тотальний негатив (TN): випадок, коли фактичний клас точки даних є хибним, а прогнозований клас також є хибним;
- хибно-позитивний (FP): випадок, у якому фактичний клас точки даних дорівнює 0 (неправда), а прогнозований клас дорівнює 1 (правда);

– хибно-від’ємний (FN): випадок, у якому фактичний клас точки даних істинний, а прогнозований клас хибний.

2.3 Гібридний підхід

За допомогою методів глибокого навчання слова представлені у вигляді вектору, а класифікація настроїв на основі глибокого навчання буде застосовано шляхом передачі векторів до моделей класифікації. Рекурентні нейронні мережі (RNN), які є найпоширенішими методами моделювання послідовності, здатні отримувати відповідні характеристики з даних і можуть бути провідним вибором для захоплення семантики довгих текстів. В результаті LSTM був створений для подолання недоліків у підтримці довгострокових залежностей у моделях RNN [36]. Модель LSTM разом із згортковою нейронною мережею (CNN) для класифікації речень забезпечує точні результати і активно використовується в сучасних дослідженнях NLP. Моделі CNN використовують шари згорткового та максимального об’єднання для вилучення найбільш релевантних функцій, тоді як моделі LSTM підтримують зв’язок між словами протягом більш тривалого періоду, використовуючи комірки пам’яті. Тому їх краще використовувати для класифікації тексту. Двонаправлена LSTM була розроблена для вирішення проблем класифікації послідовностей. Як приклад, вона була застосована в аналізі подій трафіку разом з OLDA з використанням даних соціальних мереж і досягла надзвичайно високої точності [24]. На основі результатів було запропоновано двонаправлений LSTM (Bi-LSTM), який використовує два блоки LSTM, які працюють як у лівому, так і в правому напрямках, щоб об’єднати інформацію про минуле та майбутнє з наших зібраних даних із соціальних мереж. Bi-LSTM також зберігає довгостроковий зв’язок між словами разом із дубльованою інформацією про контекст, як показано на рисунку 2.4.

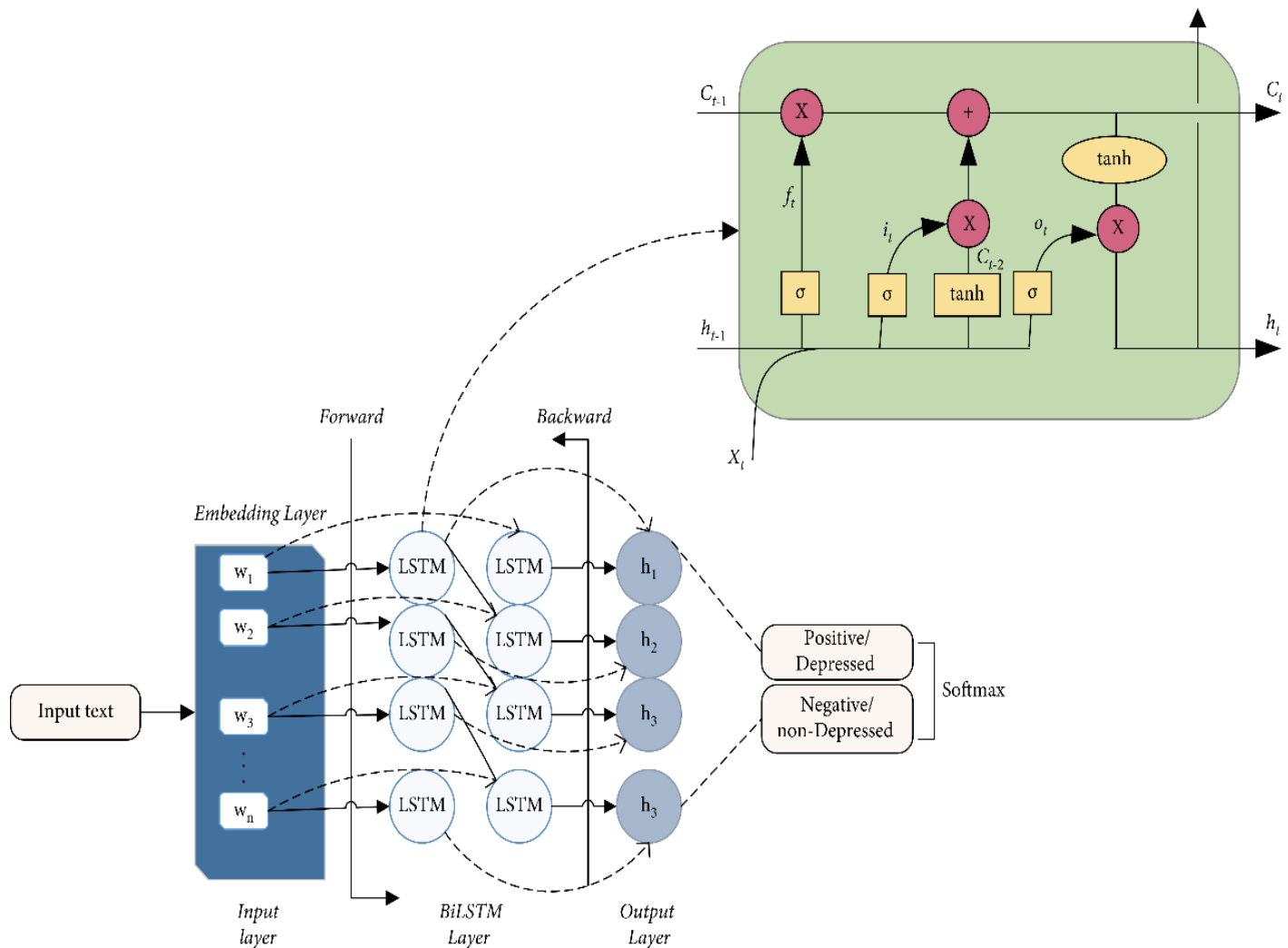


Рисунок 2.4 – Процес Bi-LSTM, що використовується для захоплення послідовних ознак

У запропонованій структурі був використаний аналіз головних компонентів (PCA) перед подачею векторів вбудовування до класифікаторів.

PCA – це метод, основною метою якого є зменшення розмірності набору даних, що складається з багатьох пов'язаних змінних.

На рисунку 2.5 зображено загальну структуру роботи гібридного підходу, який складається з BERT та Bi-LSTM.

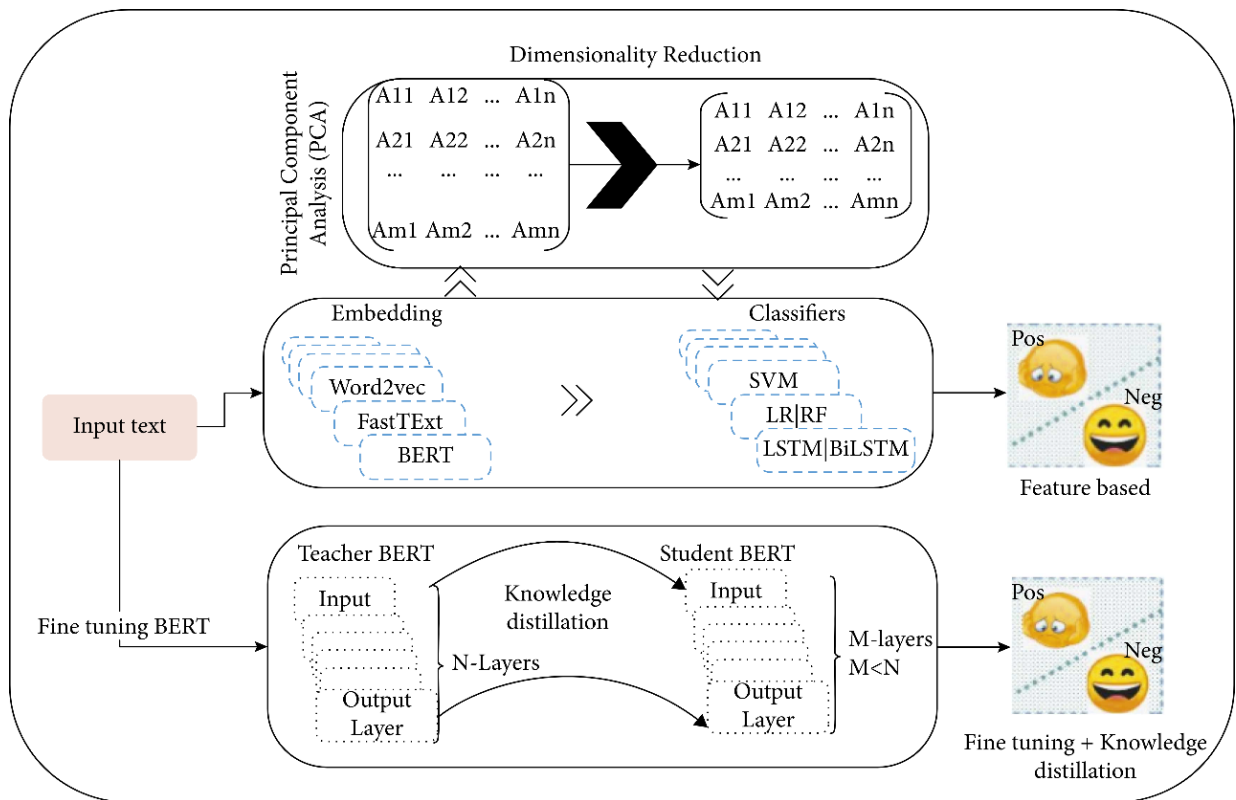


Рисунок 2.5 – Загальна структура для представлення тексту, зменшення розмірності та виявлення депресії/тривоги на основі методів дистиляції знань разом із загальною класифікацією на основі ознак

BERT – це величезна модель, понад 100 мільйонів параметрів. Для тонкого налаштування потрібен не тільки графічний процесор, а й час виводу, а процесора (або навіть багатьох із них) недостатньо. Це означає, що для використання BERT скрізь, потрібно встановлювати GPU, а це недоцільно здебільшого. У 2015 році було представлено спосіб перерозподілу знань з великої нейронної мережі в набагато меншу, таку як вчитель і учень. Було використано передбачення великої нейронної мережі для навчання малого.

Основна ідея полягає у використанні прогнозів, тобто прогнози перед кінцевою функцією активації (зазвичай softmax або сигмовидна). Передбачається, що, використовуючи необроблені значення, модель здатна вивчити внутрішні уявлення краще, ніж за допомогою жорстких передбачень. Sotmax нормалізує значення до 1,

зберігаючи максимальне значення високим і зменшує інші значення до значення, близького до нуля. У нулях інформації мало, тому, використовуючи необроблені прогнози, ми також навчаємось у непередбачуваних класів. Автори показують хороші результати у кількох завданнях, включаючи MNIST та розпізнавання мови.

Були проведені інтенсивні експерименти, щоб використати та порівняти ключові переваги вбудовування та класифікації кожного слова, які останнім часом стали надзвичайно популярними.

Експеримент було проведено на даних Reddit і Twitter, які є класифікацією постів, пов'язаних з депресією/тривогою (позитивні/негативні), з використанням методів аналізу тексту на основі глибокого навчання.

Було запропоновано двосторонній підхід до реалізації перший з яких – це впровадження найсучасніших технік машинного/глибокого навчання для ідентифікації депресії/тривоги разом із зменшенням розмірності для максимальної точності.

Другою найважливішою частиною є застосування налаштування та оптимізації моделі (перебудови знань) для створення легшої та розумнішої моделі виявлення депресії/тривоги. Дистиляція знань тут полягає в створенні стиснутої моделі, навчаючи її, що саме робити, у послідовному порядку, використовуючи більшу попередньо навчену модель BERT [25].

Дистиляція даних – це істотне зменшення вибору, шлях створення штучних об'єктів (синтетичних даних), які агрегують корисну інформацію, зберігаються в даних і дозволяють налаштувати алгоритми машинного навчання не менш ефективно, ніж на всіх даних [26]. Запропонована модель аналізу настроїв на основі дистиляції знань додатково описана на рисунку 2.6.

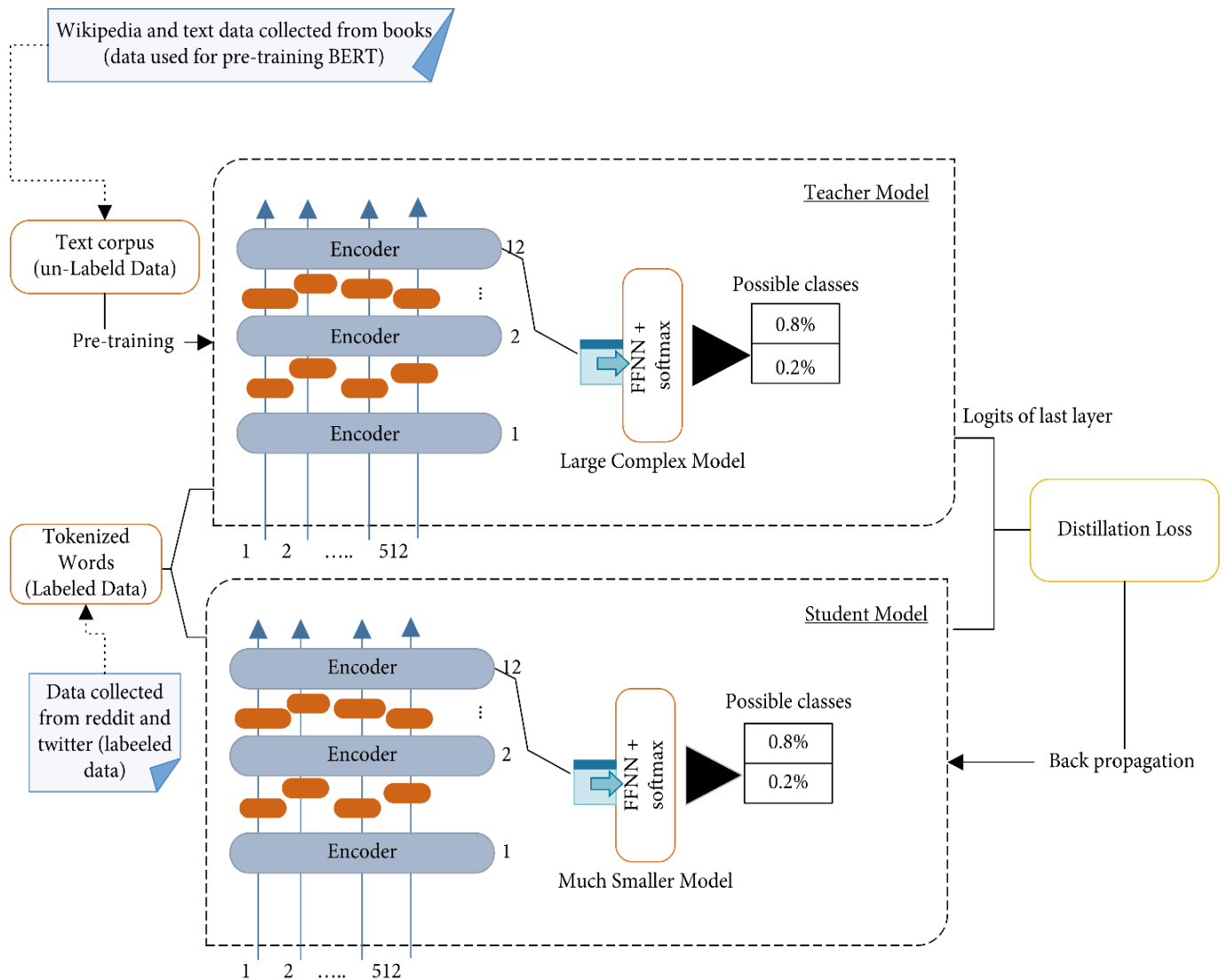


Рисунок 2.6 – Детальна архітектура застосування дистиляції знань із попередньо підготовленого BERT для побудови меншої моделі для визначення проблем, пов'язаних зі здоров'ям

Створений підхід показує, як використовуються попередньо навчені ваги від BERT, які точно налаштовані на завдання виявлення депресії/тривоги, щоб побудувати більш розумну та легшу модель, яка зможе показати кращі показники під час експерименту.

3 ЕКСПЕРИМЕНТАЛЬНІ ДОСЛІДЖЕННЯ

3.1 Порівняння підходу з іншими алгоритмами

У експерименті запропонований гібридний підхід на основі Bi-LSTM та BERT порівнювався з алгоритмами NB, SVM, RF, BI-LSTM, LR, MLP та LSTM для визначення настроїв, пов'язаних із депресією/тривогою, були використані дані власного датасету. Було використано RF із 150 ітераціями, 100 оцінками та SVM з оцінкою гребня навчального параметра та радіальною базисною функцією (ядро = rbf) відповідно. Було використано три відомі методи вбудовування слів, такі як word2vec і gloVe, які застосовуються NB, SVM, RF, BI-LSTM, LR, MLP, LSTM як показано на рисунку 3.1.

Performances of classical machine learning algorithms with our collected data (A, accuracy; P, precision; and R, recall in %).																					
Embedding techniques	Classification models																				
	NB			SVM			RF			BI-LSTM			LR			MLP			LSTM		
	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R
word2vec	67	67	67	73	73	74	73	73	73	82	82	82	70	70	70	68	69	68	74	74	74
gloVe	67	69	65	84	84	83	78	79	78	86	86	86	80	79	78	77	76	77	83	83	82

Рисунок 3.1 – Результати експерименту з використанням вбудовуванням слів

Основною метою тут було вивчити, яка функція є з найкращою ефективністю виявлення депресії та тривоги.

Було реалізовано SVM з оцінкою хребта навчального параметра та радіальною базисною функцією (ядро = rbf), отримавши точність 84%. SVM працює відносно

добре, оскільки існує чітке поділ між мітками, наданими кожному зі зібраних текстових корпусів після попередньої обробки та векторизації.

Продуктивність більшості класичних алгоритмів машинного навчання на рисунку 3.1 знижується під час застосування word2vec, оскільки алгоритм word2vec відкидає невидимі слова під час фази навчання.

Потім було зроблено другий етап тестування, де до всіх алгоритмів NB, SVM, RF, BI-LSTM, LR, MLP та LSTM було використано PCA, щоб подивитись, як зміняться результати тестування.

Було використано аналіз головних компонентів (PCA), щоб зменшити розмірність набору даних, що складається з багатьох пов'язаних змінних, як сильно, так і незначно, зберігаючи при цьому варіації, доступні в текстових даних.

Метод головних компонентів – метод факторного аналізу в статистиці, який використовує ортогональне перетворення множини спостережень з можливо пов'язаними змінними (сутностями, кожна з яких набуває різних числових значень) у множину змінних без лінійної кореляції, які називаються головними компонентами. Метод головних компонент – один з основних способів зменшити розмірність даних, втративши найменшу кількість інформації.

Обчислення головних компонент може бути зведене до обчислення сингулярного розкладу матриці даних або до обчислення власних векторів і власних чисел коваріаційної матриці початкових даних.

Було використано зменшення розмірності, яке передбачає зменшення кількості вхідних змінних, коли ми генерували вектори зі зібраного списку маркерів. Менша кількість вхідних змінних може призвести до простішої прогнозовної моделі, яка має кращу продуктивність під час прогнозування нових даних.

Як показано на рисунку 3.2, точність майже всіх класифікаторів підвищилася, коли було застосовано аналіз головних компонентів як зменшення розмірності.

Performances of classical machine learning algorithms with our collected data after applying a PCA.

Embedding techniques	Classification models																				
	NB			SVM			RF			BI-LSTM			LR			MLP			LSTM		
	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R	A	P	R
word2vec	60	57	59	75	74	74	71	72	68	80	80	80	69	68	70	71	71	68	75	75	75
gloVe	69	67	64	85	85	85	80	80	80	86	86	86	80	79	78	77	76	77	82	82	82

Рисунок 3.2 – Результати експерименту з використанням вбудовуванням слів та PCA

На таблиці 3.1 представлені результати, отримані від запропонованого Bi-LSTM, перегонки знань та інших алгоритмів. Ці моделі класифікації були використані для прогнозування проблем, пов'язаних із психічним здоров'ям, за допомогою трьох типів методів векторизації тексту, таких як word2vec, GloVe та BERT, а результати всіх базових підходів порівнювалися для оцінки ефективності моделей векторизації слів. Bi-LSTM і дистильований BERT досягли вищої точності класифікації 96% і 98% відповідно. Основна причина полягає в тому, що модель студента (дистильована BERT) імітує модель викладача, яку спочатку навчали на загальному текстовому корпусі, такому як Вікіпедія та BookCorpus. Таким чином, дистильований BERT отримує конкурентоспроможну або навіть чудову продуктивність, якщо точно налаштувати на нашу область даних, пов'язаних із депресією та тривогою. Вивчення цієї маленької моделі з більшої попередньо підготовленої моделі в запропонованій нами структурі називається дистилляцією знань. Крім того, BERT перевершує word2vec і GloVe як для завдань прогнозування тривоги, так і для депресії, оскільки BERT здатний розрізняти та охоплювати два різних семантичних значення, створюючи два різних вектори для одного слова в даному текстовому корпусі. Як результат Bi-LSTM з BERT може точно визначити депресію та тривожність із даного великого текстового корпусу.

Таблиця 3.1 – Результати експерименту з гібридним методом

Embedding methods	MLP	SVM	LSTM	NB	RF	LR	BI-LSTM	distilled_BERT
	accuracy	accuracy	accuracy	accuracy	accuracy	accuracy	accuracy	accuracy
BERT	90.32%	90.35%	89%	87%	87.8%	88%	91.34%	97%
GloVe	77%	85%	82%	69%	80%	80%	86%	
Word2vec	71%	75%	75%	60%	71%	69%	80%	

Весь процес відображень за допомогою Juoyter, на рисунку 3.3 зображено створення регресії.

Logistic Regression

```
#build tf-idf model
vec=TfidfVectorizer(preprocessor=preprocessor, ngram_range=(1,3), stop_words=total_stopwords, max_features=10000)
vec_train_data=vec.fit_transform(train_data['tweets'])
vec_test_data=vec.transform(test_data['tweets'])
```

```
# train Logistic Regression
logit=LogisticRegression(penalty='l2')
logit.fit(vec_train_data, train_data['target'])
pred_labels=logit.predict(vec_test_data)

# assess model
f1=f1_score(test_data['target'], pred_labels, average="weighted")
accuracy=accuracy_score(test_data['target'], pred_labels)
confusion=confusion_matrix(test_data['target'], pred_labels)
print("logistic regression f1 score: %.3f" %(f1))
print("logistic regression accuracy score: %.3f" %(accuracy))
print("logistic regression confusion matrix:")
print(confusion)
```

```
#try Keras
from keras.preprocessing.text import one_hot
from keras.preprocessing.text import Tokenizer
from keras.preprocessing.sequence import pad_sequences
from keras.models import Sequential
from keras.layers import Dense
from keras.layers import Flatten
from keras.layers.embeddings import Embedding
from keras.preprocessing import sequence
```

Рисунок 3.3 – Створення регресії

Усі експерименти були зроблені з допомогою Keras.

3.2 Дата-сети

Була розроблена система збору даних за допомогою API Twitter і Reddit (Tweepry і PRAW відповідно) для збору даних про депресію і тривожність, була використана модель афекту для використання ключових слів, що належать до різних емоцій, як показано на рисунку 3.4 [27].

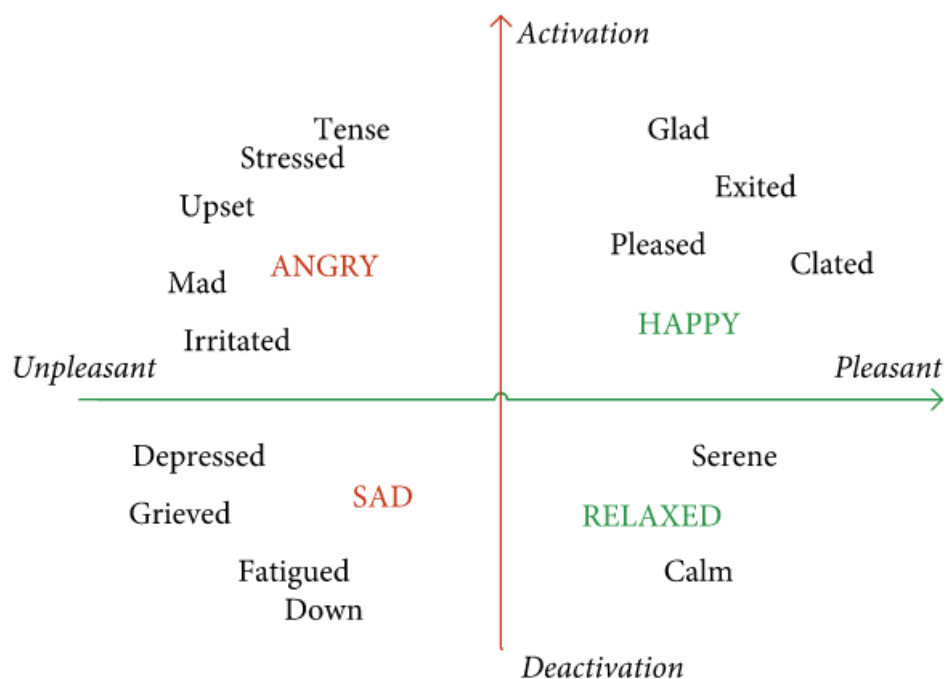


Рисунок 3.4 – Приклад створеної моделі афекту

Tweepry – це пакет, з допомогою якого можна отримати доступ до Twitter API, на рисунку 3.5 зображений скрипт для витягування твітів за допомогою Tweepry.

Було використано підхід на основі ключових слів, який ґрунтується на попередньому визначенні набору термінів для класифікації тексту за такими категоріями емоцій, як радісний, сердитий та сумний, які були використані для збору

найбільш значущих даних про депресію та тривожність із платформ соціальних мереж. Було зібрано 100 000 твітів і 95 000 дописів з Twitter і Reddit відповідно.

```

search_words = 'happiness+love+family'
search_terms = search_words + " -filter:retweets"
date_since = "2010-03-01"

tweets = tw.Cursor(api.search,q=search_terms,lang="en",since=date_since).items(1000)

users_tweets = [[tweet.user.screen_name, tweet.user.location, tweet.text] for tweet in tweets]
tweet_df1 = pd.DataFrame(data=users_tweets,columns=['user', "location", "text"])

tweet_df1.head(20)

```

	user	location	text
0	AmirahStyles_	Las Vegas , NV	I love my family. I love my life. \nHappiness.
1	seejhay_	I couldn't care less	"HoMOSeXUal wILI nOT iNhErIT tHE kINGdom Of gO...
2	AKpuresoul		We r here to shower our unconditional love \nN...
3	_itspjmluv	방탄소년단 <3	you are and always be my happiness and my fami...
4	Smile57383709		Praying the Almighty God showers you with good...
5	VJMiraX	Philippines	worked hard to improve on or correct entirely....
6	sairungs146gma2		True me and true make it happen for my heart ...
7	imlifealtering	Jonesboro, GA	I want a queen who isn't trying to live her li...
8	uptownatadvant	Advant Navis Park Sec-14 Noida	Good times are back again! It's time to get to...
9	123josnow	Cape Breton, Nova Scotia	They know where they do you believe in the hap...
10	RoshBeingRosh	India	@sidharth_shukla @Sid_ShuklaFC Take care @sidh...
11	carmimartin8	Makati City	I wish you love, hope and everlasting joy and ...
12	waifujade	sydney	11:11 to love and be loved, health, happiness,...
13	PhantomUFCfan		No #TheSin5 tomorrow - have been feeling a lit...
14	yeshdayesh	Quezon City, National Capital	Remove three thing in your life\n\nhttps://t.c...
15	miketse	Staten Island, NY	school mates - step up #love #girdad #happine...
16	miketse	Staten Island, NY	stepping up pre-k4 #love #girdad #happiness #...

Рисунок 3.5 – Скрип зі збором даних

PRAW – це пакет Python, який надає простий доступ до API reddit. PRAW є простим у використанні та розроблений відповідно до всіх правил для використання API Reddit. Приклад скрипту, який витягує пости з Reddit зображений на рисунку 3.6. Як видно з рисунка пости витягуються з емодзі та з усіма специфічними символами.

	Title	Id	Upvotes
0	Training Deep NN be like	jrkw6i	507
1	Taking code from GitHub and running on your data	gfpdsc	412
2	Exploring MNIST Latent Space	jkci6f	409
3	Holy crap! Some guy shouted "Machine learning ...	ipfe2l	381
4	I work with models	jwempm	374
5	The merits of having many monitors	g45qwn	325
6	GANs	j5n787	313
7		lqknmb	314
8	The truth  	hu6c0d	315
9	It happens 	lu430g	309

Рисунок 3.6 – Результати скрипта, який витягує пости з Reddit

Результати скрипту відображені на рисунку 3.7.

```
# Create sub-reddit instance
subreddit_name = "deeplearning"
subreddit = reddit.subreddit(subreddit_name)

df = pd.DataFrame() # creating dataframe for displaying scraped data

# creating lists for storing scraped data
titles=[]
scores=[]
ids=[]

# looping over posts and scraping it
for submission in subreddit.top(limit=21):
    titles.append(submission.title)
    scores.append(submission.score) #upvotes
    ids.append(submission.id)
```

Рисунок 3.7 – Скрипт з витягуванням даних з Reddit

Однією з поширених проблем під час підготовки даних є використання незбалансованого набору даних, коли кількість даних, що належать одному класу, значно більша або менша, ніж кількість даних, що належать іншим класам. В результаті алгоритм навчання може ігнорувати ці недопредставлені класи. Тому потрібно було збалансувати кількість позитивних і негативних класів (52% і 48%

Кінець таблиці 3.2

Data sources	Mental health problems	Collected posts	Description
Twitter	Depression	25,000	Contents and posts related to depression
Twitter	Anxiety	15,000	Contents and posts related to anxiety

Як видно з аналізу, було витягнуто більше постів з Reddit, бо там є специфічні каталоги, з яких можна витягнути данні за темами, що значно спрощує роботу. Для витягування даних, потрібно були зареєструватись та створити акаунти розробників та пройти валідацію для користування API Twitter та Reddit.

Текст був векторизований із застосуванням бібліотеки Keras. Кожне слово в тексті асоціювалось із числовим індексом у словнику, і всі зразки даних були представлені як вектори чисел. Кожен вектор доповнювався нулями до постійної довжини, щоб довжина тексту не впливала на остаточну здатність нейронної мережі узагальнювати.

Також була створена хмара слів з постів, які містять слова, які не належать депресії та тривоги.

3.3 Порівняння алгоритму з вже існуючими системами

Було порівняно запропоновану систему з найсучаснішими системами, які були розроблені для аналізу та прогнозування проблем, пов'язаних із психічним здоров'ям. На основі даних соціальних мереж Kim et al. використовував модель word2vec з XGBoost і CNN для класифікації тексту, пов'язаного з депресією і тривогою, і досяг

точності 75,13% для даних, пов'язаних з депресією, і 77,81% для даних, пов'язаних з тривогою [29]. Мікаель та ін. представив методики виявлення депресії з використанням комбінованих ознак (LIWC + LDA + біграм), класифікувавши їх за допомогою SVM та MLP та отримав точність 90% та 91% відповідно [30]. Хатун та ін. застосував лінгвістичний запит і підрахунок слів (LIWC) з лінійним SVM для прогнозування депресії з даних Twitter і отримав точність 82,50% [31]. Зоган та ін. використовували вбудовування слів (word2vec) з рекуррентно-згортковою нейронною мережею (BiGRU-CNN) для виявлення депресії з даних Twitter і досягли точності 85% [32]. У 2019 році Кумар та ін. використав матрицю ознак з класифікацією з використанням RF, NB та градієнтного підвищення та отримав точність 85,09% [33]. Точність запропонованої системи показана на рисунку 3.6. Під час цього експерименту прийнято концепцію дистиляції знань за допомогою BERT і точно налаштовано задачу виявлення депресії та тривоги, отримавши точність 98% для тривожних даних і 97% для даних, пов'язаних з депресією. Крім того, застосовано механізм уваги для збереження довготривалого зв'язку між словами, використовуючи архітектуру трансформатора, і застосували Bi-LSTM для прогнозування тривоги, отримавши точність 96%, як зображено у таблиці 3.3.

Таблиця 3.3 – Опис порівняння вже існуючих систем з гібридним підходом

Author (year)	Data source	Word embedding model	Classifiers	Overall accuracy
Jina Kim (2020)	Reddit	Word2vec	XGBoost and CNN	75.13% depression
Michael (2019)	Reddit	LIWC + LDA + bigram	SVM and CNN	90 % or 91%
Hatoon (2020)	Twitter	LIWC	Linear SVM	82.50%

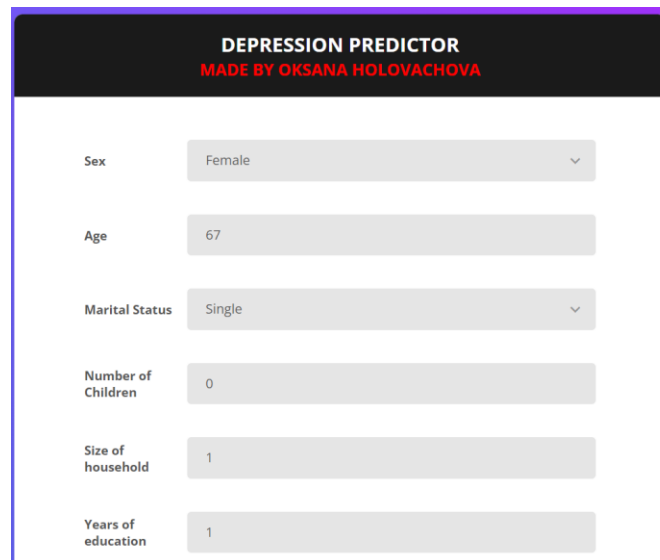
Кінець таблиці 3.3

Tadesse (2020)	Reddit	Word2vec	LSTM-CNN	93.80%
Akshi Kumar(2019)	Twitter		RF, NB and B	85.09%
Humad (2020)	Twitter	Multimodalities + word2vec	BiGRU-CNN	85%
Approach with BERT	Reddit and Twitter	Word2vec, GloVe and Bert	BI-LSTM, BERT + distillation	94%, 96% and 98%

Отримані результати показують, що запропоновані методи можуть допомогти у розробці розумних та ефективних систем для виявлення депресії, тривоги та інших проблем, пов'язаних зі здоров'ям, за допомогою текстового вмісту соціальних мереж, створеного користувачами.

Отриманий гібридний підхід був використаний для створення програмного забезпечення, яке аналізує вхідні параметри та на вихід видає інформацію про стан людини.

Ця програма буде протестована у психологічній спілці Харкова, де психологи зможуть вводити дані, використовувати різні датасети з інформацією про людину, та краще аналізувати стан людей, або ця програма буде протестована у майбутньому для аналізу свого ментального здоров'я, щоб зробити висновок чи треба звернутись до спеціаліста. Тести будуть виконуватись на реальних сеансах, з реальними людьми для більш точної оцінки створеного програмного забезпечення. Вигляд інтерфейсу зображений на рисунку 3.10.



DEPRESSION PREDICTOR
MADE BY OKSANA HOLOVACHOVA

Sex: Female

Age: 67

Marital Status: Single

Number of Children: 0

Size of household: 1

Years of education: 1

Рисунок 3.10 – Інтерфейс програми, які показує чи є ймовірність депресії.

Ця програма написана з використанням Python та створеного гібридного підходу. Приклад коду з відображенням результату зображений на рисунку 3.11. Приймає багато параметрів для аналізу, щоб отримати більш точний результат. Але, програма має бути протестована на відповідність з допомогою психологів, які зможуть оцінити справність програми.

```
@app.route('/result', methods=['POST'])
def result():
    if request.method == 'POST':
        to_predict_list = request.form.to_dict()
        to_predict_list = list(to_predict_list.values())
        to_predict_list = list(map(int, to_predict_list))
        result = ValuePredictor(to_predict_list)
        if int(result) == 1:
            prediction = 'You may be diagnosed with depression'
        else:
            prediction = 'No sign of Depression'
        return render_template("result.html", prediction=prediction)
```

Рисунок 3.11 – Код, який відображує результат діагностики депресії

Усі тести будуть виконані на реальних сеансах, щоб перевірити справність програми.

Параметри, які використовуються для аналізу ментального стану людини:

- стать;
- рік;
- сімейний стан;
- кількість дітей;
- кількість років навчання;
- кількість дітей, яким менше 18;
- вартість товарів тривалого користування;
- витрати;
- земля у власності;
- борги;
- витрати на алкоголь;
- витрати на тютюн;
- витрати на їжу;
- витрати на медичне лікування кожний місяць;
- витрати на медичне лікування дітей кожний місяць;
- витрати на навчання;
- витрати на розваги;
- інші витрати;
- кількість днів без їжі;
- кількість днів без їжі у дітей;
- кількість разів коли їли рибу або м'ясо за останній тиждень;
- чи достатня кількість їжі у холодильнику;
- скільки відкладено на навчання дітей;
- чи засинали голодним.

Функція передбачення зображена на рисунку 3.12.

```

# prediction function
def ValuePredictor(to_predict_list):
    print(to_predict_list)
    to_predict_list = np.array(to_predict_list, dtype=np.float32)
    X_scaler = MinMaxScaler().fit(to_predict_list.reshape(34, 1))
    to_predict_list_scaled = X_scaler.transform(to_predict_list.reshape(34, 1))
    print(to_predict_list_scaled)
    # result = loaded_model.predict_classes([to_predict_list])
    to_predict_list_scaled_reshaped = to_predict_list_scaled.reshape(1, 34)
    print(to_predict_list_scaled_reshaped)
    result = loaded_model.predict_classes([to_predict_list_scaled_reshaped])
    print(result)
    return result[0]

```

Рисунок 3.12 – Функція передбачення

Параметри були створені психологічною спільнотою Харкова, це авторський підхід для визначення ментального стану людини. Також у програмі важливим показником є соціальні мережі, бо там враховується дата сет з твітеру.

ВИСНОВКИ

У ході роботи було проаналізовано проблеми та методи вирішення задачі розпізнавання ментальних розладів за текстом, використовуючи технології глибокого вивчення.

У ході аналізу предметної області було проведено аналіз існуючих рішень та робіт. Більшість з існуючих рішень спеціалізовані під конкретні вхідні дані і не характеризують їх продуктивність на реальних текстових даних.

Загалом всі методи базуються на використанні методів глибокого вивчення та нейронних мереж. Серед методів найбільш часто використовуються підходи пов'язані з рекурентними нейронними мережами та векторним уявленням.

Як результат, у цьому дослідженні була розроблена чітка структура для виявлення проблем з психічним здоров'ям за допомогою методів глибокого навчання, таких як BERT, Bi-LSTM, і переробки знань на основі вмісту соціальних мереж, створеного користувачами. Запропонована структура підвищує точність розумних систем охорони здоров'я для виявлення проблем, пов'язаних із психічним здоров'ям, таких як депресію та тривогу. Цю роботу можна використовувати для створення системи в режимі реального часу для раннього виявлення проблем, пов'язаних із психічним здоров'ям, в основному на основі публікацій користувачів на Reddit і Twitter.

Були проаналізовані різні ключові функції, включаючи збір текстових даних, пов'язаних із психічним здоров'ям, із соціальних мереж за допомогою інтерфейсів прикладного програмування та модуля попередньої обробки, який зосереджується на перетворенні неструктурованих даних у змістовну форму за допомогою різних методів фільтрації даних. Також була застосована техніка позначення тексту на основі ключових слів. Крім того, створена система виявлення проблем із психічним здоров'ям застосовує найновішу техніку вбудовування тексту, засновану на

глибокому навчанні, що забезпечує охоплення семантичного та контекстного значення слів, включених у публікації користувачів.

Запропонована модель представлення тексту на основі BERT перетворює зібрані слова у вектори, які фіксують семантичне значення в зібраному текстовому корпусі, щоб підвищити точність завдання класифікації за допомогою механізму уваги. Також була запропонована дистиляція знань на основі відповідей, яка заснована на нейронній реакції останнього вихідного рівня моделі у BERT для побудови меншої та розумнішої моделі для виявлення та класифікації депресії/тривоги.

Був розроблений власний модуль збору даних, щоб підготувати набір даних із Twitter і Reddit шляхом видобутку найбільш релевантних текстових даних, які можна використовувати для побудови інтелектуальної моделі для розумних систем охорони здоров'я. У кінці, був проведений інтенсивний експеримент, на основі якого запропонована модель BERT-Bi-LSTM покращує точність класифікації настроїв тексту з дописів користувачів. Основна причина, чому модель перевершує інші моделі класифікації машинного навчання, полягає в тому, що вона поєднує сильні сторони моделей BERT і Bi-LSTM, щоб зрозуміти синтаксичну та контекстну інформацію кожного слова.

На основі створеного гібридного підходу була реалізована система для передбачення депресивного стану у людини, яка приймає більше двадцяти параметрів та дата-сети та аналізує чи є депресія у пацієнта. Програмна реалізація створена за допомогою Python та Keras.

Результати, отримані у даній кваліфікаційній роботі можуть бути використані у різних галузях: освіта, медицина, психологія. Створений підхід може стати гарним підґрунтям для створенні швидких та якісних систем для оцінки ознак депресії для подальшого запобігання виявлення депресивних ознак.

У наш час це дуже актуальна тема, тому ця дослідницька робота знайшла відгук у психологічній спільні Харкова. Створене програмне забезпечення буде

використовуватись для допомоги в аналізі стану людей, у психологічній допомозі. Зараз проводяться тести на реальних сеансах, щоб перевірити наскільки влучно працює система з реальними показниками під час сеансів.

Також створений підхід може бути використаний як покращення для вже існуючих систем, які існують у медицині або різних галузях, бо створений гібридний підхід зможе значно підвищити точність виявлення ментальних проблем під час аналізу стану людини.

Цей алгоритм гарно підходить для створення системи для набору персоналу, де дуже важливе ментальне здоров'я кандидата, або система для країн, яка буде допомагати оцінювати весь стан людей та їх думки стосовно якихось подій, та країна зможе врегульовувати ці показники.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. All answers related to depression // Forum, URL: <https://phc.org.ua/news/depresiya-vidpovidi-na-osnovni-zapitannya-pro-ce-zakhvoryuvannya> (дата звернення 19.11.2021).
2. Про схвалення Концепції розвитку охорони психічного здоров'я в Україні на період до 2030 року // Forum, URL: <https://zakon.rada.gov.ua/laws/show/1018-2017-%D1%80#Text> (дата звернення 21.11.2021).
3. Соціальна мережа Reddit // Social media, URL: <https://www.reddit.com/> (дата звернення 21.11.2021).
4. Соціальна мережа Twitter // Social media, URL: <https://twitter.com/?lang=ru> (дата звернення 21.11.2021).
5. Smys, S.; Raj, J.S. Analysis of Deep Learning Techniques for Early Detection of Depression on Social Media Network—A Comparative Study. J. Trends Comput. Sci. Smart Technol. 2021, – C., 24–39.
6. Orabi, A.H.; Buddhitha, P.; Orabi, M.H.; Inkpen, D. Deep learning for depression detection of twitter users. In Proceedings of the Fifth Workshop on Computational Linguistics and Clinical Psychology: From Keyboard to Clinic // New Orleans, LA, USA, 5 June 2018; – C. 88–97
7. Word Embeddings // URL: <https://habr.com/company/ods/blog/329410/> (дата звернення 12.12.2021).
8. Hierarchical Probabilistic Neural Network Language Model / Morin, F., & Bengio, Y // Aistats. – 2005. – №5. – C. 120-144.
9. Aladağ, A.E.; Muderrisoglu, S.; Akbas, N.B.; Zahmacioglu, O.; Bingol, H.O. Detecting suicidal ideation on forums: Proof-of-concept study. J. Med. Internet Res. 2018, – C. 20

10. Kim, J.; Lee, D.; Park, E. Machine Learning for Mental Health in Social Media: Bibliometric Study. *J. Med. Internet Res.* 2021, – C. 23
11. Wongkoblap, A.; Vadillo, M.; Curcin, V. Depression Detection of Twitter Posters using Deep Learning with Anaphora Resolution: Algorithm Development and Validation. *JMIR Ment. Health.* 2021
12. Un Nisa, Q.; Muhammad, R. Towards transfer learning using BERT for early detection of self-harm of social media users. In *Proceedings of the Working Notes of CLEF 2021—Conference and Labs of the Evaluation Forum, Bucharest, Romania, 21–24 September 2021.*
13. Ibitoye, A.O.; Famutimi, R.F.; Olanloye, D.O.; Akioyamen, E. User Centric Social Opinion and Clinical Behavioural Model for Depression Detection. *Int. J. Intell. Inf. Syst.* 2021, 10, 69.
14. Aldarwish, M.M.; Ahmad, H.F. Predicting depression levels using social media posts. In *Proceedings of the 2017 IEEE 13th International Symposium on Autonomous Decentralized System (ISADS), Bangkok, Thailand, 22–24 March 2017; pp. 277–280.*
15. Shen, G.; Jia, J.; Nie, L.; Feng, F.; Zhang, C.; Hu, T.; Chua, T.S.; Zhu, W. Depression Detection via Harvesting Social Media: A Multimodal Dictionary Learning Solution. In *Proceedings of the IJCAI, Melbourne, Australia, 19–25 August 2017; pp. 3838–3844*
16. Xiao, L.; Xue, Y.; Wang, H.; Hu, X.; Gu, D.; Zhu, Y. Exploring fine-grained syntactic information for aspect-based sentiment classification with dual graph neural networks. *Neurocomputing* 2022, 471, 48–59.
17. K. Kowsari, K. J. Meimandi, M. Heidarysafa, S. Mendu, L. Barnes, and D. Brown, “Text classification algorithms: a survey,” *Information*, vol. 10, no. 4, pp. 1–68, 2019.
18. D. Dansana, J. Adhikari, M. Mohapatra, and S. Sahoo, “An approach to analyse and forecast social media data using machine learning and data analysis,” no. 1–5, 03 2020

19. G. Hinton, Improving neural networks by preventing co-adaptation of feature detectors, 2012, p. 18.
20. T. Mikolov, A. Deoras, S. Kombrink, L. Burget, J. Černocký, Empirical evaluation and combination of advanced language modeling techniques, in: Twelfth Annual Conference of the International Speech Communication Association, 2011.
21. G. Hinton, Improving neural networks by preventing co-adaptation of feature detectors, 2012, p. 18.
22. Hinton, G., Vinyals, O., Dean, J.: Distilling the Knowledge in a Neural Network // NIPS (2015)
23. Sucholutsky, I., Schonlau M. «Soft-Label Dataset Distillation and Text Dataset Distillation»
24. Medvedev D., D'yakonov A. «New Properties of the Data Distillation Method When Working With Tabular Data», 2020
25. J. Devlin, M. W. Chang, K. Lee, and K. Toutanova, “BERT: pre-training of deep bidirectional transformers for language understanding,” in Proceedings of the NAACL HLT 2019 - 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, vol. 1, no. Mlm, pp. 4171–4186, North American, 2 June 2019.
26. Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In NAACL-HLT, pages 4171–4186
27. G. Valenza, L. Citi, A. Lanatá, E. P. Scilingo, and R. Barbieri, “Revealing real-time emotional responses: a personalized assessment based on heartbeat dynamics,” Scientific Reports, vol. 4, pp. 1–13, 2014.
28. K. Sailunaz, M. Dhaliwal, J. Rokne, and R. Alhaji, “Emotion detection from text and speech: a survey,” *Social Network Analysis and Mining*, vol. 8, no. 1, pp. 1–8, 2018.

29. J. Kim, J. Lee, E. Park, and J. Han, "A deep learning model for detecting mental illness from user content on social media," *Scientific Reports*, vol. 10, no. 1, pp. 11846–6, 2020.
30. M. M. Tadesse, H. Lin, B. Xu, and L. Yang, "Detection of depression-related posts in reddit social media forum," *IEEE Access*, vol. 7, pp. 44883–44893, 2019.
31. H. S. Alsagri and M. Ykhlef, "Machine learning-based approach for depression detection in twitter using content and activity features," *IEICE - Transactions on Info and Systems*, vol. E103.D, no. 8, pp. 1825–1832, 2020.
32. H. Zogan, X. Wang, S. Jameel, and X. U. Guandong, "Depression detection with multi-modalities using a hybrid deep learning model on social media," pp. 1–23, 2020
33. A. Kumar, A. Sharma, and A. Arora, "Anxious depression prediction in real-time social data ARTICLE INFO," in *Proceedings of the Accepted for publication in the proceeding of International Conference on Advanced Engineering, Science, Management and Technology – 2019 (ICAESMT19)*, pp. 1–7, Uttarakhand University, Uttarakhand, India, 4-Mar-2019.

**ПЕРЕЛІК ДЖЕРЕЛЬ ПОСИЛАННЯ ЗА НАУКОВИМИ НАПРЯМАМИ
КЕРІВНИКА ТА НАУКОВЦІВ КАФЕДРИ ПРОГРАМНОЇ ІНЖЕНЕРІЇ**

34. Afanasieva I.V., et al. Neural network approach for emotional recognition in text / I.V. Afanasieva, D.S. Nazarenko, N.V. Golian // Біоніка інтелекту, Харків: ХНУРЕ, 2019. – 1(92). – С. 9-14.
35. Maksym Bekuzarov, Oleksandr Samantsov, Oksana Mazurova, Mariia Shirokopetleva. Neural Network Architecture Editor With Code Generation. Problem of Infocommunications. Science and Technology (PIC S&T'2020), Kharkiv, Ukraine.- 6- 9 October 2020.
36. Maksym Shopynskyi, Nataliia Golian, Iryna Afanasieva. Long short-term memory model appliance for generating music compositions // 2020 IEEE International Scientific-Practical Conference Problems of Infocommunications, Science and Technology (PIC S&T), 6-9 Oct. 2020, Kharkiv, Ukraine.