

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ
ТЕХНОЛОГІЙ
КАФЕДРА ІНФОРМАЦІЙНИХ СИСТЕМ ТА ТЕХНОЛОГІЙ**

КВАЛІФІКАЦІЙНА РОБОТА

на тему: «Нейромережева модель для діагностики психологічних розладів на
основі аналізу тексту»

на здобуття освітнього ступеня бакалавра
зі спеціальності 126 Інформаційні системи та технології
(код, найменування спеціальності)
освітньо-професійної програми Інформаційні системи та технології
(назва)

*Кваліфікаційна робота містить результати власних досліджень.
Використання ідей, результатів і текстів інших авторів мають посилання
на відповідне джерело*

(підпис)

Катерина ДАВИДЕНКО
Ім'я, ПРІЗВИЩЕ здобувача

Виконала: здобувачка вищої освіти гр. ІСД-41
Катерина ДАВИДЕНКО
Ім'я, ПРІЗВИЩЕ

Керівник: Каміла СТОРЧАК
Ім'я, ПРІЗВИЩЕ

Рецензент:

Ім'я, ПРІЗВИЩЕ

Київ 2025

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
Навчально-науковий інститут інформаційних технологій**

Кафедра Інформаційних систем та технологій

Ступінь вищої освіти Бакалавр

Спеціальність 126 Інформаційні системи та технології

Освітньо-професійна програма Інформаційні системи та технології

ЗАТВЕРДЖУЮ

Завідувач кафедрою ІСТ

_____ Каміла СТОРЧАК

« _____ » _____ 2025 р.

ЗАВДАННЯ НА КВАЛІФІКАЦІЙНУ РОБОТУ

Давиденко Катерині Олексіївні

(прізвище, ім'я, по батькові здобувача)

1. Тема кваліфікаційної роботи: Нейромережева модель для діагностики психологічних розладів на основі аналізу тексту.

керівник кваліфікаційної роботи Каміла СТОРЧАК
(Ім'я, ПРІЗВИЩЕ)

затверджені наказом Державного університету інформаційно-комунікаційних технологій від «24» лютого 2025р. №56

2. Строк подання кваліфікаційної роботи «30» травня 2025 р.

3. Вихідні дані до кваліфікаційної роботи:

1. Науково-технічна література.
 2. Методичні матеріали з обробки природної мови (NLP).
 3. Програмні засоби для розробки та тестування нейромережових моделей.
4. Зміст розрахунково-пояснювальної записки (перелік питань, які потрібно розробити)
1. Огляд існуючих нейромережових моделей.
 2. Підбір моделей для порівняння та реалізації.
 3. Розробка та тестування створених моделей.

5. Перелік графічного матеріалу: *презентація*

6. Дата видачі завдання «24» лютого 2025 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів роботи	Примітка
1	Аналіз наявної науково-технічної літератури	24.02 – 03.03.2025	
2	Обґрунтування актуальності роботи	04.03 – 13.03.2025	
3	Аналіз існуючих аналогічних рішень	14.03 – 17.03.2025	
4	Інструменти та прийоми створення та навчання нейромережових моделей	18.03 – 24.03.2025	
5	Підбір моделей для порівняння та реалізації	24.03 – 01.04.2025	
6	Реалізація обраних моделей	02.04 – 23.04.2025	
7	Тестування моделей	24.04 – 28.04.2025	
8	Оформлення роботи: вступ, висновки, реферат	29.04 – 19.05.2025	
9	Розробка демонстраційних матеріалів	20.05 - 26.05.2025	

Здобувач вищої освіти

(підпис)

Катерина ДАВИДЕНКО
(Ім'я, ПРІЗВИЩЕ)

Керівник

кваліфікаційної роботи

(підпис)

Каміла СТОРЧАК
(Ім'я, ПРІЗВИЩЕ)

РЕФЕРАТ

Текстова частина кваліфікаційної роботи на здобуття освітнього ступеня бакалавра: 56 стор., 28 рис., 1 табл., 23 джерел.

Мета роботи – дослідження та визначення оптимального підходу до створення нейромережевої моделі для діагностики психологічних захворювань на основі аналізу текстових даних.

Об'єкт дослідження – процес автоматизованої діагностики психологічних розладів на основі аналізу тексту, що передбачає збір, обробку та аналіз мовленнєвих особливостей пацієнтів з психіатричними розладами .

Предмет дослідження – методи та технології, що застосовуються для створення нейромережевої моделі для діагностики психологічних розладів на основі аналізу тексту.

Короткий зміст роботи: У роботі визначено та проаналізовано найпоширеніші психологічні захворювання, існуючі можливості та рішення для їх діагностики з застосуванням штучного інтелекту та обробки людської мови. Визначено ключові вимоги до моделі та метрики оцінки ефективності. Реалізовано та порівняно три варіанти різних моделей: два класичних, а саме логістична регресія та машина опорних векторів, а також трансформерна модель DistilBERT. Проведено тестування й оцінку реалізованих моделей із визначенням найефективнішої з них. Також визначено потенційні можливості застосування та бар'єри впровадження.

КЛЮЧОВІ СЛОВА: ШТУЧНИЙ ІНТЕЛЕКТ, ОБРОБКА ПРИРОДНОЇ МОВИ, ДАТА СЕТ, МАШИННЕ НАВЧАННЯ, НЕЙРОМЕРЕЖЕВА МОДЕЛЬ, МЕТРИКИ ЕФЕКТИВНОСТІ, ОЦІНКА, ТЕСТУВАННЯ.

ABSTRACT

Text part of the bachelor level qualification work: 56 pages, 28 pictures, 1 table, 23 sources.

The purpose of the work is to investigate and determine the optimal approach to creating a neural network model for diagnosing psychological diseases based on text data analysis.

The object of research is the process of automated diagnosis of psychological disorders based on text analysis, which involves the collection, processing and analysis of speech features of patients with psychiatric disorders.

The subject of the study is the methods and technologies used to create a neural network model for the diagnosis of psychological disorders based on text analysis.

Summary of the work: The paper identifies and analyzes the most common psychological diseases, existing opportunities and solutions for their diagnosis using artificial intelligence and human language processing. The key requirements for the model and performance evaluation metrics are identified. Three variants of different models are implemented and compared: two classical ones, namely logistic regression and support vector machine, as well as the transformer model DistilBERT. The implemented models were tested and evaluated to determine the most effective one. Potential applications and barriers to implementation are also identified.

KEYWORDS: ARTIFICIAL INTELLIGENCE, NATURAL LANGUAGE PROCESSING, DATA SET, MACHINE LEARNING, NEURAL NETWORK MODEL, PERFORMANCE METRICS, EVALUATION, TESTING.

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ.....	10
ВСТУП	11
РОЗДІЛ 1. ТЕОРЕТИЧНІ ВІДОМОСТІ ПРО ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПСИХІАТРІЇ	14
1.1. Психіатричні розлади: класифікація та загрози для здоров'я	14
1.2. Основні поняття в сфері обробки людської мови (NLP)	19
1.3. Відомості про застосування штучного інтелекту в психіатрії.....	23
РОЗДІЛ 2. АНАЛІЗ СУЧАСНИХ ПІДХОДІВ РОЗРОБКИ НЕЙРОННИХ МЕРЕЖ.....	27
2.1. Огляд сучасних бібліотек та фреймворків для створення нейромережових моделей.....	27
2.2. Аналіз можливих алгоритмів навчання та підготовки даних.....	30
2.3. Визначення основних методів для нейромережових класифікаційних моделей.....	34
РОЗДІЛ 3. ПРАКТИЧНА РЕАЛІЗАЦІЯ НЕЙРОМЕРЕЖЕВОЇ МОДЕЛІ	39
3.1. Підбір датасету	39
3.2. Підбір моделей для класифікації	42
3.2.1. Дослідження моделей SVM, Logistic Regression, DistilBERT	42
3.2.2. Створення моделі DistilBERT	46
3.2.3. Створення моделей SVM та Logistic Regression	48
3.3. Оцінка ефективності створених моделей	51
3.4. Порівняння створених моделей за базовими характеристиками	57

3.5. Можливості покращення створеної моделі	59
РОЗДІЛ 4. ЕКОНОМІЧНО – ПРАКТИЧНЕ ВИКОРИСТАННЯ	61
4.1. Практичне застосування нейромережевої моделі для діагностики психологічних розладів	61
4.2. Аналіз потенційних витрат на впровадження та обслуговування моделі	64
4.3 Ризики та бар'єри для впровадження.....	65
ВИСНОВКИ.....	67
ПЕРЕЛІК ПОСИЛАНЬ	68
ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (ПРЕЗЕНТАЦІЯ)	71

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ

NLP — Natural Language Processing (Обробка природної мови) — напрямок штучного інтелекту, що працює з аналізом і генерацією людської мови.

ШІ — Штучний інтелект — система або програма, яка імітує інтелектуальні функції людини.

ПТСР — Посттравматичний стресовий розлад — психічний розлад, що виникає після переживання травматичних подій.

API — API (Application Programming Interface) — інтерфейс програмування додатків для взаємодії між програмним забезпеченням.

ADHD (РДУГ) — Attention Deficit Hyperactivity Disorder (Розлад дефіциту уваги з гіперактивністю) — неврологічний розлад, що характеризується порушеннями уваги і гіперактивністю.

SVM — Support Vector Machine (Машина опорних векторів) — алгоритм машинного навчання для класифікації та регресії.

GPU — Graphics Processing Unit (Графічний процесор) — апаратний компонент, що прискорює обчислення, зокрема для паралельної обробки даних.

TPU — Tensor Processing Unit — спеціалізований процесор від Google для прискорення операцій машинного навчання.

ML — Machine Learning (Машинне навчання) — підгалузь ШІ, що навчає комп'ютери робити передбачення на основі даних.

Дата сет (Dataset) — це структурований або неструктурований набір даних, який використовується для навчання, тестування або оцінки моделей машинного навчання або для аналізу в різних сферах. У машинному навчанні дата сет містить приклади (записи) з ознаками (атрибутами) і, зазвичай, відповідними мітками (лейблами), якщо це задача з навчання з вчителем.

ВСТУП

Актуальність теми:

Станом на 2022 рік в світі налічувалося 900 мільйонів людей, що страждають від розладів психіки, одними з найпоширеніших розладів були тривожний та депресивний розлади, і такими вони лишаються і на сьогоднішній день світові проблеми такі як військові конфлікти та пандемія у 2024 році спричинили різке зростання випадків психологічних розладів: за попередніми даними, кількість осіб із тривожними станами збільшилася на 26%, а з депресивними розладами — на 28%. Незважаючи на існування якісних та ефективних методів профілактики та лікування, більшість людей із психічними розладами не отримують належної допомоги або й не підозрюють про свій стан [1].

Саме за допомогою технологій штучного інтелекту людство може забезпечувати людей навіть дистанційною діагностикою хвороб і надавати якісну допомогу та підтримку, що не потребуватиме великих витрат.

Метою роботи є дослідження та визначення оптимального підходу до створення нейромережевої моделі для діагностики психологічних захворювань на основі аналізу текстових даних.

Завданнями дослідження є:

1. аналіз літературних джерел;
2. підготовка та підбір даних;
3. проектування нейромережевих моделей;
4. порівняльний аналіз створених моделей;
5. тестування моделей;
6. оцінка моделей.

Об'єктом дослідження є процес автоматизованої діагностики психологічних розладів на основі аналізу тексту, що передбачає збір, обробку та аналіз мовленнєвих особливостей пацієнтів з психіатричними розладами .

Предметом дослідження є методи та технології, що застосовуються для створення нейромережевої моделі для діагностики психологічних розладів на основі аналізу тексту.

Методи дослідження - збір теоретичної інформації щодо нейромережевих моделей та методів їх створення, аналіз існуючих методів застосування штучного інтелекту в психіатрії, підготовка даних для навчання моделей, створення нейромережевих моделей, порівняльний аналіз створених моделей, тестування моделей, оцінка ефективності створеної моделі, визначення потенційних покращень та розвитку моделі.

Наукова новизна полягає в порівнянні декількох технологій нейромережевих моделей, а саме моделі: логістична регресія, система опорних векторів та DistilBERT, та їх навчання на основі конкретизованої вибірки даних, а саме текстових даних з маркерами психологічних розладів. Визначення критичних відмінностей у чутливості, специфічності, апаратних потребах та часових термінах створених моделей. Також було запропоновано, розроблено та апробовано найбільш адаптовану та пристосовану під конкретну задачу діагностики модель, оптимізовану під задачі психолінгвістичної діагностики.

Апробація результатів та публікації

Апробація роботи проходила у формі участі у всеукраїнських науково-технічних конференціях «Сучасний стан та перспективи розвитку IoT» та «Технологічні горизонти: дослідження та застосування інформаційних технологій для технологічного прогресу України і світу». Також була пройдена апробація в журналі «Телекомунікаційні та інформаційні технології».

Теоретична значущість даної роботи полягає в тому що проведене порівняння та аналіз реалізованих моделей надає значущі результати при підборі класифікаційних моделей під конкретні задачі, такі як діагностика психологічних розладів на основі тексту.

Практична значущість полягає в тому що отримані під час дослідження результати щодо часових рамок, специфіки та апаратних ресурсів надають рекомендації щодо вибору моделі під конкретні потреби. Також запропонована нейромережева модель, що може бути основою для впровадження автоматизованої діагностики психологічних розладів до застосування в клінічній практиці.

РОЗДІЛ 1. ТЕОРЕТИЧНІ ВІДОМОСТІ ПРО ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПСИХІАТРІЇ

1.1. Психіатричні розлади: класифікація та загрози для здоров'я

Психічні розлади є надскладною медичною проблемою, що не залишає осторонь людей по всьому світу. Деякі розлади можуть бути не помітними, а деякі мають більш різкі прояви і можуть впливати не тільки на носія розладу, а ще й на оточуючих людей. Зважаючи на особливості психічних розладів, люди майже весь час перебувають в стані стресу, що різко негативно може впливати на якість життя.

Всесвітня Організація Охорони Здоров'я виділила перелік основних критеріїв психологічної врівноваженості та здоров'я:

- усвідомленість власної особистості та цілісності «я»;
- стабільність почуттів та реакцій в повторюваних та подібних ситуаціях;
- розвинене критичне мислення по відношенню до себе та власних вчинків;
- адекватність психологічних відгуків на зовнішні подразники за інтенсивністю й частотою;
- навичка корегування своєї поведінки відповідно до норм встановлених соціумом;
- навички до планування власної діяльності;
- адаптування власної поведінки відповідно до контексту ситуації [2].

Психічні хвороби — це порушення функціонування нервової системи, що впливають на емоційний стан, когнітивні процеси та поведінку людини. До них належать депресія, тривожні розлади, шизофренія та інші психіатричні стани.

Соматичні ж захворювання охоплюють фізичні порушення організму, такі як серцево-судинні хвороби, діабет, артрит тощо.

Проте межа між цими двома категоріями часто розмивається. Існують психосоматичні розлади, при яких психологічні фактори спричиняють або загострюють фізичні хвороби. Наприклад, стрес і депресія можуть погіршувати перебіг серцево-судинних захворювань, а хронічні соматичні недуги можуть призводити до розвитку тривожних і депресивних станів. Дослідження ВООЗ показали, що депресія у пацієнтів із хронічними соматичними хворобами зустрічається у 9,3–23% випадків, а її вплив на загальний стан здоров'я може бути навіть значнішим, ніж вплив самої соматичної патології.

На даний момент є різні класифікації та категорії розладів психіки. Нижче наведено перелік основних психологічних розладів (рис. 1.1), їх характеристик та проявів.

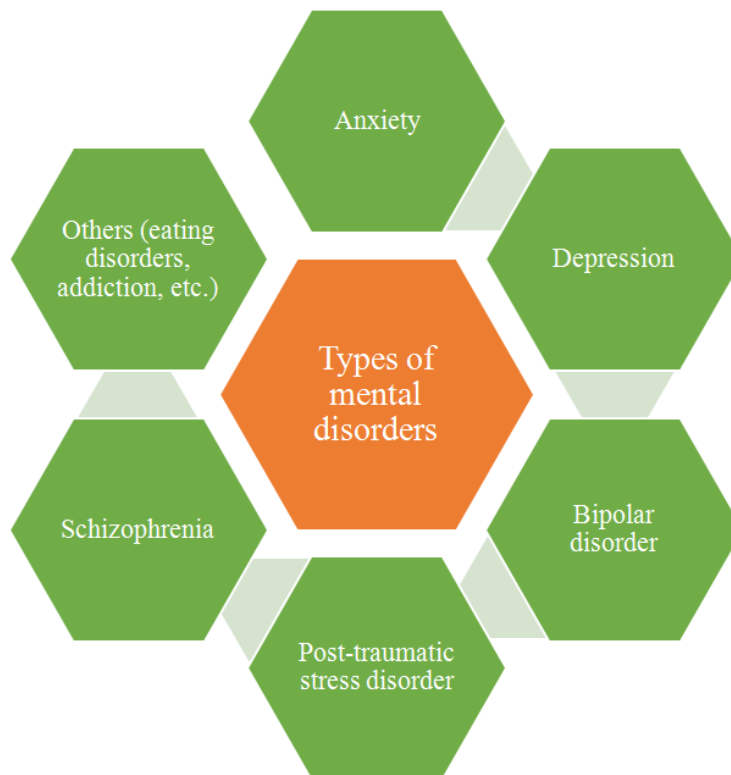


Рис. 1.1 Види розладів психіки [3]

Тривожний розлад

Тривожний розлад характеризується станом постійного занепокоєння, страху як перед ситуаціями, що потребують соціалізації так і перед виконанням завдань, що здаються за складними. Також цей розлад характеризується страхом перед майбутнім та панічними атаками. Симптоматика розладів даної категорії проявляється не тільки у вигляді симптомів з боку психіки, а й фізичних таких як: озноб, нудота, тахікардія та інші. Ця симптоматика досить складно поєднуються з нормальним функціонуванням людини та сильно ускладнює життя. До різновидів тривожних розладів відносять такі стани як:

- генералізований тривожний розлад або ГТР;
- соціофобія або соціальний тривожний розлад;
- obsесивно-компульсивний розлад (ОКР) та інші.

Депресія

Депресивний розлад, який зазвичай називають депресією, є одним із найпоширеніших психічних станів. Він проявляється стійким пригніченням настрою та втратою здатності отримувати задоволення або зацікавленості в раніше приємних заняттях протягом тривалого часу.

За різними даними, від депресії страждає близько 3,8 % населення: серед дорослих цей показник сягає 5 % (4 % у чоловіків та 6 % у жінок), а серед людей старше 60 років — близько 5,7 %. Загалом у світі приблизно 280 млн осіб живуть із діагнозом депресії. Жінки схильні до цього розладу майже вдвічі частіше, ніж чоловіки. Крім того, понад 10 % вагітних та недавно народивших жінок стикаються з депресивними станами. Кожного року самогубство стає причиною смерті для понад 700 000 осіб. Самогубство є четвертою провідною причиною смерті у віці 15-29 років [4,5].

Наявність певних факторів ризику підвищує ймовірність розвитку депресії. Нейродегенеративні захворювання, такі як хвороба Альцгеймера та хвороба

Паркінсона, інсульт, розсіяний склероз, судомні розлади, онкологічні захворювання та хронічні болі можуть сприяти розвитку депресивних розладів та станів.

Біполярний розлад

Біполярний розлад – це розлад, що супроводжується різкими чергуваннями депресивних та маніакальних епізодів. Депресивний епізод зазвичай характеризується апатією, сумом або роздратованістю, часто зникає бажання займатися навіть найпростішими буденними справами. Маніакальні епізоди – це стан повністю протилежний до депресивного, різка зміна емоцій, ейфорія, підвищена швидкість думок, безрозсудна поведінка.

Протягом маніакальної фази біполярного розладу людина може коїти ризиковані вчинки, які несуть загрозу її здоров'ю, репутації та фінансам: наприклад, бездумно витрачати великі кошти або братися за азартні ігри, а також здійснювати небезпечні маневри за кермом. Іноді при цьому з'являються психотичні прояви — у вигляді нав'язливих ідей або зорових/слухових галюцинацій — через що біполярний розлад інколи плутають із шизофренією чи шизоафективним розладом. Крім того, ризик суїциду серед таких пацієнтів значно вищий, тому вони потребують ретельного медичного супроводу та регулярного лікування.

Посттравматичний стресовий розлад (ПТСР)

Протягом останніх двох років у системі охорони здоров'я відзначається стрімке збільшення кількості людей з посттравматичним стресовим розладом. При порівнянні між 2021 та 2023 роками кількість людей з ПТСР зросла майже вчетверо. Більше того, лише за січень–лютий 2024 року було зареєстровано стільки ж нових діагнозів ПТСР, скільки за весь 2021 рік.

Враховуючи статистичні дані кількість пацієнтів з ПТСР за роками була в 2021 році - 3 167 пацієнтів, в 2022 - 7 051 пацієнтів, в 2023 - 12 494 пацієнтів і в 2024 - 3 292 пацієнтів і це тільки за перші два місяці (станом на 06.03.2024) [6].

Посттравматичний стресовий розлад (ПТСР) - є психічним розладом, що виникає як наслідок впливу надзвичайно стресової або травматичної події, яку людина пережила особисто або стала свідком такої події. Симптоми, пов'язані з цим станом, можуть включати повторне переживання події через спогади, нічні кошмари та відчуття тривоги; а також неконтрольовані думки про подію.

Шизофренія

Шизофренія – це розлад, що впливає на думки, почуття та поведінку людей. Пацієнти з цим розладом зазвичай не мають можливості розрізняти реальні події від галюцинацій, і це може призводити до серйозних травм та самогубств.. Зазвичай шизофренію виявляють у віці від 16 до 30 років, після перших проявів психозу. Попередниками першого випадку психозу є різка зміна мислення, настрою та соціального функціонування. Не вчасна або не коректна діагностика часто призводить до затримок початку лікування, що сильно впливає на ефективність та результативність лікування.

Розлади харчової поведінки

Порушення харчової поведінки проявляються як стійкі відхилення у способі споживання їжі, супроводжувані нав'язливими думками та тривогою щодо їжі, своєї зовнішності та тіла. Вони можуть серйозно позначатися на фізичному здоров'ї, емоційному стані та соціальній адаптації. До основних форм таких розладів належать нервова анорексія, нервова булімія, розлади з уникненням обмежень у харчуванні, інші специфічні варіанти порушень харчової поведінки, а також синдром «пика» і синдром «румінації». Часто ці стани призводять до самопошкоджень чи спроб самогубства.

Нервова анорексія часто зустрічається в підлітковому віці, але і в дорослому це теж не є рідкістю. Люди, що стикаються з нервовою булімією мають значно підвищений ризик вживання психоактивних речовин та ускладнень зі здоров'ям.

1.2. Основні поняття в сфері обробки людської мови (NLP)

Обробка природної мови або NLP — це напрямок штучного інтелекту, присвячений створенню систем, які взаємодіють із людиною за допомогою звичайної мови. Головне завдання NLP полягає в тому, щоб надати комп'ютерам здатність сприймати, тлумачити та адекватно відповідати на тексти й усні висловлювання.

Ця технологія слугує для полегшення безперешкодної взаємодії між людським спілкуванням і машинним розумінням, тим самим дозволяючи більш інтуїтивно зрозумілу та ефективну взаємодію з цифровими системами.

NLP використовує різноманітні методи, які допомагають комп'ютерам розуміти природну мову подібно до того, як це робить людина. Обробка природної мови складається з двох основних етапів: попередня обробка даних і розробка алгоритму. Етап попередньої обробки даних називають етапом очистки текстової інформації для того, щоб машина мала змогу її проаналізувати.

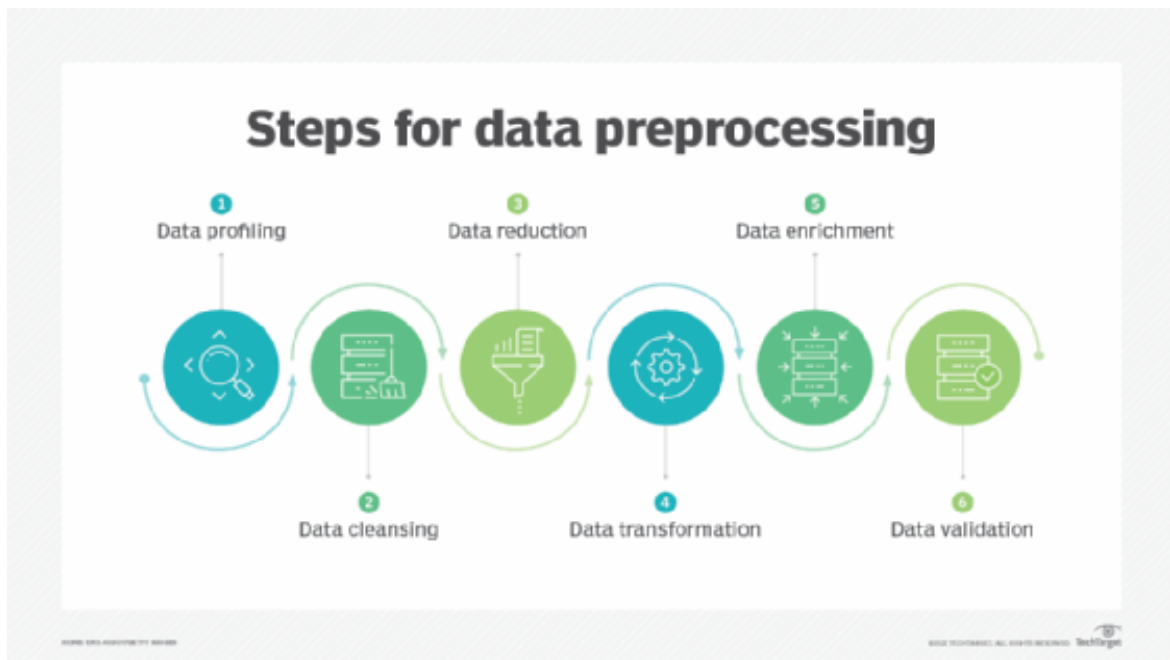


Рис. 1.2 Кроки попередньої обробки даних [7]

Першим етапом попередньої обробки даних є очищення даних. Очищення даних - це процес виявлення пошкоджених даних і неточних записів у наборі записів або таблиці бази даних. Основне призначення етапу очищення полягає у виявленні неповних, неточних, непослідовних і нерелевантних даних та застосуванні методів для модифікації або видалення цих непотрібних даних.

Другим етапом є інтеграція даних, яка зосереджена на об'єднанні даних, що знаходяться в різних джерелах, і представленні їх в єдиному вигляді. Дані з різним представленням об'єднуються разом і вирішуються будь-які конфлікти, що виникають в результаті цього. Цей процес стає життєво важливим у ряді наукових і комерційних застосувань. Зі збільшенням обсягу та експоненціальним зростанням даних, їх інтеграція стає ще більш значущою.

Трансформація даних є третім етапом NLP і відіграє ключову роль у перетворенні необроблених даних у зрозумілу для машини форму. Цей етап, в свою чергу, складається з нормалізації, агрегації та узагальнення даних. Основний фокус нормалізації даних на тому щоб впорядкувати стовпці та таблиці

бази даних таким чином, щоб надмірність або була мінімальною, або була повністю відсутня. Це сильно зменшує час і складність обробки даних. Агрегація даних займається коротким резюмуванням даних. Процес узагальнення даних також відомий як згортання даних. Він допомагає узагальнити дані і створює послідовні шари підсумків в оціночній базі даних.

Скорочення даних - це процес перетворення цифрової інформації в упорядковану та спрощену форму. Ці дані зазвичай отримують емпіричним та експериментальним шляхом. Він передбачає скорочення великих обсягів даних до менших і змістовних фрагментів.

Останнім етапом є дискретизація даних. Дискретизація даних є важливою концепцією, коли ви маєте велику кількість числових даних, але хочете класифікувати їх лише на основі номінальних значень. У цьому випадку безперервні дані розбиваються на дискретні форми, а значення цих дискретних наборів називаються номінальним значенням. По суті, це процес перетворення атрибутів безперервних даних у кінцевий набір інтервалів з мінімальною втратою інформації.

Іноді нормалізацію та дискретизацію даних розглядають як загальні частини більш об'ємного етапу перетворення, а зменшення розмірності може охоплювати як дискретизацію, так і інші методи зменшення. Хоча точна послідовність і можливості цих етапів можуть змінюватися залежно від завдання і використовуваних інструментів, наведена вище п'ятиетапна структура залишається стандартним підходом в машинному навчанні та аналізі даних.

Обробка природної мови передбачає п'ятиетапний процес, який комп'ютери використовують для аналізу, класифікації та інтерпретації як усного, так і письмового тексту. Ці етапи спираються на методи навчання на основі глибоких нейронних мереж, що мають на меті відтворити здатність мозку здобувати знання і точно обробляти інформацію (рис. 1.3).

- а) Лексичний аналіз – ідентифікує структуру слів і фраз у мові.
- б) Синтаксичний аналіз – виконує аналіз слів у реченні, виділяє граматику та розташування слів у спосіб, який показує зв'язок між словами.
- в) Семантичний аналіз – витягує з тексту або словника точне значення.
- г) Інтеграція дискурсу – значення будь-якого речення залежить від речення перед ним.
- д) Прагматичний аналіз – під час цього кроку сказане переосмислюється відповідно до того, що воно насправді означало.

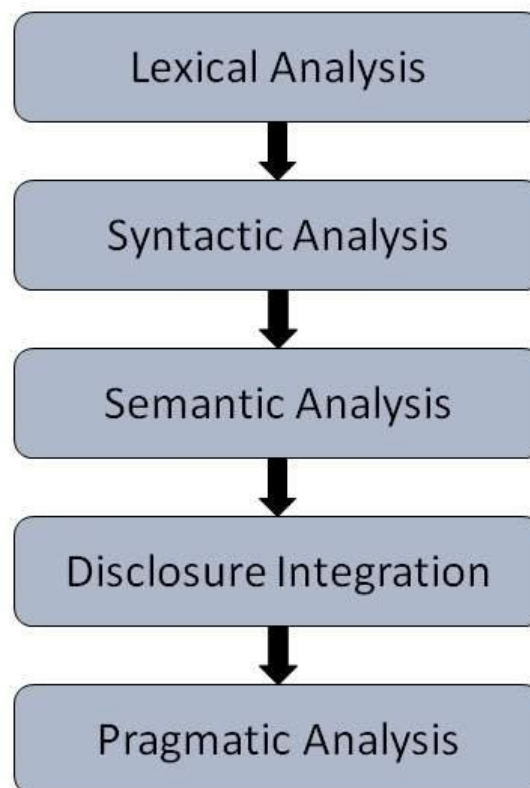


Рис. 1.3 Кроки в NLP [8]

NLP вже є невідомною частиною цифрового життя, сприяючи значному прогресу в різних галузях. Одним з найвідоміших способів використання NLP є

перекладач Google Translate. Постійне підвищення точності перекладу підкреслює значний вплив NLP у сфері подолання мовних кордонів.

Іншим важливим застосуванням NLP є розпізнавання розмовної мови шляхом запису голосу. Такі віртуальні помічники, як Siri від Apple, Alexa від Amazon і Google Assistant, використовують складні алгоритми NLP, щоб розуміти голосові команди і реагувати на них.

Чат-боти – це ще один напрямок, що швидко розвивається. Застосування обробки людської мови за допомогою використання чат-ботів особливо популярне у сфері обслуговування клієнтів.

Ще одним потужним застосуванням NLP є аналіз настроїв, який широко використовується в маркетингу та моніторингу соціальних мереж.

Ефективність пошуку інформації підвищується за допомогою NLP шляхом розробки та впровадження вдосконалених пошукових систем і рекомендаційних систем. Такі компанії, як Netflix і Amazon, використовують NLP для аналізу вподобань і поведінки користувачів, надаючи персоналізовані рекомендації щодо контенту, які підвищують залученість користувачів і покращують користувацький досвід.

1.3. Відомості про застосування штучного інтелекту в психіатрії

За останні роки можливо відслідкувати суттєвий прогрес у галузі штучного інтелекту (ШІ), а особливо в його застосуванні у сфері психіатрії, що викликає все більший інтерес у суспільства та розробників. Використання ШІ в цій галузі має великі перспективи для покращення результатів лікування пацієнтів і надання цінної інформації медичним працівникам. Однак необхідно визнати виклики та етичні міркування, які супроводжують інтеграцію ШІ в цю сферу.

Штучний інтелект наразі використовується для прискорення виявлення захворювань, покращення розуміння прогресування хвороби, оптимізації дозування ліків та лікування, а також для відкриття інноваційних методів лікування. Оскільки ШІ є галуззю науки про дані, однією з його найбільших переваг є здатність швидко аналізувати закономірності та ефективно обробляти великі масиви даних. Галузі медицини, що досягли найбільшого успіху у використанні ШІ для розпізнавання образів, включають офтальмологію, онкологію та радіологію. У цих галузях алгоритми ШІ продемонстрували продуктивність на рівні або навіть вище, ніж досвідчені клініцисти, у виявленні відхилень або тонкощів, які можуть бути непомітними для людського ока.

Наприклад, згорткові нейронні мережі були використані для автоматизованого виявлення пухлин печінки на КТ-сканах і показали високу точність і достовірність результатів - 93% і 67% відповідно. Машинне навчання використовувалося для сегментації щільності молочних залоз на мамографічних знімках і показало результати, які за точністю не поступаються ручним [9,10].

Хоча штучний інтелект набув широкого застосування в медицині, особливо в галузі фізичного здоров'я, його роль у сфері психічного здоров'я залишається значно обмеженішою. Охорона психічного здоров'я значною мірою покладається на міжособистісні навички, такі як встановлення контакту, формування значущих стосунків з пацієнтами та уважне спостереження за їхніми емоціями та поведінкою. Крім того, клінічні дані в галузі психічного здоров'я часто є більш суб'єктивними і якісними, оскільки складаються з розповідей пацієнтів і письмових оцінок. Однак, незважаючи на ці виклики, дослідники все більше і більше вводять ШІ у використання в психіатрії.

Штучний інтелект використовується в психіатрії вже кілька десятиліть, а його коріння сягає ELIZA (рис. 1.4), комп'ютерної програми 1966 року, розробленої для імітації роботи психотерапевта. Відтоді штучний інтелект досяг

значних успіхів у застосуванні до психічного здоров'я, зокрема у виявленні та прогнозуванні психологічних станів.

```

=====
EEEEEEEE L      IIIIII ZZZZZZZ AAA
E         L      I      Z      A  A
E         L      I      Z      A  A
EEEEEE    L      I      Z      A  A
E         L      I      Z      AAAAAA
E         L      I      Z      A  A
EEEEEEEE LLLLLLL IIIIII ZZZZZZ A  A
=====
ELIZA > (PLEASE TYPE IN ALL CAPS) WHAT'S YOUR NAME DEAR ?
AMIT
ELIZA > HI, I'M ELIZA. WHAT DO YOU WANT TO TALK ABOUT ?
AMIT > CAN I TALK ABOUT MY PROBLEM ?
ELIZA > SURE... !
AMIT > DUE TO ON GOING EXAMS IN COLLEGE, I HAVE A LOT OF STRESS.
ELIZA > PLEASE RELAX A LITTLE. SLEEP WELL.
AMIT > THANKS FOR YOUR ADVICE.
ELIZA > NO MENTION.
AMIT > BYE
ELIZA > BYE AND KEEP IN TOUCH...
=====

```

Рис. 1.4 Перший в світі чат-бот ELIZA [11]

Нещодавні дослідження продемонстрували неабияку точність ШІ у виявленні осіб із суїцидальними думками. Алгоритми машинного навчання, що аналізують результати фМРТ, досягли 91% успішності у розрізненні людей з суїцидальними думками та без них і 94% точності у виявленні тих, хто намагався покінчити життя самогубством. ШІ перевершив фахівців у галузі психічного здоров'я у розрізненні справжніх передсмертних записок від сфабрикованих, правильно класифікувавши їх у 78% випадків порівняно з 63% точності експертів.

Прогностичні моделі також застосовуються до електронних медичних записів: алгоритми ШІ в Університеті Вандербільта визначають ризик самогубства з точністю до 84%. В іншому дослідженні, проведеному в дитячій лікарні Цинциннаті, лінгвістичний і акустичний аналіз використовувався для

класифікації пацієнтів на суїцидальні, психічно хворі, але не суїцидальні, або контрольні групи з точністю 85% [12,13].

Аналіз мови за допомогою штучного інтелекту також виявився ефективним у прогнозуванні психозу. Дослідження, що аналізувало мовленнєві патерни людей з групи ризику, показало, що зниження семантичної зв'язності та синтаксичної складності може передбачити психоз з точністю 82%. В іншому дослідженні, присвяченому аналізу молоді з групи ризику, було досягнуто рекордно високої точності - 100% - у прогнозуванні розвитку психозу [14].

Ці результати підкреслюють зростаючий потенціал ШІ в психологічних та психіатричних дослідженнях, пропонуючи нові інструменти для раннього виявлення, оцінки ризиків і втручання в охорону психічного здоров'я.

РОЗДІЛ 2. АНАЛІЗ СУЧАСНИХ ПІДХОДІВ РОЗРОБКИ НЕЙРОННИХ МЕРЕЖ

2.1. Огляд сучасних бібліотек та фреймворків для створення нейромережових моделей

Більшість нейронних мереж створюються за допомогою Python, що є однією з найвідоміших мов програмування.

Порівнюючи з іншими мовами програмування, Python користується високим попитом на ринку. Він має постійно активну та прогресуючу спільноту користувачів, яка постійно розширюється та поповнюється. Ця мова є універсальною і може застосовуватися в різних сферах і для різних цілей, від Інтернету речей до аналізу даних та веб-розробки, від науки про дані та машинного навчання до дизайну ігор.

Нейронні мережі вже давно не створюються з нуля, а зазвичай ґрунтуються на певних шаблонах, бібліотеках та фреймворках, які допомагають швидко та ефективно розвивати та масштабувати нейромережі. Нижче наведено перелік основних бібліотек, що застосовуються для створення нейромережових моделей (рис. 1.5):

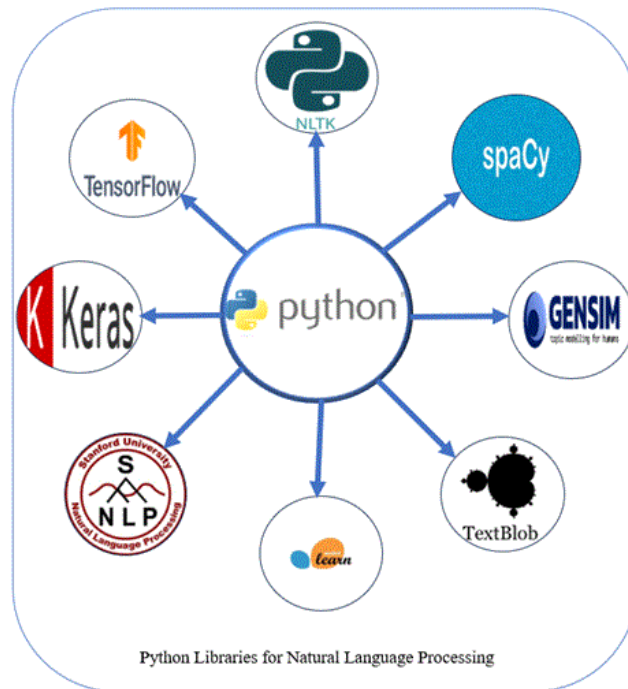


Рис. 2.1 Бібліотеки Python для обробки людської мови [15]

– TensorFlow;

TensorFlow був вперше розроблений у 2015 році командою Google Brain, дослідницької групи з машинного навчання в Google. Він був розроблений для полегшення досліджень у галузі побудови моделей машинного навчання та нейронних мереж. Таким чином, TensorFlow пропонує просту побудову моделей і надійне середовище для розгортання моделей машинного навчання. Він також може похвалитися тим, що має платформу, якою можуть користуватися як початківці, так і експерти, і може використовуватися на платформах Windows, Linux, Android і macOS. Серед компаній, які використовують TensorFlow, є Google, DeepMind, AirBnB, Coca-Cola та Intel.

– CNTK;

CNTK також відома як Microsoft Cognitive Toolkit. Вперше випущена в 2016 році і зараз застаріла, CNTK - це бібліотека з відкритим вихідним кодом для навчання нейронних мереж для глибокого навчання. Він дозволяє користувачам

створювати і комбінувати широко використовувані нейронні мережі, такі як згорткові нейронні мережі (CNN), а також рекурентні нейронні мережі (RNN). На платформах Linux, Windows та macOS підтримується CNTK.

– PyTorch;

PyTorch, випущений у 2016 році, був розроблений дослідницькою лабораторією штучного інтелекту Facebook. PyTorch був побудований на основі Torch - бібліотеки машинного навчання та наукових обчислень. За словами розробників PyTorch, це фреймворк машинного навчання, що прискорює процес створення дослідницьких прототипів до впровадження у виробництво.

PyTorch має розподілене навчання нейромережевих моделей, хмарну підтримку, надійну екосистему та підтримується на платформах Windows, Linux та macOS. Його використовують такі організації, як Стенфордський університет та Udacity.

– Keras;

Keras був розроблений Франсуа Шолле, інженером Google у 2015 році. Він надає зручний інтерфейс для побудови нейронних мереж на мові Python. Keras містить такі модулі, як функції активації та шари для реалізації нейронних мереж за кілька кроків для швидкого експериментування.

Крім того, бібліотека може працювати поверх інших бібліотек, таких як TensorFlow та CNTK. Завдяки простому, але потужному інтерфейсу, її використовують такі організації, як NASA та CERN.

Також виділяють окремі бібліотеки для обробки природної мови:

Natural Language Toolkit (NLTK) одна з найбільших бібліотек Python для виконання різноманітних завдань з обробки людської мови. Від елементарних завдань, таких як попередня обробка тексту, до таких завдань як векторизоване представлення тексту - API NLTK охоплює всі можливі завдання.

sprasy - це бібліотека для розширеної NLP на Python та Cython. Вона заснована на передових дослідженнях і з самого початку була розроблена для використання в реальних продуктах. sprasy постачається з попередньо навченими конвеєрами і наразі підтримує токенизацію та навчання для 70+ мов. Він має найсучасніші швидкісні та нейромережеві моделі для тегування, синтаксичного аналізу, розпізнавання іменованих об'єктів, класифікації тексту тощо, багатозадачне навчання з попередньо навченими трансформаторами, такими як BERT, а також готову до виробництва систему навчання та просте пакування, розгортання та управління робочим процесом.

Hugging Face це онлайн-спільнота, де люди можуть об'єднуватися, досліджувати та працювати разом над проектами машинного навчання. Hugging Face Hub - це платформа з понад 350000 моделей, 75000 наборів даних та 150000 демо-додатків, безкоштовних і відкритих для всіх.

Gensim - це бібліотека з відкритим вихідним кодом на Python, написана Радимом Рехуреком, яка використовується для неконтрольованого моделювання тем і обробки природної мови. Вона призначена для вилучення семантичних тем з документів. Вона може обробляти великі колекції текстів. Це відрізняє його від інших програмних пакетів машинного навчання, які орієнтовані на обробку пам'яті. Gensim також надає ефективні багатоядерні реалізації різних алгоритмів для збільшення швидкості обробки. Він надає більш зручні засоби для обробки тексту, ніж інші пакети, такі як Scikit-learn, R тощо.

2.2. Аналіз можливих алгоритмів навчання та підготовки даних

Дані є невід'ємною частиною проектів, що пов'язані з штучним інтелектом. Без гарно підібраних та оброблених даних проект не буде виконувати свої функції та працювати. Створення дата сету з нуля – є не простим завданням.

Навчання нейромережевої моделі – це процес навчання для виконання прогнозів або прийняття рішень, піддаючи їх впливу великих обсягів даних. Модель вивчає закономірності та кореляції на основі цих даних, що дозволяє їй узагальнювати та виконувати завдання, з якими вона раніше не стикалася.

Фундаментом навчання нейромережевої моделі є навчальні дані, що є великим обсягом матеріалу, на якому навчаються моделі штучного інтелекту. Дані можуть бути як текстові, зображення, аудіо або показання датчиків, дані залежать від характеру завдання штучного інтелекту.

Є чотири різні алгоритми навчання моделей:

- навчання з учителем (Supervised learning) - навчальні дані є важливими і повинні супроводжуватися мітками. Ці мітки дозволяють моделі зрозуміти взаємозв'язок між певними атрибутами та відповідними їм мітками;
- навчання без учителя (Unsupervised learning) – немає потреби в мітках у навчальному наборі даних. При неконтрольованому навчанні модель машинного навчання шукає притаманні атрибутам закономірності або структури, щоб сформулювати узагальнені групування або прогнози;
- напівконтрольоване навчання (Semi-supervised learning) – це метод, що використовує гібридний навчальний набір даних, що містить суміш немаркованих і маркованих ознак, для вирішення унікальних завдань, пов'язаних з напівкеруваним навчанням;
- ще один метод, який можна застосувати до цих трьох моделей, - це підкріплене навчання (reinforced learning). Підкріплене навчання означає надання заохочень або покарань за результати, які дає модель ШІ, таким чином навчаючи її в повторюваному процесі.

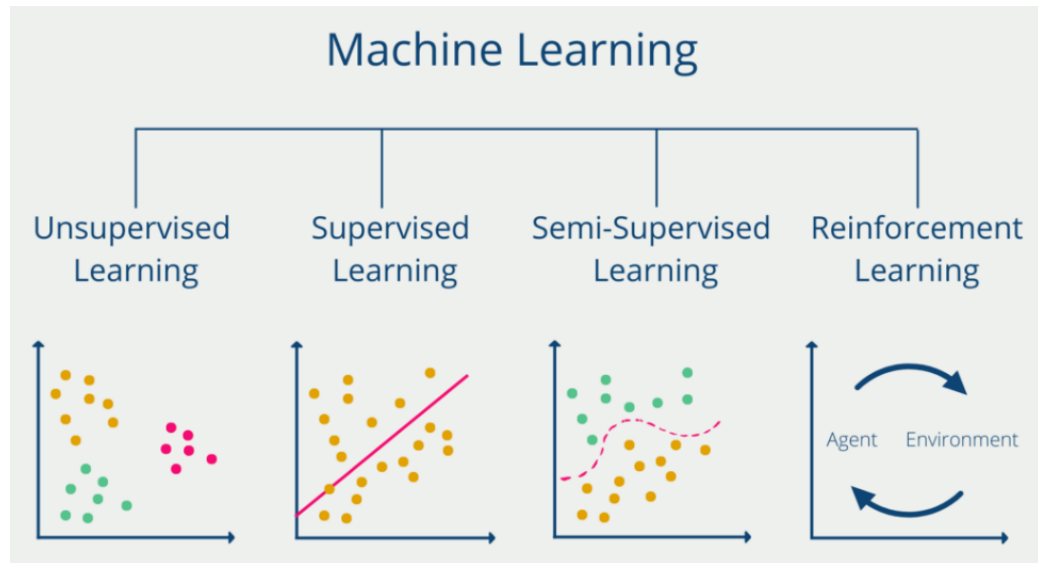


Рис. 2.2 Алгоритми машинного навчання [16]

Різні методи навчання підходять під різні потреби та задачі, але є ще деякі дуже важливі деталі в процесі навчання моделей, наприклад марковані дані.

Марковані дані неймовірно цінні для навчання ШІ-моделей, оскільки вони надають наочні приклади завдань, які має виконувати модель. Вони завжди використовуються для контрольованого або напівконтрольованого навчання моделей.

Марковані дані - це підмножина навчальних даних ШІ, які анотовані або позначені відповідною інформацією. Іншими словами, кожна точка даних супроводжується міткою або тегом, який вказує, що ці дані представляють.

Також важливою деталлю любого машинного навчання є концепція «людина в циклі» або Human-in-the-loop (HITL). Ця технологія передбачає нагляд і втручання в процес навчання моделі. Хоча ШІ-моделі можуть навчатися на чистих даних, вони не є безпомилковими і можуть припускатися помилок. Люди-експерти часто залучаються до перевірки та виправлення прогнозів моделі, особливо коли наслідки помилок є значущими.

Обов'язковим для розгляду є підготовка навчальних даних, що складається з двох етапів – попередньої обробки даних та їх тестування та валідації.

Попередня обробка даних – це очищення даних для виділення та викорінення помилок, невідповідностей або нерелевантної інформації. Також відбувається перетворення даних у формат, придатний для навчання ШІ-моделі. Наприклад, текстові дані можуть бути токенізовані, а зображення - змінені та нормалізовані.

Токенізація - розбиття тексту на менші об'єкти, які називаються токенами. Токени можуть бути абсолютно різними, залежно від типу токенізатора, який використовується. Токенізація - це один з головних кроків у кожній операції обробки природної мови. Враховуючи різноманітність існуючих мовних структур, токенізація в кожній мові відбувається по-різному.

Лематизація – теж невід'ємна частина підготовки даних для обробки природної мови. Лематизація полягає у виведенні слова до його основної форми, тобто інфінітиву, наприклад «є» буде лематизовано до форми «бути»;

Тестування та валідація даних – це процес, коли перед початком навчання нейромережевої моделі дані розділяються на тестові та перевірочні, зазвичай «test» та «train». Навчальна вибірка слугує для безпосереднього тренування алгоритму, а тестова — для оцінки його точності та якості роботи. Дані валідації/тестування допомагають точно налаштувати гіперпараметри моделі та запобігти надмірному (або недостатньому) пристосуванню. Надмірне пристосування відбувається, коли модель стає занадто спеціалізованою на своїх навчальних даних і погано працює на нових, небачених даних. Недостатня пристосованість - це протилежна ситуація.

Тестові дані повинні відрізнятися від навчальних даних і не повинні використовуватися в процесі навчання моделі. Вони слугують еталоном для вимірювання рівня продуктивності моделі, точності, запам'ятовування та інших

метрик. Якщо модель добре працює на тестових даних, вона, швидше за все, буде добре працювати і в реальних умовах.

Існують різні можливі варіанти отримання даних, можливо або самостійно збирати дані - збір за допомогою опитувань, методів вилучення даних або датчиків. Отримувати дані з публічних дата сетів – це теж дуже зручний спосіб отримання даних. Багато організацій та дослідницьких установ надають загально доступні набори даних для різних завдань ШІ.

2.3. Визначення основних методів для нейромережових класифікаційних моделей

Класифікація вчить машину сортувати речі за категоріями. Вона вчиться на прикладах з мітками (наприклад, електронних листах, позначених як «спам» або «не спам»). Після навчання вона може вирішувати, до якої категорії належать нові елементи, наприклад, визначати, чи є новий лист спамом, чи ні. Класифікуючи машинне навчання можна виділити два типи класифікаторів: лінійні та активні.

Активні класифікатори є спеціалізованими алгоритмами машинного навчання, котрі спершу займаються формуванням загальної моделі спираючись на тренувальні дані, і тільки після цього можуть займатися прогнозуванням даних. Через процеси детального налаштування вагових коефіцієнтів під час фази навчання

Через детальне налаштування ваг під час фази навчання вони потребують більше часу на підготовку, але значно пришвидшують процес прогнозування.

Більшість алгоритмів машинного навчання - це алгоритми, що піддвюються швидкому навчанню, і нижче наведено деякі основні приклади:

- машина опорних векторів;
- логістична регресія;

- дерева рішень;
- штучні нейронні мережі.

Натомість «ліниві» або прикладні класифікатори не будують узагальнену модель під час етапу навчання — звідси й назва. Вони зберігають усі тренувальні зразки та при необхідності прогнозування перебирають їх, щоб знайти найближчий екземпляр. Через це швидкість видачі результату суттєво знижується.

Прикладами ліневих класифікаторів є:

- K-Nearest Neighbor;
- міркування на основі конкретних випадків.

Для класифікаційних моделей використовують навчання з учителем, яке в свою чергу поділяється на методи класифікації та регресії (рис. 2.3).

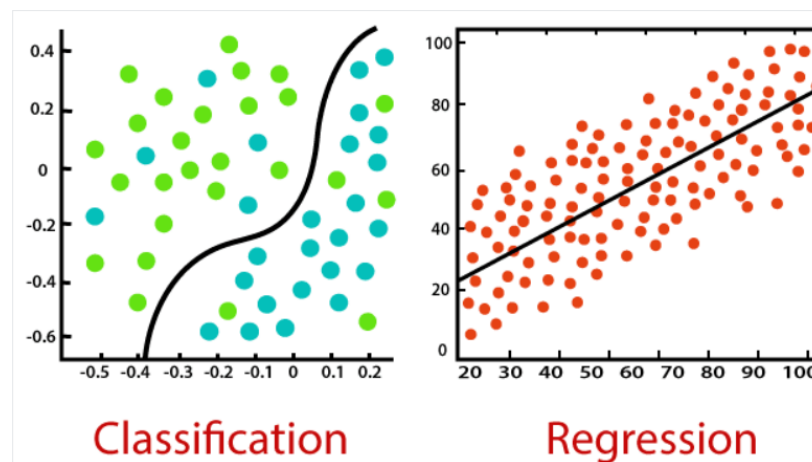


Рис. 2.3 Методи класифікації та регресії [17]

Регресія - це процес пошуку кореляції між залежними та незалежними змінними. Це допомагає у прогнозуванні безперервних змінних, таких як прогнозування ринкових тенденцій, прогнозування цін на житло тощо.

Не зважаючи на те, що немає двох абсолютно однакових алгоритмів класифікації машинного навчання, вони всі застосовують один і той самий загальний двоетапний процес класифікації даних.

Класифікаційне навчання традиційно є різновидом керованого машинного навчання, де для навчання моделей використовуються мічені дані. У керованому навчанні кожна точка навчальних даних містить вхідні змінні (також відомі як незалежні змінні або ознаки) і вихідну змінну, або мітку.

При навчанні класифікації завдання моделі полягає в тому, щоб зрозуміти зв'язки між ознаками і мітками класів, а потім застосувати ці критерії до інших наборів даних. Моделі класифікації використовують ознаки кожної точки даних разом з міткою класу, щоб розшифрувати, які ознаки визначають кожен клас. З математичної точки зору, модель розглядає кожну точку даних як кортеж x . Кортеж - це впорядкована числова послідовність, яка представляється у вигляді $x = (x_1, x_2, x_3, \dots, x_n)$.

Кожне значення в кортежі є характеристикою точки даних. Зіставляючи навчальні дані з цим рівнянням, модель дізнається, які ознаки пов'язані з кожною міткою класу.

Метою навчання є мінімізація помилок під час прогнозного моделювання. Алгоритми градієнтного спуску навчають моделі, мінімізуючи розрив між прогнозованими та фактичними результатами. Пізніше моделі можуть бути доопрацьовані за допомогою додаткового навчання для виконання більш специфічних завдань.

Класифікація. Другим кроком у задачах класифікації є власне класифікація. На цьому етапі користувачі розгортають модель на тестовому наборі нових даних. Раніше невикористані дані використовуються для оцінки продуктивності моделі, щоб уникнути надмірного пристосування: коли модель занадто сильно

покладається на свої навчальні дані і стає нездатною робити точні прогнози в реальному світі.

Модель використовує вивчену функцію прогнозування для класифікації нових даних за різними класами відповідно до особливостей кожної вибірки. Потім користувачі оцінюють точність моделі за кількістю правильно передбачених зразків тестових даних.

В машинному навчанні розрізняють чотири різні види класифікації: бінарна, мульти – класова, класифікація з багатьма маркуваннями і незбалансована класифікація.

Бінарна класифікація - це алгоритм навчання під контролем, який відносить нові спостереження до одного з двох класів. Багатокласова класифікація - це процес віднесення об'єктів до більш ніж двох класів. Кожному об'єкту присвоюється один клас без будь-якого перетину. Різниця між бінарною та мульти – класовою класифікаціями наведена на рис. 2.4.

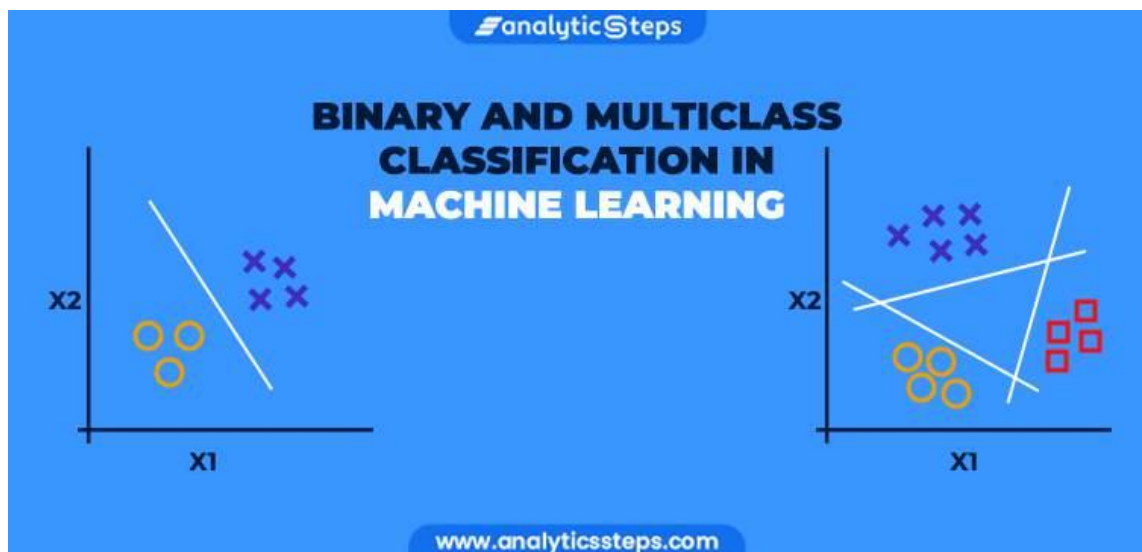


Рис. 2.4 Різниця між бінарною та мульти – класовою класифікаціями [18]

Класифікація з багатьма мітками - це задача керованого навчання, де екземпляр може бути пов'язаний з декількома мітками. Це розширення класифікації з однією міткою (тобто багатокласової або бінарної), де кожен екземпляр асоціюється лише з однією міткою класу. Різниця між класифікацією з багатьма мітками та багатокласовою класифікацією наведена на рис. 2.5.

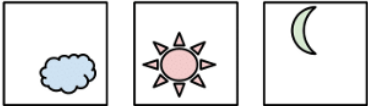
Multi-Class		Multi-Label
C = 3	Samples	Samples
		
	Labels (t) [0 0 1] [1 0 0] [0 1 0]	Labels (t) [1 0 1] [0 1 0] [1 1 1]

Рис. 2.5 Різниця між мульти – класовою класифікацією та класифікацією з багатьма мітками [19]

Незбалансована задача класифікації - це приклад задачі класифікації, в якій розподіл прикладів між відомими класами є упередженим або перекошеним. Розподіл може варіюватися від незначного зміщення до серйозного дисбалансу, коли на сотні, тисячі або мільйони прикладів у класі більшості припадає один приклад у класі меншості.

РОЗДІЛ 3. ПРАКТИЧНА РЕАЛІЗАЦІЯ НЕЙРОМЕРЕЖЕВОЇ МОДЕЛІ

3.1. Підбір датасету

Нами було розглянуто можливі варіанти відкритих для використання датасетів, що мають достатню кількість даних для якісного навчання моделі, містять в собі саме текстові описи, цитування, діалоги або коментарі людей (соціальні мережі/веб-сторінки), за ознаками характерні для психологічних розладів.

Було визначено як найбільш оптимальний для даного дослідження – датасет «Reddit Mental Health Dataset». Цей набір даних містить дописи з 28 суббредитів (15 груп підтримки психічного здоров'я) за 2018-2020 роки. Розробники застосовували цей набір даних, щоб зрозуміти вплив COVID-19 на групи підтримки психічного здоров'я з січня по квітень 2020 року, а також включили старіші часові рамки, щоб отримати базові дописи до COVID-19. [20]

В нашому дослідженні з цього датасету ми використовуємо тільки три суббредіти, за якими і відбувається розподіл на класи. Нами використовуються класи:

- ADHD – розлад дефіциту уваги та гіперактивності;
- anxiety - тривожність;
- depression – депресія;

Дані кожного класу розподілені по чотирьох файлах, а саме це дані за 2018 – 2019 рік та дані до розповсюдження COVID-19 та після розповсюдження хвороби (рис. 3.1).

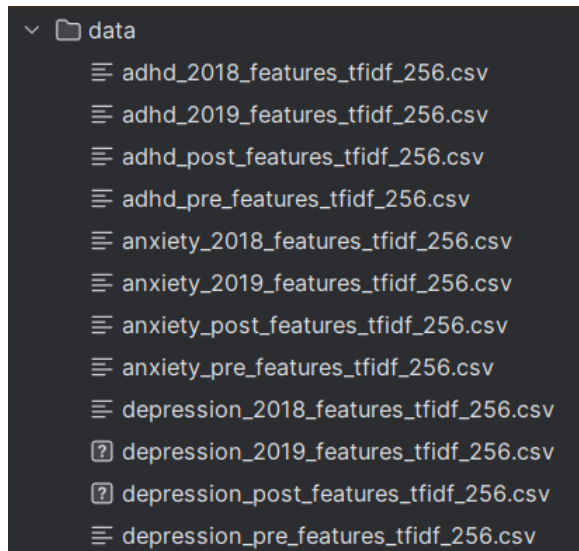


Рис. 3.1 Завантажені файли з даними

Для більшої об'єктивності моделі і розширення функціоналу нам знадобився дата сет «Sentiment Analysis Dataset». Текстові дані з колонки «text» цього дата сету були винесені як окремий файл для створення категорії «healthy» тобто «здорові» дані, відповідну колонку «category» було додано і перейменовано колонку «text» на «posts» для простішого об'єднання в один дата сет.

Для більш зручного використання даних, було прийнято рішення об'єднати усі дані у один новий файл(рис. 3.2). На рисунку 3.2 ми застосовуємо бібліотеку pandas. Pandas - це бібліотека Python для роботи з наборами даних. Вона має функції для аналізу, очищення, дослідження та маніпулювання даними. Наступним кроком ми створили списки файлів за категоріями для коректного об'єднання. Далі було застосовано функцію для завантаження файлів, категоризації та маркування файлів для кожної категорії.

```

import pandas as pd
import os
healthy_file_path = r'C:\Users\katie\PycharmProjects\PythonProject5\final_healthy_data.csv'
adhd_files = [
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\adhd_2018_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\adhd_2019_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\adhd_post_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\adhd_post_features_tfidf_256.csv'
]
anxiety_files = [
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\anxiety_2018_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\anxiety_2019_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\anxiety_post_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\anxiety_pre_features_tfidf_256.csv'
]
depression_files = [
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\depression_2018_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\depression_2019_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\depression_post_features_tfidf_256.csv',
    r'C:\Users\katie\PycharmProjects\PythonProject5\data\datasetparts\depression_pre_features_tfidf_256.csv'
]

def process_files(file_paths, category_name):
    processed_data = []
    for file_path in file_paths:
        df = pd.read_csv(file_path)
        if 'post' in df.columns:
            df = df[['post']].dropna(subset=['post'])
            df['category'] = category_name
            processed_data.append(df)
        else:
            print(f"Warning: 'post' column not found in {file_path}. Skipping.")

```

Рис. 3.2 Об'єднання даних у один загальний файл

Далі ми поєднали усі класи в один загальний файл «final_combined_dataset.csv», але з певними обмеженнями з урахуванням потужності робочого пристрою і збалансованості даних. За результатами підготовки даних ми отримали загальний файл, що містить чотири категорії:

- тривожність;
- РДУГ;
- депресія;
- здорові дані.

За допомогою бібліотеки `matplotlib.pyplot` нами було візуалізовано баланс класів фінального дата сету (рис.3.3) з обмеженнями по 25000 одиниць на кожен клас, задля забезпечення об'єктивного порівняння.

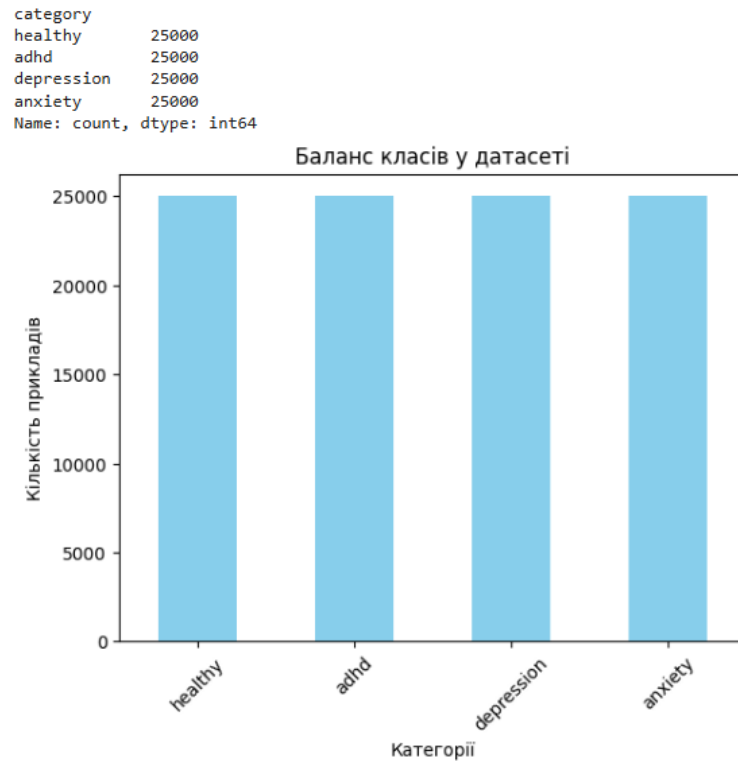


Рис. 3.3 Баланс класів

3.2. Підбір моделей для класифікації

3.2.1. Дослідження моделей SVM, Logistic Regression, DistilBERT

Нами було вирішено розглянути різні моделі класифікації та регресії, і за результатами метрик оцінювання навчання класифікації визначити ту – яка підійде найкраще.

Дослідники та розробники обирають певні метрики оцінки для моделей класифікації залежно від конкретного завдання класифікації. Всі вони вимірюють точність, з якою учні або класифікатори точно передбачають класи моделі.

Нами було протестовано три різні види моделей, а саме: Logistic Regression, Support Vector Machine (SVM) та DistilBERT.

Моделі SVM та Logistic Regression - це керовані алгоритми машинного навчання. Вони обидва використовуються для вирішення завдань класифікації (сортування даних за категоріями).

Логістична регресія спрямована на вирішення проблем класифікації. Вона робить це, прогножуючи категоричні результати, на відміну від лінійної регресії, яка прогнозує безперервний результат.

У найпростішому випадку є два результати, які називаються біноміальними, прикладом яких є передбачення того, чи є пухлина злоякісною або доброякісною. Інші випадки мають більше двох результатів для класифікації, в цьому випадку вони називаються мультиноміальними.

Машина опорних векторів (SVM) – є керованим алгоритмом машинного навчання, що використовується для таких задач як: класифікація та регресія. Хоча SVM може впоратися з завданнями регресії, він найкраще підходить для завдань класифікації.

SVM знаходить оптимальну гіперплощину в N-вимірному просторі для розділення точок даних на різні класи. Алгоритм максимізує відстань між найближчими точками різних класів.

Різниця між SVM та логістичною регресією:

1. SVM намагається знайти «найкращу» межу (відстань між лінією та опорними векторами), яка розділяє класи, і це зменшує ризик помилки в даних, тоді як логістична регресія цього не робить, натомість вона може мати різні межі рішень з різними вагами, які знаходяться поблизу оптимальної точки;
2. SVM добре працює з неструктурованими і напівструктурованими даними, такими як текст і зображення, тоді як логістична регресія працює з уже визначеними незалежними змінними;

3. SVM базується на геометричних властивостях даних, тоді як логістична регресія базується на статистичних підходах;
4. ризик перенавчання є меншим у SVM, тоді як логістична регресія є вразливою до перенавчання.

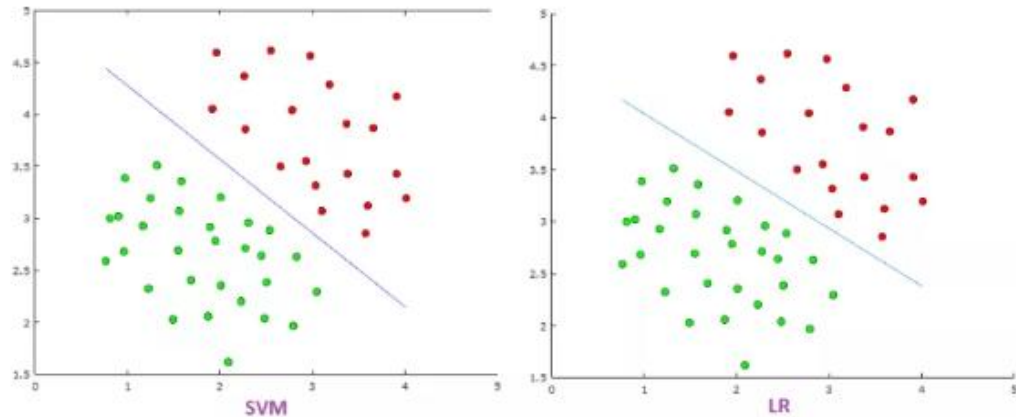


Рис. 3.4 Різниця між SVM та логістичною регресією [21]

DistilBERT попередньо тренує меншу і швидшу універсальну модель представлення мови, яку можна точно налаштувати, щоб отримати хорошу продуктивність для багатьох наступних завдань.

Основною мотивацією для створення DistilBERT є потреба в компактних і більш ефективних моделях для NLP-додатків, які працюють в режимі реального часу на периферійних пристроях. Мовні моделі, такі як BERT, часто вимагають значних обчислювальних ресурсів, що обмежує їх розгортання в середовищах з обмеженими ресурсами.

Студентська модель DistilBERT має таку ж загальну архітектуру, як і BERT, але з деякими змінами. Вбудовування типу токенів і пулер вилючено, а кількість шарів зменшено в 2 рази. BERT Base має 12 шарів, а DistilBERT - лише 6. Ці архітектурні зміни роблять DistilBERT більш ефективною, і під час навчання були зроблені подальші вдосконалення такі як:

Ініціалізація студента: Важливою частиною навчання DistilBERT є визначення того, як ініціалізувати модель учня, щоб вона добре навчалася. Оскільки BERT (вчитель) і DistilBERT (учень) мають однакову розмірність, учень ініціалізується шляхом взяття одного шару з двох з моделі вчителя.

Дистиляція: DistilBERT було дистильовано дуже великими партіями (до 4 тис. прикладів на партію) з використанням динамічного маскування. Завдання передбачення наступного речення (NSP) було видалено, щоб модель могла більше зосередитися на розумінні замаскованих слів у кожному реченні. Ця концепція була взята з RoBERTa для забезпечення більш надійної роботи.

Дані та обчислювальна потужність: DistilBERT навчалася на тому ж наборі даних, що й оригінальна модель BERT, за допомогою конкатенації англійської Вікіпедії (2500 мільйонів слів) і Торонтського книжкового корпусу (800 мільйонів слів), на 8 16-гігабайтних графічних процесорах V100 протягом приблизно 90 годин.

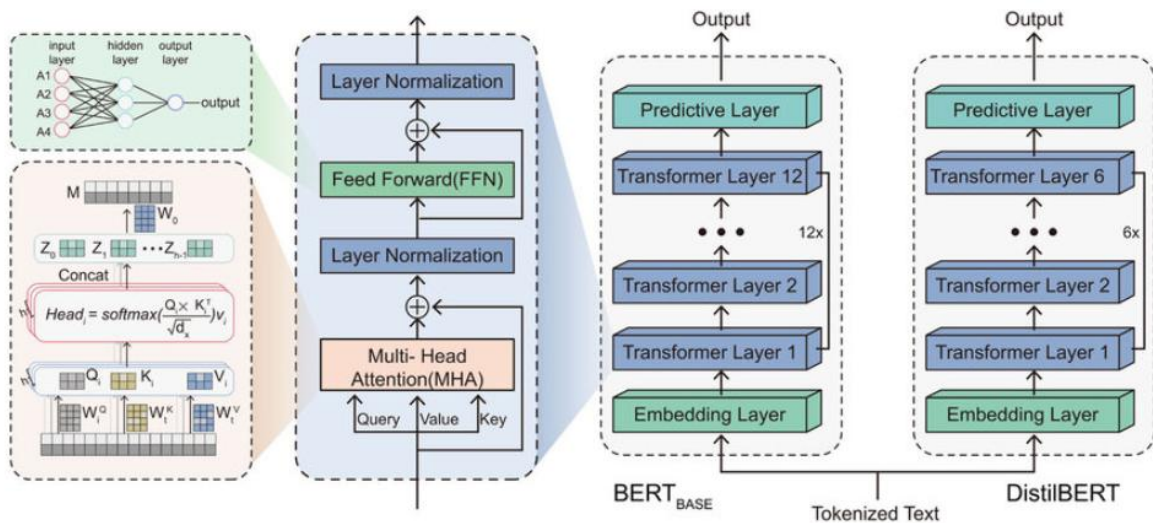


Рис. 3.5 Різниця між BERT та DistilBERT [22]

3.2.2. Створення моделі DistilBERT

Створена нами модель DistilBERT розроблялася в умовах обмеженої потужності на власній локальній машині. Спершу нами було завантажено та підготовлено файл `final_combined_dataset.csv`, та впевнилися в наявності колонок, що потрібні для тренування і відсутності пустих значень. Нами було закодовано мітки за допомогою `LabelEncoder`, щоб текстові мітки (`adhd`, `anxiety`, `depression`, `healthy`) були перетворені у числові (рис. 3.6.).

```
label_map = {  
    "LABEL_0": "adhd",  
    "LABEL_1": "anxiety",  
    "LABEL_2": "depression",  
    "LABEL_3": "healthy"  
}
```

Рис. 3.6 Мапа міток

На рис. 3.7. нами було розділено дані на тренувальний та тестовий набори (80% на тренування та 20% на тестування), при цьому було збережено баланс класів за допомогою параметру `stratify`.

```
# Розділення на тренувальну та тестову вибірки  
train_texts, test_texts, train_labels, test_labels = train_test_split(  
    *arrays: df['post'].tolist(), df['label'].tolist(), test_size=0.2, random_state=42, stratify=df['label'])
```

Рис. 3.7 Розподіл на тренувальний та тестовий набори

Також, тексти було токеноізовано за допомогою DistilBertTokenizerFast, що модель вільно приймала значення. Токенізація включає обрізку, паддінг та обмеження максимальної довжини тексту. Було створено клас TextDataset, що дозволив зберігати токеноізовані тексти та мітки для роботи з DataLoader у PyTorch і завантажено модель попередньо натреновану модель DistilBERT з можливістю класифікації тексту. Налаштувавши оптимізатор AdamW (рис.3.8) ми виставили для тренування оптимальну кількість епох – 3, а згодом для покращення результатів моделі було додано ще 2 епохи.

```
from torch.optim import AdamW

train_loader = DataLoader(train_dataset, batch_size=8, shuffle=True)
test_loader = DataLoader(test_dataset, batch_size=8)

optimizer = AdamW(model.parameters(), lr=5e-5)
```

Рис. 3.8 Налаштування оптимізатору AdamW

Для обробки батчів ми використали DataLoader і здійснили кроки зворотного поширення помилки та оновлення параметрів моделі. Після завершення тренування модель було переведено в режим eval для оцінки її ефективності на тестовому наборі. Збираються прогнози та виводиться classification_report, що виводить метрики точності та F1-міри для кожного класу. Останнім кроком нами було збережено модель, задля спрощення подальшого деплойменту.

3.2.3. Створення моделей SVM та Logistic Regression

За допомогою бібліотеки `pandas` ми завантажили файл «final_combined_dataset.csv», де кожен запис містить текстові повідомлення (стовпець `post`) та відповідні категорії (стовпець `category`). Далі відбувається перевірка та видалення порожніх рядків, перевірка на значення NaN за методом `dropna`. Наступним етапом проведено очищення тексту, для приведення кожної строки до стандартного формату, процес складається з певних етапів (рис. 3.9):

```
stop_words = set(stopwords.words('english'))

def clean_text(text):
    text = text.lower()
    text = re.sub(pattern: r"http\S+|www\S+|https\S+", repl: '', text)
    text = re.sub(pattern: r"[^a-z\s]", repl: '', text)
    text = re.sub(pattern: r'\s+', repl: ' ', text).strip()
    tokens = text.split()
    tokens = [t for t in tokens if t not in stop_words]
    return ' '.join(tokens)

df['clean_post'] = df['post'].apply(clean_text)
```

Рис. 3.9 Очищення тексту

- перетворення тексту в нижній регістр;
- видалення URL-адрес з тексту (за допомогою регулярних виразів);
- видалення символів, які не є літерами (залишаючи лише букви англійського алфавіту);
- видалення зайвих пробілів та нормалізація тексту;

- видалення стоп-слів (використовуючи бібліотеку `nlTK` для доступу до списку стоп-слів англійської мови). Стоп-слова — це слова, які не несуть значної інформації для задачі, наприклад, "і", "в", "на".

Наступним етапом текстові дані перетворюються на числові вектори за допомогою TF-IDF векторизатора з максимальним числом ознак 5000. Це дозволяє моделі працювати з числовими даними замість сирого тексту. У результаті, кожен текст стає вектором ознак, які відображають важливість слів у тексті, враховуючи їх частоту в документі та у всьому корпусі.

Далі, так само як і при створенні моделі DistilBERT, розділяються на тренувальну (80%) і тестову (20%). Переходячи до навчання моделі (рис. 3.10):

```
log_model = LogisticRegression(max_iter=1000)
log_model.fit(X_train, y_train)
log_preds = log_model.predict(X_test)

print("Logistic Regression:")
print(classification_report(y_test, log_preds))

svm_model = SVC(kernel="linear", probability=True)
svm_model.fit(X_train, y_train)
svm_preds = svm_model.predict(X_test)

print(" Linear SVM:")
print(classification_report(y_test, svm_preds))
```

Рис. 3.10 Навчання моделей

- логістична регресія: Модель логістичної регресії тренується на тренувальних даних. Після цього проводиться прогноз на тестових даних, і виводиться `classification report`, який містить метрики якості моделі (точність, відчутність, F1-міра);

- лінійний SVM: Створюється модель лінійного Support Vector Machine (SVM), яка також навчається на тих самих тренувальних даних. Після тренування SVM також робить прогнози на тестових даних і виводиться classification report.

Після навчання моделей та векторизатора вони зберігаються у файли за допомогою pickle (рис. 3.11):

```
with open("vectorizer.pkl_try", "wb") as f:
    pickle.dump(vectorizer, f)

with open("logistic_regression_model_try.pkl", "wb") as f:
    pickle.dump(log_model, f)

with open("svm_model_try.pkl", "wb") as f:
    pickle.dump(svm_model, f)

print(" Моделі та векторизатор збережено.")
```

Рис. 3.11 Збереження навчених моделей

- модель векторизатора (vectorizer.pkl_try);
- моделі логістичної регресії (logistic_regression_model_try.pkl) і SVM (svm_model_try.pkl).

Важливою деталлю є використання різних векторизаторів при створенні моделей. Для моделі DistilBERT ми використали DistilBertTokenizerFast, а для моделей SVM та Logistic Regression — TF-IDF, щоб забезпечити об'єктивне порівняння ефективності різних підходів до векторизації тексту.

3.3. Оцінка ефективності створених моделей

Для якісної оцінки та порівняння створених моделей нами було визначено метрики за якими ми проведемо оцінку:

- accuracy - відповідає на питання «Скільки з усіх прогнозів, які ми зробили, справдилися?»;
- precision - вимірює точність позитивних прогнозів. Відповідає на питання: «З усіх елементів, які модель позначила як позитивні, скільки з них насправді виявилися позитивними?»;
- recall - вимірює здатність моделі знаходити всі позитивні приклади. Відповідає на питання: «Скільки з усіх фактичних позитивних елементів модель правильно ідентифікувала?»;
- f1-score - середнє гармонійне значення точності та пригадування. Він збалансовує дві метрики в одне число, що робить його особливо корисним, коли між точністю та пригадуванням існує компроміс;
- confusion Matrix (з обговоренням де модель помиляється) - Матриця помилок - це проста таблиця, яка показує, наскільки добре працює модель класифікації, порівнюючи її прогнози з фактичними результатами. Вона розбиває прогнози на чотири категорії: правильні прогнози для обох класів (істинно позитивні та істинно негативні) та неправильні прогнози (хибно позитивні та хибно негативні).

Метрики ефективності для кожного класу окремо:

- логістична регресія;

	precision	recall	f1-score	support
adhd	0.93	0.90	0.91	4909
anxiety	0.90	0.85	0.87	5013
depression	0.86	0.87	0.86	5072
healthy	0.92	0.99	0.95	5006
accuracy			0.90	20000
macro avg	0.90	0.90	0.90	20000
weighted avg	0.90	0.90	0.90	20000

Рис. 3.12 Метрики ефективності моделі логістична регресія

- SVM (Support Vectors Machines);

	precision	recall	f1-score	support
adhd	0.92	0.90	0.91	4909
anxiety	0.89	0.85	0.87	5013
depression	0.85	0.87	0.86	5072
healthy	0.93	0.98	0.96	5006
accuracy			0.90	20000
macro avg	0.90	0.90	0.90	20000
weighted avg	0.90	0.90	0.90	20000

Рис. 3.13 Метрики ефективності SVM (Support Vectors Machines)

– DistilBERT;

	precision	recall	f1-score	support
adhd	0.94	0.93	0.93	3750
anxiety	0.88	0.88	0.88	3750
depression	0.88	0.86	0.87	3750
healthy	0.97	1.00	0.98	3750
accuracy			0.92	15000
macro avg	0.92	0.92	0.92	15000
weighted avg	0.92	0.92	0.92	15000

Рис. 3.14 Метрики ефективності DistilBERT

За рисунками наведеними вище, модель DistilBERT перевершує інші моделі по всіх метриках. Логістична регресія та SVM показали мінімальну різницю в результатах.

DistilBERT демонструє кращі результати порівняно з традиційними моделями завдяки можливості створювати контекстуально-залежні векторні представлення слів, базуючись на трансформерній архітектурі. Такий підхід дає змогу враховувати значення слова в конкретному мовному оточенні, відображати складні нелінійні взаємозв'язки між словами та глибше розуміти зміст тексту.

Натомість Logistic Regression і SVM застосовують лінійні методи разом із фіксованими векторами TF-IDF, які не враховують контекстуальні, синтаксичні та семантичні особливості, що обмежує їхню здатність до точного класифікування, призводячи до схожих показників точності та F1-оцінок.

Матриці плутанини показують, наскільки добре кожна модель класифікувала відповідні класи.

Порівняємо матриці плутанини:

1. особливості матриці плутанини моделі логістична регресія:

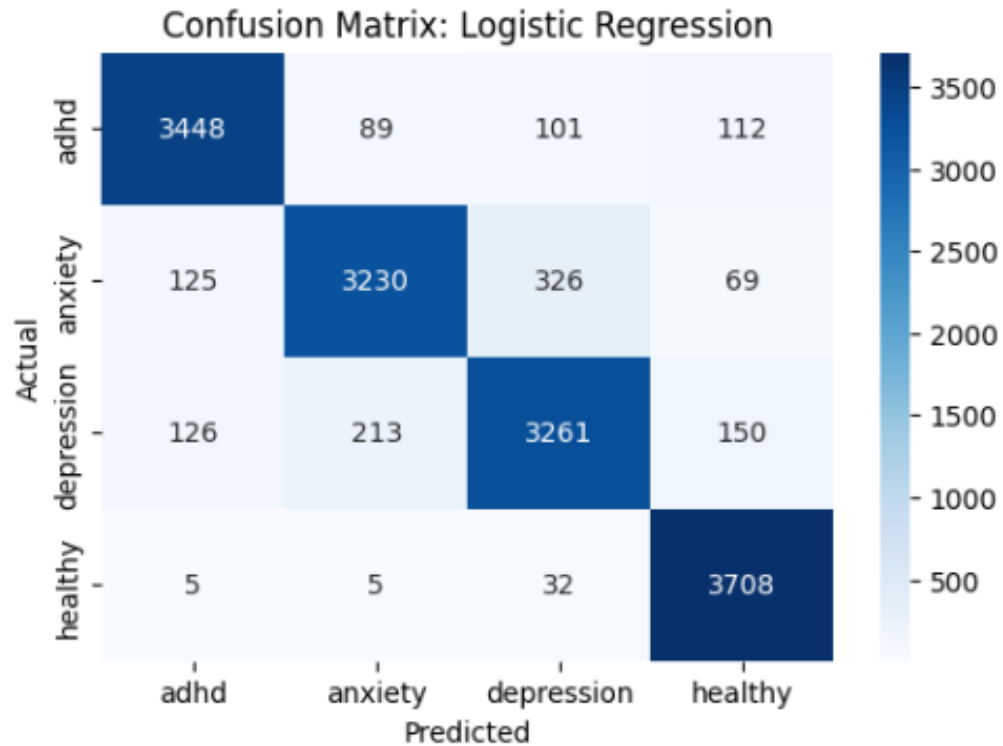


Рис. 3.15 Матриця плутанини моделі логістична регресія

- має стабільні метрики, з високою точністю та F1-мірою, що свідчить про ефективність на рівні всіх класів;
- точність майже відповідає 90% для кожного класу, що є досить гарним показником для загального застосування;
- більшість плутанини створюють класи "anxiety" і "depression". Зокрема, частина прикладів "anxiety" класифікуються як "depression", і навпаки. Це може свідчити про схожість цих класів за змістом у текстах;
- слабкіше розрізняє клас "healthy", хоча ця помилка не критична, тому що клас має високий recall (99%).

2. матриця плутанини моделі Support Vectors Machine (SVM);

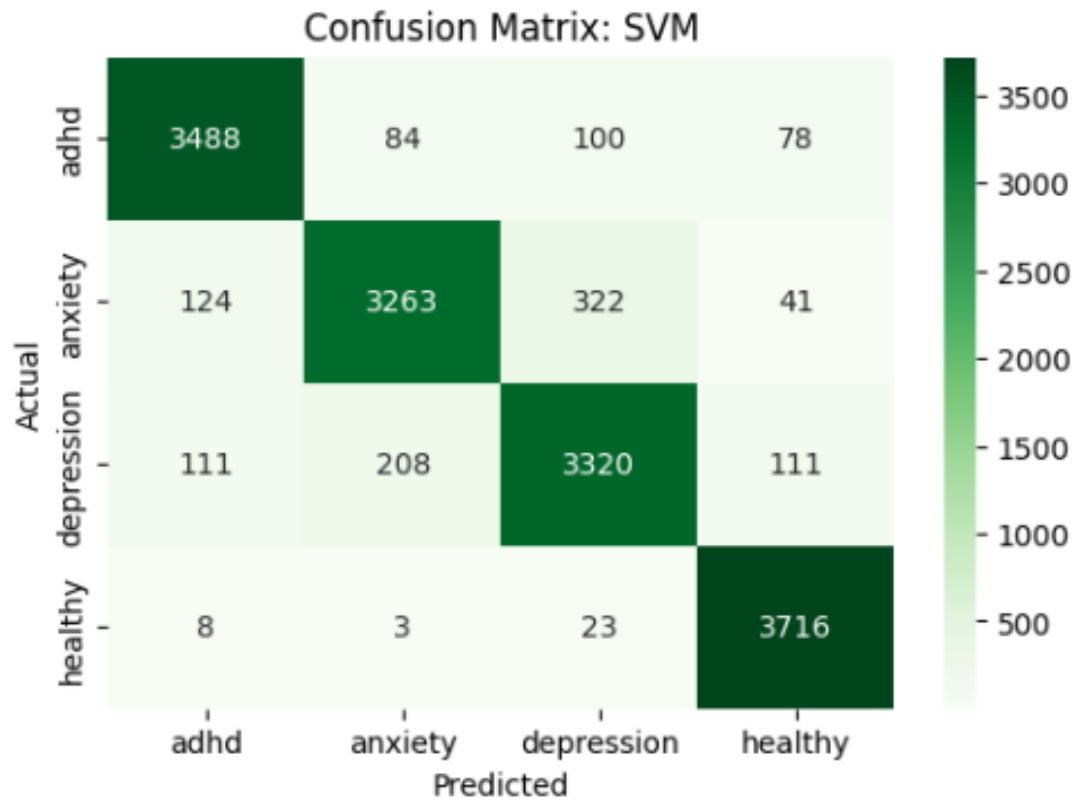


Рис. 3.16 Матриця плутанини моделі Support Vectors Machine (SVM)

- модель добре справляється з класами "adhd" і "healthy", особливо останній має відмінні показники (100% recall);
- як і в логістичній регресії, модель часто плутає класи "anxiety" і "depression". Однак тут також є й інші помилки, зокрема певна кількість помилок у класі "healthy" — особливо у випадках, коли "healthy" плутається з "adhd" або "anxiety";
- однією з можливих причин є те, що SVM схильний до створення лінійних розмежувальних гіперплощин, що може не так добре працювати для текстових даних, де класи можуть бути складно лінійно роздільними.

3. матриця плутанини моделі Support Vectors Machine (SVM);

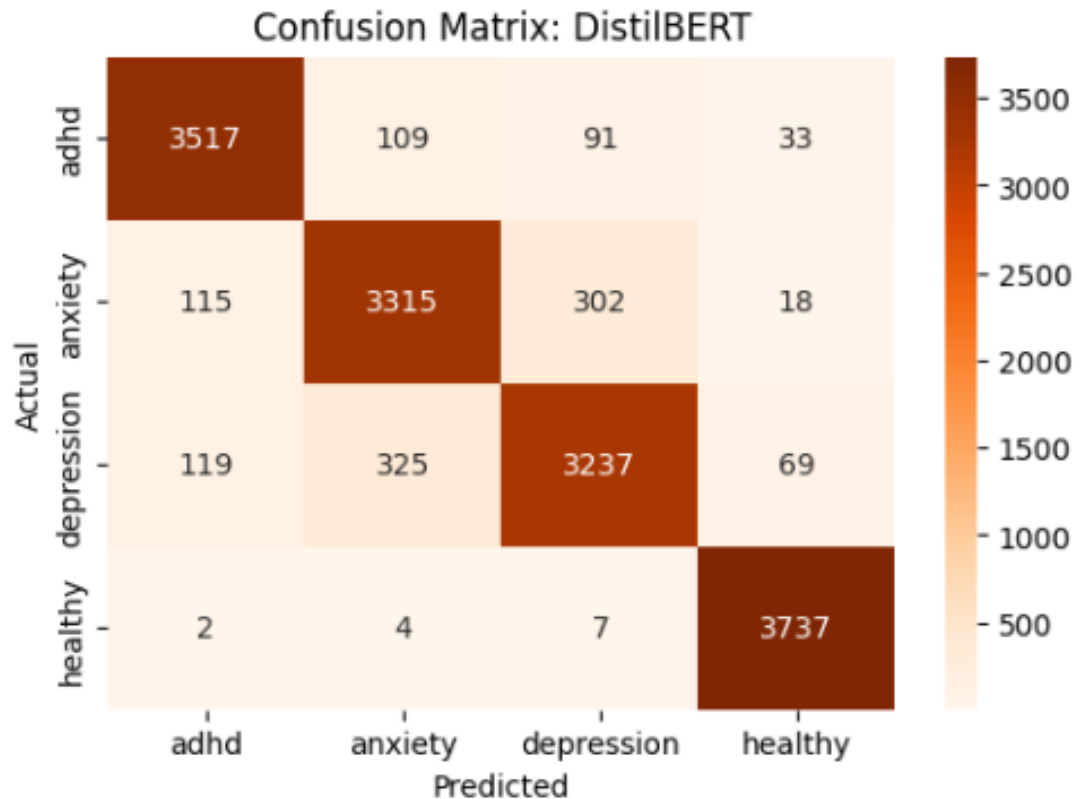


Рис. 3.17 Матриця плутанини DistilBERT

- найкращі показники серед усіх моделей, з дуже високими показниками для класу "healthy", де recall досягає 100%;
- в цілому, цей метод має кращу здатність до класифікації завдяки використанню трансформерної архітектури, яка добре працює з текстами та контекстом;
- виявлено деякі плутанини між "adhd" та "anxiety", а також між "depression" та "anxiety". Це може свідчити про схожість цих класів за змістом текстів, і в силу їх контекстної подібності модель може помилково класифікувати;
- однак загалом модель демонструє найкращі результати, тому ці помилки мають менший вплив на загальні показники.

Отже, DistilBERT - найточніша модель, з меншою кількістю помилкових класифікацій та кращою загальною ефективністю. Показники моделі SVM порівняно з логістичною регресією, але з дещо більшою кількістю помилкових класифікацій для тривоги та депресії. Логістична регресія забезпечує прийнятні результати, але має більше помилкових класифікацій, особливо для тривоги та депресії.

3.4. Порівняння створених моделей за базовими характеристиками

Під час створення моделі нами фіксувалися базові характеристики кожної моделі (таблиця 3.1).

Таблиця 3.1

Параметр	Logistic Regression	SVM(linear)	DistilBERT
Датасет	100000 записів	100000 записів	100000 записів
Векторизатор/ токенізатор	TF-IDF (max_features=5 000)	TF-IDF (max_features = 5 000)	HuggingFace DistilBERT- tokenizer, max_length=256
Апаратне забезпечення	Intel Core i7 (8 ядер @ 3.6 GHz) - CPU	Intel Core i7 (8 ядер @ 3.6 GHz) -CPU	Intel Core i7 (8 ядер @ 3.6 GHz) 32 GB RAM NVIDIA RTX 2050 (4 GB VRAM)
Час тренування	≈ 1 хвилина на одну ітерацію	≈ 2 години	≈ 50 хвилин на одну епоху, 8 епох - ≈7 годин

Продовження таблиці 3.1

Параметр	Logistic Regression	SVM(linear)	DistilBERT
Кількість параметрів	5000	(non-parametric, $SV \approx 8\,000$)	≈ 66 млн
CPU-навантаження під час тренування	$\approx 60\%$	$\approx 70\%$	$\approx 30\%$
GPU-навантаження під час тренування	Не навантажувалося	Не навантажувалося	$\approx 85\%$
Інтерференс (один зразок)	0.5 мс (CPU)	1.0 мс (CPU)	50 мс (CPU) / 10 мс (GPU)
Розмір моделі на диску	~ 1 MB	~ 5 MB	~ 250 MB

Спираючись на вищезазначені метрики нами було зроблено певні висновки. Найшвидшою моделлю виявилась модель Логістичної реакресії, яку було натреновано менше ніж на хвилину на одній ітерації вибірки. Це можливо пояснити простотою алгоритму та реалізацією за допомогою бібліотеки `scikit-learn`. У нашому випадку лінійний SVM, який навчався виключно на CPU без жодних додаткових оптимізацій, зайняв понад дві години (~ 2 год 10 хв) на один прохід. Незважаючи на просте лінійне ядро, цей алгоритм використовує квадратичне програмування, через що його продуктивність сильно знижується при обробці великих обсягів даних і багатовимірних TF-IDF векторів. DistilBERT показав середній результат, близько 50 хвилин на одну епоху. Якщо розглядати абсолютні часові витрати, то модель DistilBERT за витратами часу на навчання є швидшою ніж SVM.

За апаратним забезпеченням SVM та Logistic Regression не потребували застосування GPU, а DistilBERT потребує велику потужність графічного процесору.

Підбиваючи підсумки, якщо є потреба в швидкому навчанні та економії оперативної пам'яті – ідеальним варіантом буде модель логістичної регресії, але вона проявляє себе гірше при реальному тестуванні. У випадках коли задача потребує високої точності, як в нашому випадку, і є можливість задіювати GPU, то краще використовувати модель DistilBERT, він забезпечить відносно швидке навчання, кращу точність і розуміє контекст. Лінійний SVM добре показує себе на середній обсягах даних, а на великих без додаткової оптимізації суттєво поступається часом виконання і потребує додаткових налаштувань ядра.

3.5. Можливості покращення створеної моделі

DistilBERT модель показала кращі результати ніж усі інші моделі, але вона ще має необмежену кількість можливих варіантів модифікації та допрацювання. Спершу виділимо покращення, що можна пов'язати з датою сетом та даними:

- збільшити кількість та варіативність даних: урізноманітнити вже існуючий дата сет, зібрати дані різні за обсягом тексту від великих до малих, а не тільки узагальненого обсягу. Також додавати більше класів наприклад ПТСР та біполярний розлад;
- змінити пропорційність дати сету відповідно до результативності: на рисунку 3.17 візуалізовані помилки та плутанина між відповідними класами, тож враховуючи це потрібно глибше проаналізувати і додати дані в класи такі як anxiety та depression, щоб у моделі під час навчання було більше прикладів за допомогою яких вона зможе покращити розуміння

- відмінностей між цими класами. Також можливо використовувати метод доповнених даних, наприклад синонімізація тексту або зворотній переклад;
- мультимодальність даних. Використання не тільки текстових даних, але і аудіо формату, задля збільшення можливостей моделей.

Однією з найважливіших вимог при розгортанні моделей типу BERT є інфраструктура, тож одним з покращень буде використання не локальної машини, а спеціалізованого серверного обладнання для ефективного розгортання та підтримки роботи моделі. Також при наявності відповідної інфраструктури можливо покращувати і архітектуру самої моделі і пробувати застосовувати інші трансформерні моделі такі як: RoBERTa або спеціалізовані моделі - MentalBERT або ClinicalBERT.

Використання більш потужних векторизаторів теж може виявитися ефективним покращенням, наприклад - Sentence-BERT, Universal Sentence Encoder або InferSent.

Не менш важливою деталлю є методи навчання та оптимізація моделі, потенційним покращенням є налаштування гіперпараметрів за допомогою Optuna, Ray Tune або GridSearch. Використання побічних аналізів, а саме окрім основної діагностики впровадити функції розрізнення змін настрою, класифікації емоцій і можливість надання рекомендацій та реакцій щодо покращення стану.

РОЗДІЛ 4. ЕКОНОМІЧНО – ПРАКТИЧНЕ ВИКОРИСТАННЯ

4.1. Практичне застосування нейромережевої моделі для діагностики психологічних розладів

Нейромережеві моделі відкривають нові горизонти у діагностиці розладів психіки. Вони можуть стати надійними помічниками як для лікарів, так і для самих пацієнтів, допомагаючи у самодіагностиці та підготовці до консультації. Це робить медичні послуги доступнішими, а роботу медичних установ — менш перевантаженою. У світі, де діджиталізація та автоматизація набирають обертів, штучний інтелект займає все важливіше місце. Тож наша модель, з подальшим розвитком, має потенціал гнучко адаптуватися під різноманітні потреби та задачі.

Можливі напрямки впровадження.

Модель може бути застосована на веб-сторінках, платформах та цифрових додатках, які дозволятимуть користувачам проводити самостійну профілактичну оцінку свого стану психічного здоров'я, отримувати корисні поради, збирати статистику.

До прикладу, людина може заповнити анкетування чи ввести текстовий опис своїх емоційних переживань, на рисунках 4.1 та 4.2 зображені варіанти інтерфейсу для подібного застосунку, після чого система проаналізує ці дані й надасть попередній висновок щодо ймовірності наявності депресії, тривожних розладів або інших психічних порушень. Такі платформи можуть використовуватися як для системної допоміжної терапії та самопізнання пацієнта так і лікарями.

Mood & Mental Health Checker

Welcome! This tool can help you reflect on how you're feeling. You can write freely about **your thoughts**, or **answer some guided questions** to describe your mood. *(This is not a medical diagnosis. If you're in crisis, please seek professional help.)*

How would you like to share your feelings?

- ☒ Write my own text
☐ Answer some guided questions

Describe how you're feeling:

I don't care about anything anymore. I just feel numb, and everything seems pointless. I've been in a long-term relationship, but even when I'm with her, I don't feel like I connect anymore. I don't know how to explain it, but I feel distant from everyone and everything.

Analyze My Mood

Results:

Logistic Regression: depression

Linear SVM: depression

DistilBERT: depression (confidence: 1.00)

This tool is for informational purposes only and is not a substitute for professional diagnosis or treatment. If you're experiencing distress or thoughts of self-harm, please contact a healthcare professional or emergency services.

Рис. 4.1 Приклад інтерфейсу для аналізу тексту

Mood & Mental Health Checker

Welcome! This tool can help you reflect on how you're feeling.

You can **write freely about your thoughts**, or

answer some guided questions to describe your mood.

(This is not a medical diagnosis. If you're in crisis, please seek professional help.)

How would you like to share your feelings?

- ☐  Write my own text
☒ Answer some guided questions

✱ What's been on your mind lately?

I am really nervous about my diplomas

🧘 Have you been feeling more anxious, sad, or restless?

I am feeling nervous and restless

🛌 How has your sleep or energy been?

I don't have enough energy to even wake up at the morning

🗨 Anything else you'd like to share?

I feel like i have no people to support me

Analyze My Mood

Results:

Logistic Regression: anxiety

Linear SVM: anxiety

DistilBERT: borderline between 'adhd' and 'anxiety' (confidence: 0.48)

This tool is for informational purposes only and is not a substitute for professional diagnosis or treatment.

If you're experiencing distress or thoughts of self-harm, please contact a healthcare professional or emergency services.

Рис. 4.2 Приклад інтерфейсу для самодіагностики за опорними запитаннями

Допоміжний – консультаційний засіб для лікарів: Платформи на основі неймережевих моделей для діагностики психологічного стану можуть надавати лікарям додаткові інструменти для точнішої діагностики або консультації/отримання іншої думки у випадках, коли є сумніви або неоднозначність. Наприклад, система може надавати лікарю рекомендації щодо коректності діагнозу на основі тестувань, анкетувань та текстів. Платформа може виводити показники у вигляді діаграм або графіків, що ілюструють ймовірнісні оцінки за кожною категорією психіатричного захворювання, що може бути допоміжним засобом при постановці діагнозів та виборі терапії.

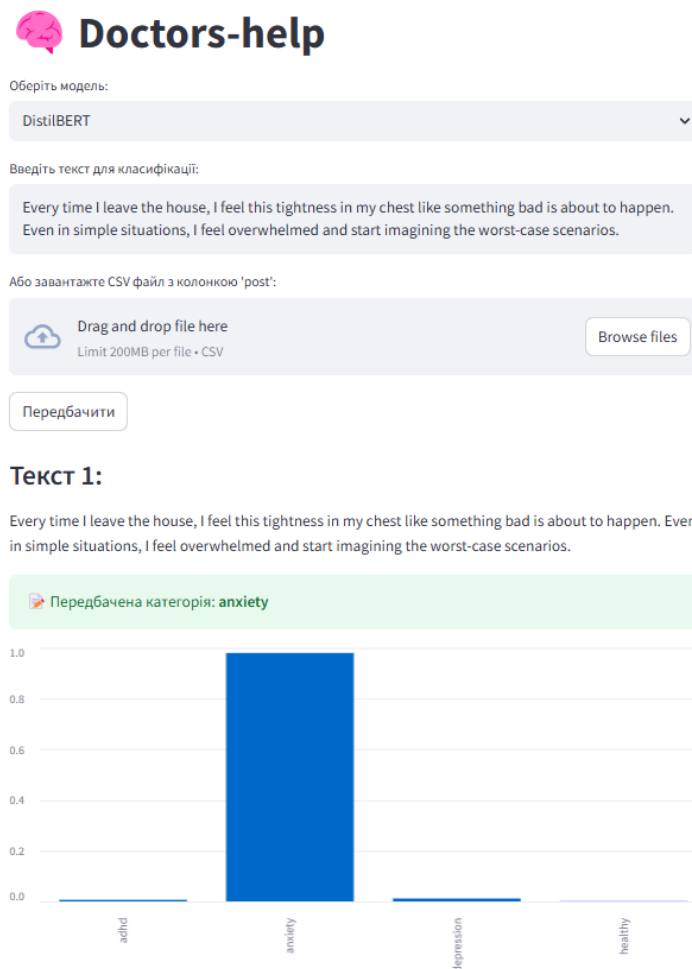


Рис. 4.3 Інтерфейс платформи допомоги лікарям

Модель може використовуватися для проведення аналізу постів та публікацій користувачів у соціальних мережах з метою виявлення ранніх ознак психологічних розладів. Це надасть можливість вчасно звертати увагу на осіб, що потребують допомоги, але не можуть отримати і звернутися по неї самостійно з певних причин.

Також моделі такого характеру можуть стати помічниками на роботі та у компаніях. Це може оптимізувати та робочий процес завдяки покращенню контролю психологічного стану, навантаження та втомлюваності співробітників, покращити атмосферу на робочому місці та в колективі в цілому.

4.2. Аналіз потенційних витрат на впровадження та обслуговування моделі

Створення якісної нейромережевої моделі може потребувати великих вкладень та ресурсів, таких як:

- кваліфіковані фахівці;

психологи та психіатри для узгодження та постійних консультацій для запобігання хибних даних та некоректності системи. Також мають проводитися регулярні консультації для коригування моделі і перевірки нових елементів.

Фахівці з анотування даних потрібні для коректного розподілу та позначення категорій для текстів, що описують психологічні розлади або мають їх ознаки.

Розробники, а саме data scientists, ML engineers – для створення, налаштування та підтримки моделі.

Маркетинг та реклама також потребують кваліфікації та знань, використання реклами у соціальних мережах, просування через спеціалізовані платформи, для цього потрібна люди з досвідом та знаннями в цій галузі.

- Обчислювальні ресурси;

нейромережеві моделі, особливо трансформерні вимагають потужних обчислювальних можливостей. Обов'язково потрібні GPU (Графічний процесор – це електронна схема, що застосовується для цифрової обробки зображень і прискорення комп'ютерної графіки, присутня або у вигляді дискретної відеокарти, або вбудована в материнські плати, мобільні телефони, персональні комп'ютери, робочі станції та ігрові приставки.) або TPU (Тензорні процесори - це спеціалізована розробка Google - інтегральні схеми (ASIC), що застосовуються для пришвидшення машинного навчання.), для цього можна користуватися або хмарними середовищами, такими як: AWS, Google Cloud, Microsoft Azure.

Вартість може варіюватися від 2 до 10 доларів на годину, при цьому все залежатиме від обсягу даних та складності моделі.

Також замість хмарного середовища можливо застосовувати власне серверне обладнання, яке потрібно купляти та налаштовувати, що потребуватиме додаткових витрат. Отже, вартість обладнання може вартувати від 5000 доларів до 20000 доларів і варто зауважити, що фізичне обладнання потребує постійного охолодження, налаштування та обслуговування.

Важливою деталлю є зберігання даних, для чого можливо використовувати або хмарні сховища, або дисковий фізичний простір.

4.3 Ризики та бар'єри для впровадження

Штучний інтелект і створена нейромережева модель в сфері психіатрії має безумовний потенціал, але для успішного і коректного впровадження в сферу охорони здоров'я є певні виклики, які мають бути вирішені.

1. Упередженість моделей ШІ;

Одним з основних викликів є те, що більшість моделей ШІ навчаються на однорідних даних, і це може обмежувати точність моделі через не врахування

особливостей різних демографічних груп населення. Тож при створенні наборів даних дослідники мають приділяти особливу увагу демографічній, віковій, етнічній, соціальній та економічній складовим.

2. Довіра до ШІ;

Через стрімкий розвиток технологій штучного інтелекту серед суспільства виникають сумніви і недовіра до машин. Головними причинами недовіри є:

- недостатня обізнаність в сфері технологій;
- страх витоку інформації;
- відсутність прозорості системи.

Для успішного використання ШІ створеної нейромережевої моделі розробники та дослідники мають зосередитися на прозорому впровадженні моделі, щоб лікарі та користувачі розуміли, що ШІ не замінить фахівця і на чому базується його робота, а фахівці в той же час мали уявлення про те звідки зібрані дані і як працює система в цілому.

3. Стандарти надійності, точності та безпеки;

Тісна співпраця між дослідниками даних, фахівцями в галузі психічного здоров'я та регуляторними органами матиме вирішальне значення при розробці рекомендацій щодо впровадження ШІ в клінічних умовах. Крім того, інтеграція ШІ з передовими технологіями, такими як нейровізуалізація та геноміка, може зробити революцію в персоналізованій психіатрії. Наприклад, використання ШІ для аналізу генетичних маркерів може поліпшити прогнозування індивідуальних реакцій на психіатричні препарати, мінімізуючи метод проб і помилок при виборі лікування [21].

ВИСНОВКИ

Основним завданням цієї роботи було створення та порівняння різних нейромережових моделей для діагностики психологічних захворювань на основі аналізу тексту, але на однакових даних.

Нами було обрано для порівняння дві класичні моделі – Logistic Regression та Support Vectors Machine і більш просунуту трансформерну модель – DistilBERT і також різні векторизатори, для класичних моделей потужний векторизатор – TD-IDF, а для трансформерної моделі - DistilBertTokenizerFast.

За результатами порівняння було визначено, що не дивлячись на відносно – невеликий для трансформерної моделі даних набір – модель DistilBERT все одно показала кращі результати ніж інші моделі за усіма метриками.

Також було виявлено основні можливості для розвитку моделі, а саме:

- збільшення кількості та покращення якості даних;
- розгортання на серверному обладнанні з використанням GPU та TPU;
- мультимодальність даних – додавання можливості обробки голосу.

Після покращення та допрацювання моделі вона матиме практично необмежений потенціал для впровадження, а саме:

- допоміжний – консультаційний засіб для лікарів;
- веб-платформи та мобільні застосунки для самодіагностики та самопізнання;
- засіб для покращення роботи працівників і покращення емоційного стану на робочому місці;
- аналіз стану користувачів через соціальні мережі.

Підбиваючи підсумки, створена модель має великий потенціал для впровадження в подальшого вдосконалення, особливо якщо система буде достатньо прозорою і зрозумілою для пересічного користувача.

ПЕРЕЛІК ПОСИЛАНЬ

1. Mental disorders. *World Health Organization (WHO)*.
URL: <https://www.who.int/news-room/fact-sheets/detail/mental-disorders>.
2. Chaban, O.s & Khaustova, O. & Asanova, A. & L., Trachuk & Assonov, Dmytro. (2019). Практична психосоматика: діагностичні шкали. Навчальний посібник.
3. Mental Illness: A Review of Causes, Types, and Medicinal Plant Based Therapeutic Interventions / S. Veeramachaneni et al. *Open Access Peer Reviewed Journals | Science and Education Publishing: Home*.
URL: <https://pubs.sciepub.com/rpbs/11/2/2/index.html>.
4. Institute of Health Metrics and Evaluation. Global Health Data Exchange (GHDx). <https://vizhub.healthdata.org/gbd-results/> (Accessed 4 March 2023).
5. Woody CA, Ferrari AJ, Siskind DJ, Whiteford HA, Harris MG. A systematic review and meta-regression of the prevalence and incidence of perinatal depression. *J Affect Disord*. 2017;219:86–92
6. В ЕСОЗ протягом останніх двох років фіксується значне збільшення пацієнтів зі встановленим діагнозом ПТСР. *Національна служба здоров'я України*. URL: <https://nszu.gov.ua/novini/v-esoz-protyagom-ostannih-dvoh-rokiv-fiksuyetsya-znachne-zbi-1203>.
7. Yasar K., Lawton G. What is Data Preprocessing? Key Steps and Techniques. *Search Data Management*.
URL: <https://www.techtarget.com/searchdatamanagement/definition/data-preprocessing>.
8. Що таке обробка природної мови (NLP)? - Oksim. *Oksim - IT допомога*.
URL: https://www.oksim.ua/2023/07/24/shho-take-obrobka-pryrodnoyi-movy-nlp/#Складові_NLP.

9. Afifi, A., Nakaguchi, T., 2015. Unsupervised detection of liver lesions in CT images, in: Proceedings of the Annual International Conference of the IEEE Engineering in Medicine and Biology Society, IEEE Engineering in Medicine and Biology Society. Annual International Conference, 2015, 2411–2414.
- 10.M. Kallenberg, K. Petersen, M. Nielsen, A.Y. Ng, Igel, C. Pengfei Diao, C.M. Vachon, K. Holland, R.R. Winkel, N. Karssemeijer, M. Lillholm Unsupervised deep learning applied to breast density segmentation and mammographic risk scoring IEEE Trans. Med. Imaging, 35 (5) (2016), pp. 1322-1331
- 11.Khyani, Divya & B S, Siddhartha & Niveditha, N. & M., Divya & Y M, Dr. (2021). An Interpretation of Lemmatization and Stemming in Natural Language Processing. Shanghai Ligong Daxue Xuebao/Journal of University of Shanghai for Science and Technology. 22. 350-357.
- 12.Walsh CG, Ribeiro JD, Franklin JC. Predicting risk of suicide attempts over time through machine learning. Clinical Psychological Science. 2017;5(3):457-469.
- 13.Pestian JP, Sorter M, Connolly B, et al; STM Research Group. A machine learning approach to identifying the thought markers of suicidal subjects: a prospective multicenter trial. Suicide Life Threat Behav. 2017;47(1):112-121.
- 14.Bedi G, Carrillo F, Cecchi GA, et al. Automated analysis of free speech predicts psychosis onset in high-risk youths. NPJ Schizophr. 2015;1:15030. doi:10.1038/npjschz.2015.30
- 15.Python Libraries for Natural Language Processing – VCET TechZette – *विसीईटी ज्ञानपत्र. VCET TechZette – विसीईटी ज्ञानपत्र.*
URL: <https://techz.vcet.edu.in/2023/09/12/python-libraries-for-natural-language-processing/>.
- 16.All M. Supervised Machine Learning. Data Camp.
URL: <https://www.datacamp.com/blog/supervised-machine-learning>.

17. Regression vs Classification in Machine Learning - Tpoint Tech. [www.tpointtech.com](https://www.tpointtech.com/regression-vs-classification-in-machine-learning). URL: <https://www.tpointtech.com/regression-vs-classification-in-machine-learning>.
18. Binary and Multiclass Classification in Machine Learning | Analytics Steps. *Analytics Steps - A leading source of Technical & Financial content*. URL: <https://www.analyticssteps.com/blogs/binary-and-multiclass-classification-machine-learning>.
19. Sharma G. Multi-Label Classification. *Medium*. URL: <https://medium.com/analytics-vidhya/multi-label-classification-a9643d221954>.
20. Bassey P. Logistic Regression Vs Support Vector Machines (SVM). *Medium*. URL: <https://medium.com/axum-labs/logistic-regression-vs-support-vector-machines-svm-c335610a3d16>.
21. Distilbert: A Smaller, Faster, and Distilled BERT - Zilliz Learn. *Zilliz: Vector Database built for enterprise-grade AI applications*. URL: <https://zilliz.com/learn/distilbert-distilled-version-of-bert>.
22. Low, D. M., Rumker, L., Torous, J., Cecchi, G., Ghosh, S. S., & Talkar, T. (2020). Natural Language Processing Reveals Vulnerable Mental Health Support Groups and Heightened Health Anxiety on Reddit During COVID-19: Observational Study. *Journal of medical Internet research*, 22(10), e22635.
23. Schwarz E. Artificial Intelligence Applications in Psychiatry - Neurotorium. *Neurotorium*. URL: <https://neurotorium.org/artificial-intelligence-applications-in-psychiatry/>.

ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ (ПРЕЗЕНТАЦІЯ)

КВАЛІФІКАЦІЙНА РОБОТА

на тему: «Нейромережева модель для діагностики психологічних розладів на основі аналізу тексту»

на здобуття освітнього ступеня **бакалавра**
зі спеціальності **126 Інформаційні системи та технології**
освітньо-професійної програми **Інформаційні системи та технології**

Студентки групи ІСД- 41
ДАВИДЕНКО Катерини Олексіївни

Вступ

2

➤ Мета роботи

Метою роботи є дослідження та визначення оптимального підходу до створення нейромережевої моделі для діагностики психологічних захворювань на основі аналізу текстових даних.

➤ Об'єкт дослідження

Процес автоматизованої діагностики психологічних розладів

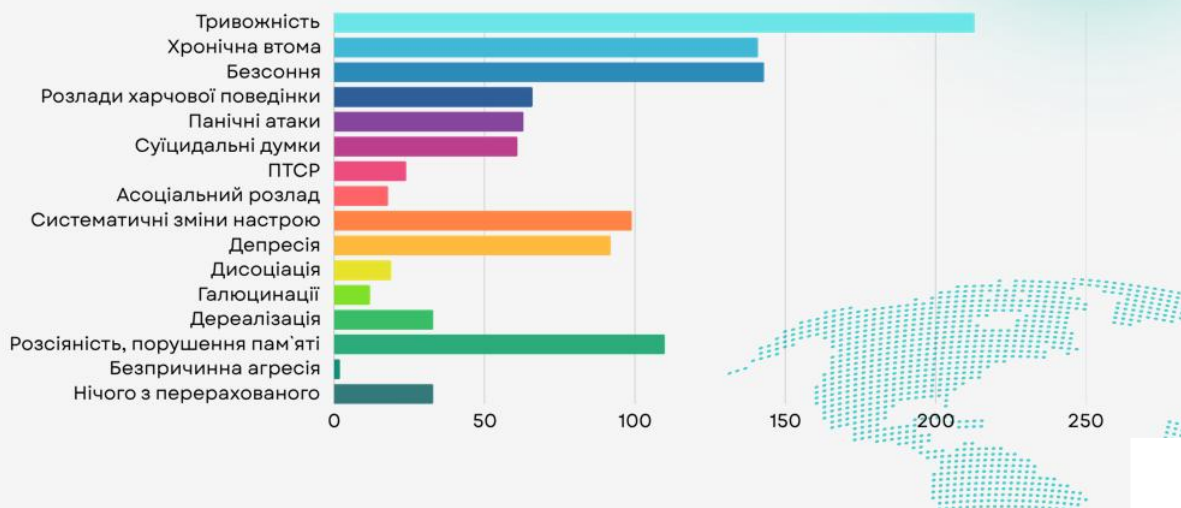
➤ Предмет дослідження

Методи та технології для створення нейромережевої моделі



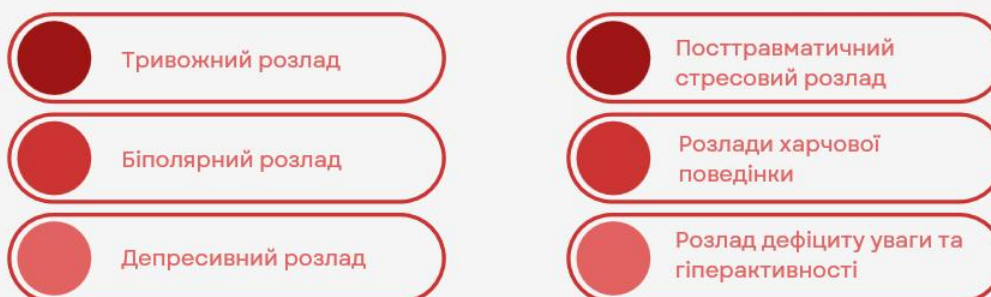
Актуальність теми

3



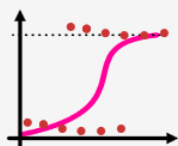
Найбільш поширені психологічні розлади

4



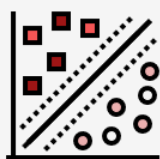
Обрані моделі

5



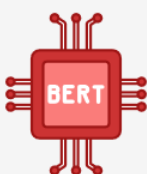
➤ Logistic Regression

Логістична регресія спрямована на вирішення проблем класифікації. Вона робить це, прогнозуючи категоричні результати, на відміну від лінійної регресії, яка прогнозує безперервний результат.



➤ Support Vector Machine (SVM)

Машина опорних векторів (SVM) – є керованим алгоритмом машинного навчання, що використовується для таких задач як: класифікація та регресія.

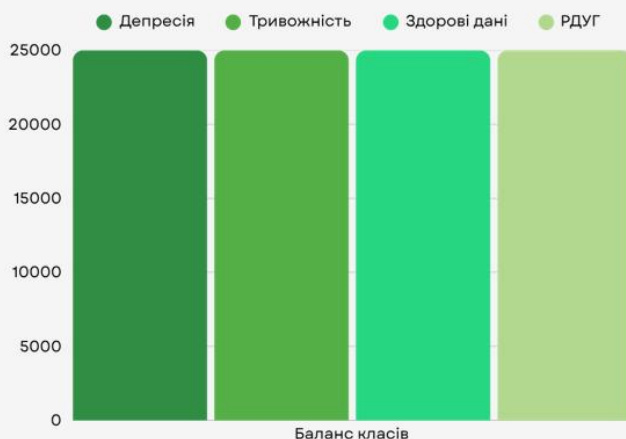


➤ DistilBERT

DistilBERT попередньо тренує меншу і швидшу універсальну модель представлення мови, яку можна точно налаштувати, щоб отримати хорошу продуктивність для багатьох наступних завдань.

Датасет

6



```

data
├── adhd_2018_features_tfidf_256.csv
├── adhd_2019_features_tfidf_256.csv
├── adhd_post_features_tfidf_256.csv
├── adhd_pre_features_tfidf_256.csv
├── anxiety_2018_features_tfidf_256.csv
├── anxiety_2019_features_tfidf_256.csv
├── anxiety_post_features_tfidf_256.csv
├── anxiety_pre_features_tfidf_256.csv
├── depression_2018_features_tfidf_256.csv
├── depression_2019_features_tfidf_256.csv
├── depression_post_features_tfidf_256.csv
├── depression_pre_features_tfidf_256.csv

```

Порівняння базових характеристик

7

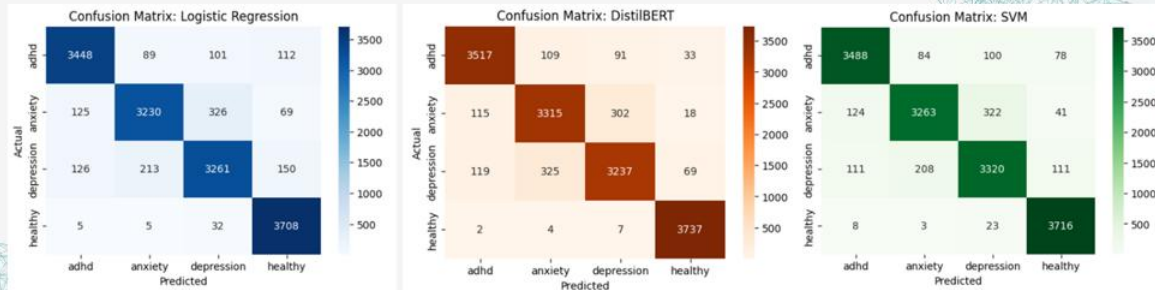
Параметр	Logistic Regression	SVM	DistilBERT
Датасет	100 000 записів	100 000 записів	100 000 записів
Векторизатор/токенізатор	TF-IDF (max 5 000)	TF-IDF (max 5 000)	DistilBertTokenizerFast
Час тренування	≈ 1 хвилина	≈ 2 години 10 хвилин	≈ 50 хв/епоха (7 год всього)
Кількість параметрів	5 000	Непараметрична, SV ≈ 8 000	≈ 66 млн
CPU-навантаження	≈ 60 %	≈ 70 %	≈ 30 %
GPU-навантаження	Немає	Немає	≈ 85 %
Інтерференс (один зразок)	0.5 мс (CPU)	1.0 мс (CPU)	50 мс (CPU) / 10 мс (GPU)
Розмір моделі	~1 MB	~5 MB	~250 MB

Оцінка створених моделей

8

Клас	Метрика	Log. Reg.	SVM	DistilBERT
ADHD	Precision	0.93	0.92	0.94
	Recall	0.9	0.9	0.93
	F1-score	0.91	0.91	0.93
Anxiety	Precision	0.9	0.89	0.88
	Recall	0.85	0.85	0.88
	F1-score	0.87	0.87	0.86
Depression	Precision	0.86	0.85	0.88
	Recall	0.87	0.87	0.86
	F1-score	0.86	0.86	0.87
Healthy	Precision	0.92	0.93	0.97
	Recall	0.99	0.98	1
	F1-score	0.95	0.96	0.98
Загалом	Accuracy	0.9	0.9	0.92

Матриці плутанини



Можливості впровадження нейромережі

- Допоміжний – консультаційний засіб для лікарів
- Системна терапія та самодіагностика
- Аналіз постів та публікацій користувачів у соціальних мережах
- Помічники на роботі та у компаніях

Mood & Mental Health Checker

Welcome! This tool can help you reflect on how you're feeling. You can write freely about your thoughts, or answer some guided questions to describe your mood.

(This is not a medical diagnosis. If you're in crisis, please seek professional help.)

How would you like to share your feelings?

☐ Write my own text.

☒ Answer some guided questions

What's been on your mind lately?

I am really nervous about my diplomas

Have you been feeling more anxious, sad, or restless?

I am feeling nervous and restless

How has your sleep or energy been?

I don't have enough energy to even wake up at the morning

Anything else you'd like to share?

I feel like I have no people to support me

Analyze My Mood

Results:

Logistic Regression: anxiety

Linear SVM: anxiety

DistilBERT: borderline between 'adhd' and 'anxiety' (confidence: 0.48)

This tool is for informational purposes only and is not a substitute for professional diagnosis or treatment. If you're experiencing distress or thoughts of self-harm, please contact a healthcare professional or emergency services.

Бар'єри для впровадження

11

- Упередженість моделей ШІ
- Довіра до ШІ
- Стандарти надійності, точності та безпеки



Висновки

12

- Основним завданням цієї роботи було створення та порівняння різних неймережових моделей для діагностики психологічних захворювань на основі аналізу тексту, але на однакових дата сетах.
- Нами було обрано для порівняння дві класичні моделі – Logistic Regression та Support Vectors Machine і більш просунуту трансформерну модель – DistilBERT і також різні векторизатори, для класичних моделей потужний векторизатор – TD-IDF, а для трансформерної моделі – DistilBertTokenizerFast.
- DistilBERT продемонстрував найвищі показники за всіма метриками, попри відносно невеликий датасет для трансформерної моделі.

Апробація роботи

Теза	Теза	Стаття
Давиденко К. О. <u>“Сучасні підходи застосування штучного інтелекту в психіатрії” та “Перспективи та виклики застосування штучного інтелекту в психіатрії”: матеріали VI Науково-технічної конференції «Сучасний стан та перспективи розвитку ІОТ».</u> Подано до друку.	Давиденко К. О. <u>Застосування штучного інтелекту у сфері охорони здоров'я : матеріали III всеукраїнської науково-технічної конференції «Технологічні горизонти: дослідження та застосування інформаційних технологій для технологічного прогресу України і світу».</u> Надруковано на стр. 147.	Давиденко К. О. <u>Застосування штучного інтелекту в психіатрії: сучасні підходи та перспективи: матеріали журналу «Телекомунікаційні та інформаційні технології».</u> Подано до друку.

Дякую за увагу!

