

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

Спеціальність: 123 «Комп'ютерна інженерія»

Освітньо-професійна програма: «Комп'ютерна інженерія»

Група: 2БКС-29

КВАЛІФІКАЦІЙНА РОБОТА

здобувача освіти денної форми навчання
БКС.29.05.000.КРБ

ГАВРИЛЮКА
ОЛЕГА СЕРГІЙОВИЧА

м. Одеса
2025 р.

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

Спеціальність: **123 «Комп'ютерна інженерія»**

Освітньо-професійна програма: **«Комп'ютерна інженерія»**

Група: **2БКС-29**

ПОЯСНЮВАЛЬНА ЗАПИСКА

До кваліфікаційної роботи бакалавра на тему: **Дослідження алгоритмів роботи нейронних мереж архітектури Transformer**

Проектний матеріал складається з пояснювальної записки на 75 сторінках та графічного (презентаційного) матеріалу на 10 аркушах (слайдах)

Виконавець _____ (Гаврилюк О.С.)

Керівник проекту _____ (Іванова Л.В.)

Консультанти:

з розділу охорони праці та техніки безпеки _____ (Чорновол Н.І.)

з нормоконтролю _____ (Петрашова В.І.)

старший консультант _____ (Кривченко Ю.В.)

До захисту допущений

Завідувач кафедри _____ (Іванова Л.В.)

Завідувач відділення _____ (Краснокутська К.Г.)

Захист «27» 06 2025 р.

Протокол ЕК № 2

Оцінка ЕК 5 (відмінно) / 90

Секретар ЕК _____

АНОТАЦІЯ

Роботу присвячено порівняльному дослідженню LLM-моделей ChatGPT-4o та Gemini 2.5 Pro на основі архітектури Transformer. На базі експериментальних тестів, що включали завдання на обробку контексту, логіки та мультимодальності, було проаналізовано продуктивність моделей.

У ході аналізу виявлено ключові відмінності у їхньому функціонуванні та сильні сторони. Результати узагальнено у вигляді порівняльної характеристики, що може бути корисною для вибору оптимальної моделі та слугувати основою для подальших досліджень.

Розглянуто питання з охорони праці та техніки безпеки.

Ключові слова: LLM, Transformer, ChatGPT-4o, Gemini 2.5 Pro, AI, механізм уваги, порівняльний аналіз.

ANNOTATION

This thesis is dedicated to a comparative study of the LLM models ChatGPT-4o and Gemini 2.5 Pro, based on the Transformer architecture. The performance of the models was analyzed through experimental tests involving tasks that required context processing, logic, and multimodality.

The analysis revealed key differences in their functionality and identified their respective strengths. The results are summarized in a comparative framework, which can be useful for selecting the optimal model and serve as a basis for further research.

Issues of occupational health and safety have been considered.

Keywords: LLM, Transformer, ChatGPT-4o, Gemini 2.5 Pro, AI, attention mechanism, comparative analysis.

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ

ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

Відділення Комп'ютерних систем Кафедра Комп'ютерної інженерії
Спеціальність 123 «Комп'ютерна інженерія»
Освітньо-професійна програма «Комп'ютерна інженерія»

ЗАТВЕРДЖУЮ:

Заст. дир.з НВР Беркань І.В.

“ 28 ” 08 20 25 р.

ЗАВДАННЯ

на кваліфікаційну роботу бакалавра

здобувачеві освіти Гаврилюку Олегу Сергійовичу
(прізвище, ім'я, по батькові)

1. Тема кваліфікаційної роботи Дослідження алгоритмів роботи нейронних мереж архітектури Transformer

затверджена наказом по коледжу від “ 14 ” 11 20 25 р. № 246

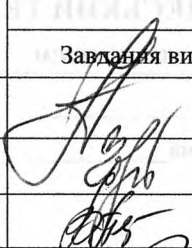
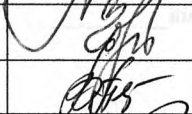
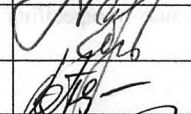
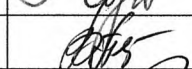
2. Термін здачі студентом кваліфікаційної роботи 15.06.2025

3. Вихідні дані до роботи Мова програмування Python; Редактор VS Code; Нейронні мережі : ChatGPT 4o, Gemini 2.5 Pro Preview 05-06.

4. Зміст розрахунково-пояснювальної записки (перелік питань, що їх належить розробити)
1. Основний розділ; Визначення та історія нейронних мереж; Типи нейронних мереж; Застосування нейронних мереж; Ключові алгоритми архітектури Transformer; Огляд досліджуваних моделей; Аналіз роботи нейронних мереж; 2. Охорона праці та техніка безпеки; Висновки; Перелік використаних інформаційних джерел; Додатки;

5. Перелік графічного матеріалу (слайдів мультимедійної презентації) Титульний слайд; Архітектура Transformer: Основи; Огляд моделі ChatGPT; Огляд моделі Gemini; Методологія дослідження; Аналіз генерації тексту (Creative Writing Coherence); Аналіз машинного перекладу (Translation); Аналіз питань-відповідей (QA); Аналіз стислого переказу (Summarization); Аналіз генерації коду.

6. Консультанти по кваліфікаційній роботі, із зазначенням розділів, що їх стосуються

Розділ	Консультант	ПІДПИС	
		Завдання видав	Завдання прийняв
Основний розділ	Іванова Л.В.		
Розділ охорони праці	Чорновол Н.І.		
Нормоконтроль	Петрашова В.І.		
Старший консультант	Кривченко Ю.В.		

7. Дата видачі завдання

29.04.2025

Керівник роботи Іванова Л.В.

Завдання прийняв до виконання

(підпис)

(підпис)

КАЛЕНДАРНИЙ ПЛАН

Пор. №	Назва етапів кваліфікаційної роботи	Термін виконання етапів роботи	Примітка
1.	Підбір технічної літератури	4.05.2025	виконав
2.	Пошук інтернет-джерел	8.05.2025	виконав
3.	Формування завдання до проєкту	10.05.2025	виконав
4.	Огляд досліджуваних моделей	12.05.2025	виконав
5.	Аналіз роботи нейронних мереж	15.05.2025	виконав
6.	Аналіз генерації тексту	18.05.2025	виконав
7.	Аналіз машинного перекладу	22.05.2025	виконав
8.	Аналіз питань-відповідей	24.05.2025	виконав
9.	Аналіз стислого переказу	26.05.2025	виконав
10.	Аналіз генерації коду	28.05.2025	виконав
11.	Перевірка на вимоги користувача	31.05.2025	виконав
12.	Виконання розділу охорони праці	2.06.2025	виконав
13.	Перевірка роботи на плагіат	6.06.2025	виконав
14.	Підготовка до малого захисту	10.06.2025	виконав
15.	Розробка мультимедійної презентації	12.06.2025	виконав
16.	Підготовка до основного захисту	15.06.2025	виконав
17.	Оформлення мультимедійної презентації	18.06.2025	виконав

Здобувач освіти

(підпис)

Керівник роботи

(підпис)

ЗМІСТ

Вступ.....	8
1 Основний розділ.....	10
1.1 Визначення та історія нейронних мереж.....	10
1.2 Типи нейронних мереж.....	11
1.3 Застосування нейронних мереж.....	13
1.4 Ключові алгоритми архітектури Transformer.....	15
1.5 Огляд досліджуваних моделей.....	18
1.5.1 Огляд моделі ChatGPT.....	18
1.5.2 Огляд моделі Gemini.....	20
1.6 Аналіз роботи нейронних мереж.....	22
1.6.1 Аналіз генерації тексту (Creative Writing Coherence).....	23
1.6.2 Аналіз машинного перекладу (Translation).....	31
1.6.3 Аналіз питань-відповідей (QA).....	37
1.6.4 Аналіз стислого переказу (Summarization).....	42
1.6.5 Аналіз генерації коду.....	49
2 Розділ охорони праці та техніки безпеки.....	54
2.1 Аналіз небезпечних і шкідливих факторів, що впливають на програміста під час розробки системи.....	54
2.2 Розробка заходів з охорони праці.....	54
2.2.1 Виробниче приміщення.....	54
2.2.2 Мікроклімат робочої зони працівників, вентиляція.....	55
2.2.3 Освітлення робочого місця, шум.....	56
2.2.4 Електробезпека.....	57
2.2.5 Організація робочого місця користувача ПК.....	57
2.3 Пожежна безпека при роботі з ВДТ.....	58
Висновки.....	59
Перелік використаних інформаційних джерел.....	60
Додаток А. Фрагменти коду на мові Python для аналізу тестів.....	61
Додаток Б. Слайди мультимедійної презентації.....	71

ВСТУП

Сучасний світ переживає епоху стрімкого розвитку штучного інтелекту, де великі мовні моделі (LLM) відіграють ключову роль, трансформуючи численні галузі – від обробки природної мови та генерації контенту до наукових досліджень та взаємодії людини з комп'ютером. В основі переважної більшості найуспішніших LLM лежить архітектура Transformer, представлена у 2017 році. Розуміння фундаментальних алгоритмів цієї архітектури, зокрема механізму уваги (attention mechanism), є критично важливим для ефективного використання, оцінки можливостей та обмежень, а також для подальшого вдосконалення цих складних систем. Поява нових поколінь моделей, таких як ChatGPT-4o від OpenAI та Gemini 2.5 Pro від Google, з розширеними можливостями (покращена мультимодальність, обробка надвеликих контекстів), підкреслює необхідність детального дослідження їхніх внутрішніх механізмів, навіть в умовах обмеженого доступу до повних архітектурних деталей. Аналіз та порівняння роботи цих передових моделей дозволяє глибше зрозуміти практичну реалізацію та еволюцію алгоритмів Transformer.

Мета роботи полягає в дослідженні та порівняльному аналізі ключових алгоритмів роботи нейронних мереж архітектури Transformer на прикладі функціонування сучасних великих мовних моделей ChatGPT-4o та Gemini 2.5 Pro для виявлення їхніх характерних особливостей та відмінностей у вирішенні типових завдань.

Для досягнення поставленої мети необхідно вирішити такі завдання дослідження:

1. Проаналізувати теоретичні основи та ключові алгоритмічні компоненти архітектури Transformer.
2. Розглянути відомі модифікації та особливості реалізації Transformer у сучасних LLM.
3. Систематизувати доступну інформацію про моделі ChatGPT-4o та Gemini 2.5 Pro.

					БКС 29. 05 000. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		8

4. Розробити методологію та набір тестових сценаріїв для порівняльного аналізу поведінки моделей, що відображає роботу базових алгоритмів.

5. Провести експериментальне порівняння роботи ChatGPT-4o та Gemini 2.5 Pro на розроблених завданнях.

6. Проаналізувати отримані результати, зробити висновки щодо ймовірних відмінностей у функціонуванні алгоритмів та їх впливу на продуктивність моделей.

Об'єктом дослідження є алгоритми роботи нейронних мереж, побудованих на архітектурі Transformer. Предметом дослідження є особливості реалізації та функціональні характеристики алгоритмів Transformer у великих мовних моделях ChatGPT-4o та Gemini 2.5 Pro, виявлені шляхом аналізу їхньої поведінки та порівняння результатів роботи.

					БКС 29. 05 000. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		9

1 ОСНОВНИЙ РОЗДІЛ

1.1 Визначення та історія нейронних мереж

Нейронні мережі - це тип алгоритму машинного навчання, який базується на структурі та функціях людського мозку. Вони складаються з взаємопов'язаних вузлів, або нейронів, які організовані в шари. Кожен нейрон отримує вхідні дані від інших нейронів і виробляє вихідний сигнал, який передається іншим нейронам наступного шару. Цей процес триває доти, доки останній шар не створить вихід мережі. Нейронні мережі навчаються за допомогою процесу, що називається зворотним поширенням, який коригує ваги зв'язків між нейронами, щоб мінімізувати похибку між прогнозованим виходом і фактичним виходом.

Ідея штучних нейронних мереж виникла в 1940-х роках, коли Уоррен Маккалох і Волтер Піттс запропонували математичну модель нейрона, яка могла обчислювати логічні функції. Вони показали, що мережа цих штучних нейронів може бути використана для обчислення будь-якої булевої функції, а отже, може слугувати універсальною обчислювальною машиною. Однак реалізація нейронних мереж була обмежена тогочасними обчислювальними технологіями.

У 1950-х і 1960-х роках кілька дослідників, включаючи Френка Розенблата, розробили тип нейронної мережі під назвою персептрон. Персептрон - це тип нейронної мережі прямого поширення, яка складається з одного шару нейронів. Його використовували для вирішення простих завдань класифікації, таких як розпізнавання рукописних цифр. Однак персептрон був обмежений у своїх можливостях і не міг вирішувати більш складні завдання.

Нейронні мережі стали популярною темою досліджень лише у 1980-х роках, коли Девід Румельхарт, Джеффри Хінтон і Рональд Вільямс розробили алгоритм зворотного розповсюдження (backpropagation). Зворотне поширення - це алгоритм керованого навчання, який регулює ваги зв'язків між нейронами, щоб мінімізувати похибку між прогнозованим і фактичним виходом. Це дозволило нейронним мережам вирішувати більш складні завдання, такі як розпізнавання мови та обробка природної мови. У 1990-х роках нейронні мережі втратили популярність,

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						10
Зм.	Арк.	№ докум.	Підпис	Дата		

оскільки інші алгоритми машинного навчання, такі як машини опорних векторів і дерева рішень, стали більш популярними. Частково це було пов'язано зі складністю навчання глибоких нейронних мереж, які мають багато шарів і важко піддаються оптимізації. Однак розробка нових алгоритмів, таких як згорткові нейронні мережі та мережі глибокого переконання, уможливила навчання глибоких нейронних мереж і призвела до відродження інтересу до нейронних мереж в останні роки.

Сьогодні нейронні мережі використовуються в широкому спектрі застосувань, включаючи комп'ютерний зір, обробку природної мови та ігри. Вони досягли найсучаснішої продуктивності на багатьох еталонних наборах даних і були використані для вирішення раніше нерозв'язних проблем. Однак такі проблеми, як інтерпретованість, масштабованість і надійність, залишаються, і для їх вирішення та розкриття повного потенціалу нейронних мереж необхідні подальші дослідження.

1.2 Типи нейронних мереж

Нейронні мережі бувають різних типів, кожен з яких має свою унікальну архітектуру та алгоритм навчання. Деякі з найпоширеніших типів нейронних мереж включають нейронні мережі прямого поширення, 8 рекурентні нейронні мережі, згорткові нейронні мережі та мережі глибокого переконання.

Прямі нейронні мережі. Нейронні мережі прямого поширення - це найпростіший і найпоширеніший тип нейронних мереж. Вони складаються з вхідного шару, одного або декількох прихованих шарів і вихідного шару. Вхідні дані проходять через приховані шари для отримання кінцевого результату. Нейронні мережі прямого поширення використовуються для таких завдань, як класифікація, регресія та розпізнавання образів.

Рекурентні нейронні мережі. Рекурентні нейронні мережі призначені для обробки послідовностей вхідних даних. Вони мають зв'язки між нейронами, які утворюють спрямований цикл, що дозволяє мережі зберігати інформацію в часі. Рекурентні нейронні мережі використовуються для таких завдань, як

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		11

розпізнавання мови, машинний переклад і підписи до зображень.

Згорткові нейронні мережі. Згорткові нейронні мережі призначені для обробки даних з сітчастою топологією, таких як зображення. Вони використовують згорткові шари для вилучення ознак із вхідних даних і об'єднання шарів для зменшення вибірки карт ознак. Згорткові нейронні мережі використовуються для таких завдань, як класифікація зображень, виявлення об'єктів і семантична сегментація.

Мережі глибокого переконання. Мережі глибокого переконання складаються з декількох шарів обмежених машин Больцмана. Вони використовуються для неконтрольованого навчання та вилучення ознак, а також можуть бути використані для попереднього навчання ваг нейронної мережі для завдання контрольованого навчання. Мережі глибокого переконання використовуються для таких завдань, як розпізнавання мови, обробка природної мови та комп'ютерний зір.

Автокодери. Автокодери - це тип нейронних мереж, які можна використовувати для неконтрольованого навчання та вилучення ознак. Вони складаються з кодера, який відображає вхідні дані в низьковимірне представлення, і декодера, який реконструює вхідні дані з низьковимірного представлення. Автокодери використовуються для таких завдань, як згладжування зображень, зменшення розмірності та виявлення аномалій.

Генеративні змагальні мережі. Генеративні змагальні мережі - це тип нейронної мережі, яка складається з двох мереж: генератора і дискримінатора. Генератор генерує фальшиві зразки, а дискримінатор намагається відрізнити справжні зразки від фальшивих. Обидві мережі навчаються разом, генератор намагається обдурити дискримінатор, а дискримінатор намагається точно класифікувати зразки. Генеративні змагальні мережі використовуються для таких завдань, як синтез зображень, генерація тексту та відео.

Загалом, існує багато різних типів нейронних мереж, кожен з яких має свої унікальні сильні та слабкі сторони. Вибір типу нейронної мережі залежить від конкретного завдання і характеристик даних. Вибір правильного типу нейронної

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						12
Зм.	Арк.	№ докум.	Підпис	Дата		

мережі може значно покращити продуктивність моделі та дозволити їй вирішувати складні завдання.

1.3 Застосування нейронних мереж

Нейронні мережі використовуються в широкому спектрі застосувань, що охоплюють різні сфери, і досягли найсучаснішої продуктивності у виконанні багатьох завдань. Деякі з найпоширеніших застосувань нейронних мереж включають

Комп'ютерний зір: Нейронні мережі широко використовуються в комп'ютерному зорі, який передбачає аналіз і розуміння візуальних 10 даних, таких як зображення і відео. Вони використовуються для таких завдань, як класифікація зображень, виявлення об'єктів і семантична сегментація. Нейронні мережі досягли найсучаснішої продуктивності на багатьох еталонних наборах даних і використовуються в реальних додатках, таких як безпілотні автомобілі, системи безпеки та медична візуалізація.

Обробка природної мови: Нейронні мережі використовуються в обробці природної мови, яка передбачає аналіз і розуміння людської мови. Вони використовуються для таких завдань, як машинний переклад, аналіз настроїв і розпізнавання мови. Нейронні мережі досягли найсучаснішої продуктивності на багатьох еталонних наборах даних і використовуються в реальних додатках, таких як віртуальні асистенти, чат-боти і мовні моделі.

Ігри: Нейронні мережі використовуються для розробки агентів, які можуть грати в ігри на надлюдському рівні. Їх використовували для гри в такі ігри, як шахи, го і покер, і вони перемагали чемпіонів світу в кожній з цих ігор. Нейронні мережі також використовувалися для розробки агентів, які можуть грати у відеоігри, і досягли найсучасніших результатів у багатьох іграх.

Фінанси: Нейронні мережі використовуються у фінансах для таких завдань, як прогнозування цін на акції, виявлення шахрайства та оптимізація портфеля. Їх використовували для розробки моделей, які можуть передбачати ціни на акції з високою точністю, і застосовували в реальних додатках, таких як алгоритмічна

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		13

торгівля та управління ризиками.

Охорона здоров'я: Нейронні мережі використовуються в охороні здоров'я для таких завдань, як діагностика, планування лікування та пошук ліків. Вони були використані для розробки моделей, які можуть діагностувати захворювання на основі медичних зображень з високою точністю, і були використані в реальних додатках, таких як виявлення раку і персоналізована медицина.

Транспорт: Нейронні мережі використовуються в транспорті для вирішення таких завдань, як автономне водіння, прогнозування трафіку та оптимізація маршрутів. Вони використовуються для розробки моделей, які можуть прогнозувати затори на дорогах і оптимізувати маршрути для мінімізації часу в дорозі. Нейронні мережі також використовуються в автономних транспортних засобах, дозволяючи їм сприймати навколишнє середовище і орієнтуватися в ньому в режимі реального часу.

Безпека: Нейронні мережі використовуються у сфері безпеки для таких завдань, як виявлення вторгнень, аналіз загроз і виявлення шкідливого програмного забезпечення. Вони були використані для розробки моделей, які можуть виявляти зловмисну активність у комп'ютерних мережах і виявляти вразливості в програмному забезпеченні. Нейронні мережі також використовуються в біометрії, дозволяючи розпізнавати людей за їхніми рисами обличчя або відбитками пальців.

Загалом, нейронні мережі використовуються в широкому спектрі застосувань і досягли найсучасніших показників у виконанні багатьох завдань. Вони мають потенціал для трансформації багатьох галузей, а їхній постійний розвиток і вдосконалення, як очікується, призведе до ще більших проривів у майбутньому. Однак залишаються такі проблеми, як інтерпретованість, масштабованість і надійність, і для їх вирішення та розкриття повного потенціалу нейронних мереж необхідні подальші дослідження.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						14
Зм.	Арк.	№ докум.	Підпис	Дата		

1.4 Ключові алгоритми архітектури Transformer

Механізм уваги та архітектура «Трансформер». Багато моделей, що створені для моделювання послідовності в послідовність використовують архітектуру кодер-декодер. Де кодер, відображає вхідну послідовність токенів ($x_1, x_2 \dots x_n$) в послідовність $z = (z_1, z_2 \dots z_n)$. Отримавши на вхід z декодер потім генерує послідовність ($y_1, y_2 \dots y_n$) для кожного елемента з x .

На кожному кроці модель використовує властивість авторегресивності – приймає на вхід створені попередніми кроками виходи, для генерації наступних. Трансформер повністю використовує описану архітектуру (Рис. 1.1), використовуючи механізм уваги та по елементні операції як для кодера, так і для декодера.

Кодер складається з шести однакових шарів, що йдуть послідовно один за одним. Кожен з великих шарів, містить в собі дві частини. Перша – це механізм «само-уваги» з декількома головами(гілками). Друга – проста нейронна мережа прямого поширення, з властивістю «fully-connected», тобто кожен нейрон з'єднаний зі всіма нейронами наступного шару. Також використовується залишковий зв'язок, яким обгорнуті обидві ці частини, після якого виконується нормалізація. Тобто результат кожної частини:

$$Res = LayerNorm (x + Sublayer (x)) \quad (1.1)$$

де $Sublayer(x)$ – функція реалізована всередині кожної частини. Для того щоб було менше проблем з розмірністю, всі залишкові з'єднання, а також внутрішні частини кодера мають виходи однакової розмірності – 512.

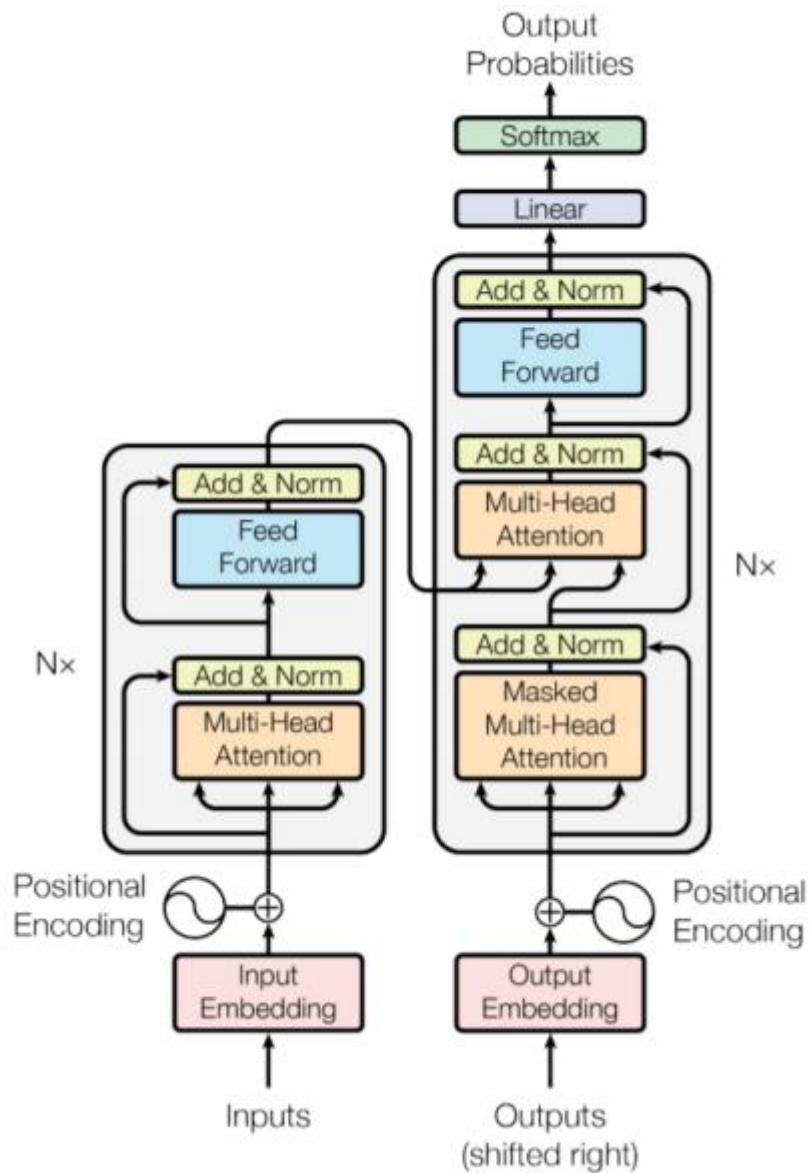


Рисунок 1.1. Оригінальна архітектура «Трансформеру»

Декодер має стільки ж однакових шарів, як і кодер – 6. Також він використовує два шари кодер, що описані вище, без змін. Але декодер містить в собі третій шар – головна функція якого механізм уваги з багатьма головами (Multi Head Attention). Подібно до кодера, також використовуються залишкові з'єднання, навколо кожної з трьох частин. Також необхідний крок для роботи 34 декодера – зробити здвиг вправо на одну позицію, для того, щоб гарантувати, що прогнози зроблені для позиції будуть залежати лише від всіх попередніх позицій, а не від поточної.

Механізм уваги можна описати як маппінг(встановлення відповідності) запиту, та набору пар ключ-значення, де всі три елементи представлені у

векторному вигляді. Результат обчислюється як зважена сума, де вага, присвоєна кожному значенню, обчислюється функцією сумісності між запитом, та відповідним йому ключем.

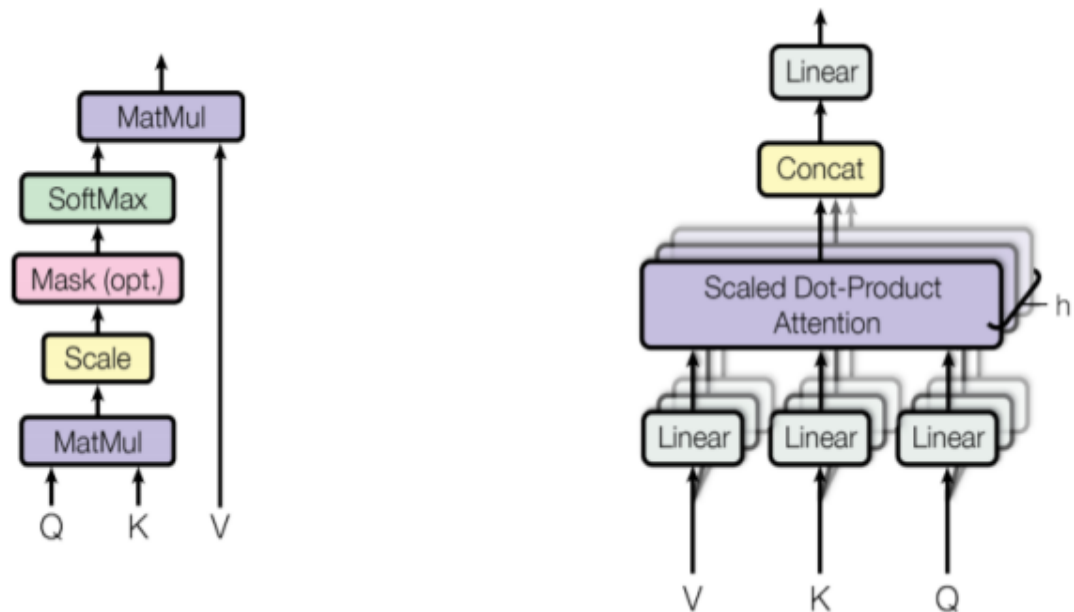


Рисунок 1.2. Скалярний механізм уваги(зліва) та механізм уваги з багатьма головами(справа)

Всього у «Трансформері» використовується два види механізму уваги – масштабований скалярний механізм уваги, та механізм уваги з декількома головами (Рис. 1.2). Масштабний скалярний механізм уваги приймає на вхід запити та ключі розмірності dk та значення розмірності dv . Виконується операція скалярного добутку між запитами та ключами, і нормується діленням на корінь з dk . Після цього застосовується SoftMax функція, щоб отримати ваги для кожного 35 значення. На практиці використовується не послідовне обчислення, а оперування матрицями, за формулою:

$$Attention(Q, K, V) = softmax(Q * K^T \sqrt{dk}) * V \quad (1.2)$$

де Q, K, V- запити, ключі і значення відповідно, у матричному вигляді;

Механізм уваги з багатьма головами створений для того, щоб використовувати одну функцію уваги паралельно, з різними вивченими лінійними проекціями. Після паралельного обчислення всіх функцій, вони об'єднуються в одне векторне представлення, і ще раз проектується в потрібну нам розмірність.

Архітектури нейронних мереж, що використовують механізм уваги. Всі моделі, які використовують механізм уваги, можна поділити на чотири групи. Модель може відноситись до однієї, або навіть декількох груп, але для простоти відображення будемо відносити модель до групи, якій вона найбільше відповідає. Всього існує три групи моделей: авторегресивні, автокодувальні, і моделі типу «послідовність в послідовність».

1.5 Огляд досліджуваних моделей

1.5.1 Огляд моделі ChatGPT

Лінійка мовних моделей GPT (Generative Pre-trained Transformer) була започаткована компанією OpenAI у 2018 році з випуском GPT-1. Кожне наступне покоління моделей демонструвало помітне збільшення обчислювальної потужності, обсягу параметрів і якості генерації природної мови.

GPT-4, випущена у березні 2023 року, стала наступником успішної моделі GPT-3.5. Новітня версія GPT-4-turbo, презентована в листопаді 2023 року, є вдосконаленою та оптимізованою реалізацією GPT-4, орієнтованою на вищу ефективність генерації, зниження вартості обчислень і розширення функціональних можливостей, зокрема у сфері мультимодальності (текст, зображення, аудіо). Розвиток GPT-4 значною мірою зосереджувався на підвищенні стабільності, якості діалогової взаємодії та можливості обробки великих обсягів контексту, що дозволяє моделі ефективно працювати із довгими документами та складними запитам.

GPT-4 базується на класичній архітектурі Transformer, однак має ряд внутрішніх вдосконалень, про які OpenAI частково не розкриває специфіку з міркувань безпеки та комерційної таємниці. Відомо, що в основі GPT-4 лежить багатомодельний стек експертних моделей (Mixture of Experts), що дозволяє активувати лише частину параметрів залежно від вхідного запиту, оптимізуючи продуктивність та зменшуючи обчислювальні витрати.

Значною відмінністю GPT-4 є здатність до більш тонкого опрацювання запитів, завдяки вдосконаленому інструктивному навчанню (Instruction tuning) та

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		18

застосуванню reinforcement learning with human feedback (RLHF). Це дозволило підвищити релевантність відповідей, логічну узгодженість та здатність до критичного аналізу даних.

Модель GPT-4o вирізняється низкою технічних характеристик, що забезпечують її ефективне застосування в аналітичних та експериментальних задачах:

- Максимальна довжина контексту — до 128 000 токенів.
- Мультимодальна підтримка — текст, зображення, аудіо, відео (в GPT-4o).
- Підтримка файлів — обробка PDF, DOCX, CSV, XLSX тощо.
- Опціональна пам'ять — адаптація до користувача на основі збережених

даних.

Модель GPT-4o реалізована в середовищі ChatGPT з підтримкою таких інструментів:

- Кодовий інтерпретатор
- Аналіз документів та зображень
- Голосова взаємодія (GPT-4o)
- Інтеграція з веб-пошуком
- Обробки природної мови (NLP);
- Оцінювання продуктивності LLM (BLEU, ROUGE, perplexity);
- Дослідження QA-систем, генерації, перекладу;
- Порівняльного аналізу моделей на benchmark-наборах (ARC, MMLU, GSM8K, HellaSwag).

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		19

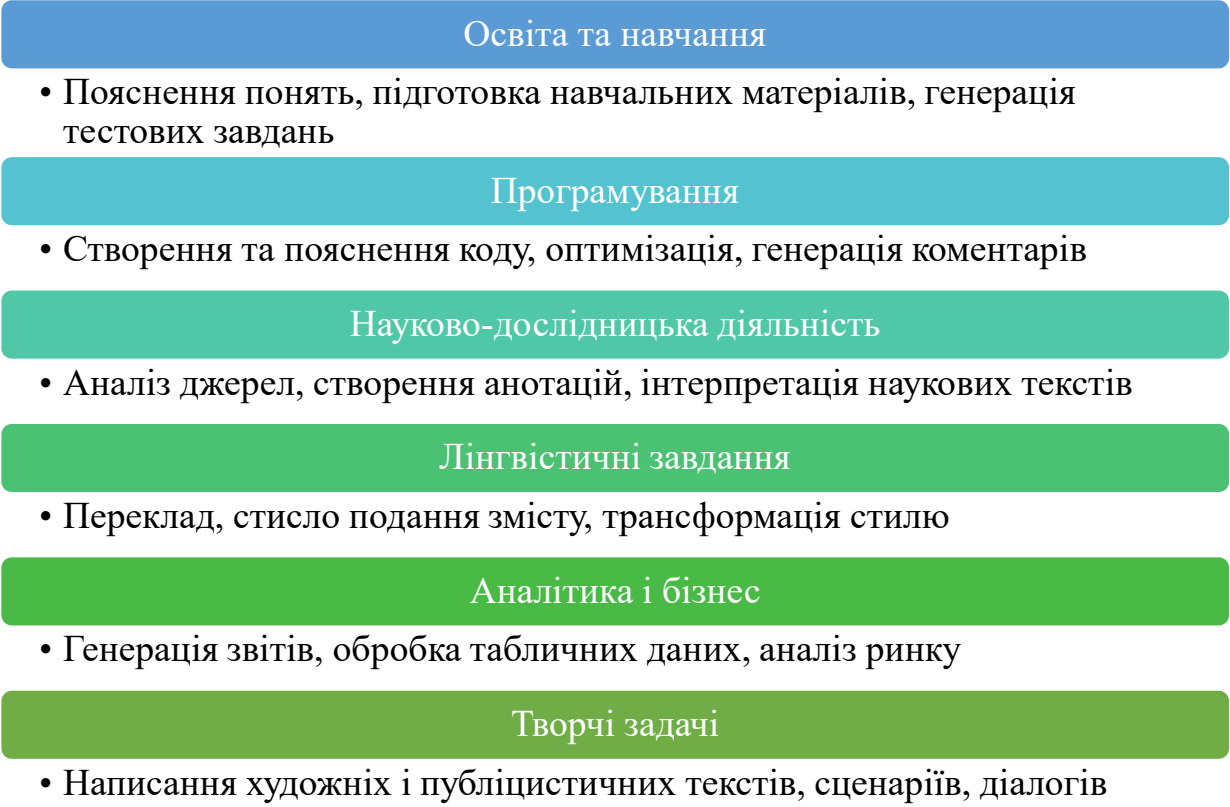


Рисунок 1.3. Функціональне застосування ChatGPT

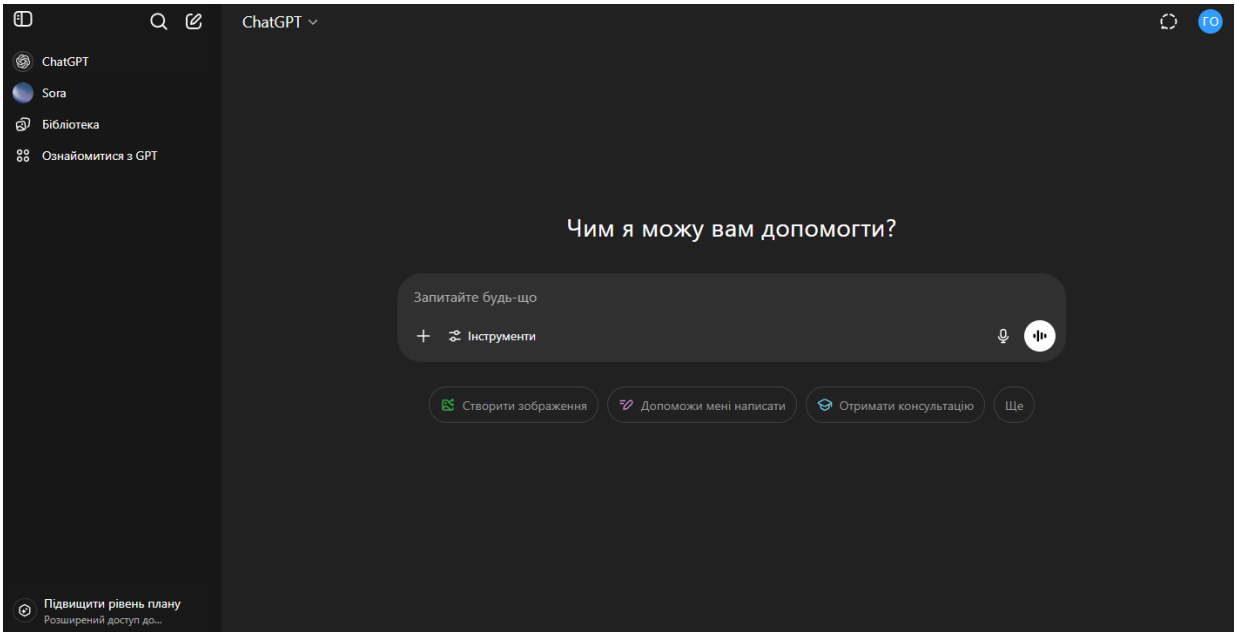


Рисунок 1.4 Огляд сторінки ChatGPT

1.5.2 Огляд моделі Gemini

Gemini 2.5 Pro Preview 05-06— це потужна велика мовна модель (LLM), розроблена компанією Google DeepMind, яка є частиною лінійки моделей Gemini. Вона належить до класу мультимодальних моделей, здатних працювати з

текстовими, зображувальними, аудіо та іншими видами даних. Gemini 2.5 Pro Preview 05-06 позиціонується як універсальний асистент для розв'язання складних задач у галузях штучного інтелекту, машинного навчання, наукових досліджень, а також розробки програмного забезпечення.

Основні технічні характеристики моделі Gemini 2.5 Pro Preview 05-06 включають:

- Максимальна довжина контексту — до 1 000 000 токенів.
- Мультимодальна підтримка — текст, зображення, аудіо.
- Підтримка коду — здатність генерувати, доповнювати й пояснювати код різними мовами програмування.
- Інтеграція в екосистему Google — використовується у Workspace, Chrome, Search, Android тощо.

Gemini 2.5 Pro Preview 05-06 інтегрована в продукти компанії Google через інтерфейси:

- Gemini в Workspace (Docs, Sheets, Slides);
- Розширення в Chrome та Android;
- Gemini Advanced — платна версія з доступом до новітніх моделей;
- API через Google AI Studio.

Gemini 2.5 Pro Preview 05-06 використовується для проведення сучасних досліджень у сфері штучного інтелекту та машинного навчання. Модель застосовується у задачах генерації текстів, аналізу даних, обробки зображень, побудови діалогових агентів та порівняльних benchmark-тестів.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						21
Зм.	Арк.	№ докум.	Підпис	Дата		

Освіта та навчання

- Допомога з теоретичним матеріалом, пояснення концепцій, створення навчальних ресурсів

Програмування

- Генерація коду, рефакторинг, пояснення функцій, тестування

Наукові дослідження

- Аналіз публікацій, формування гіпотез, побудова висновків

Мультимодальна обробка

- Розуміння вхідних зображень, діаграм, аудіо

Інтеграція з сервісами

- Автоматизація через Google-додатки (Docs, Sheets, Gmail)

Рисунок 1.5. Функціональне застосування Gemini

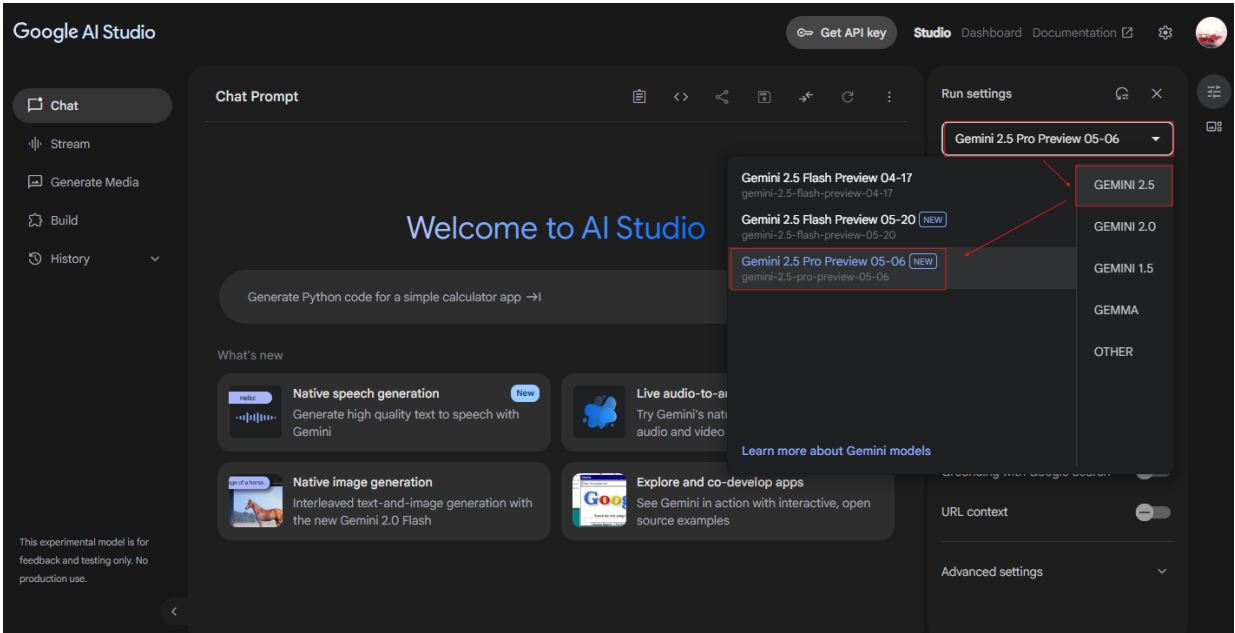


Рисунок 1.6. Огляд сторінки Gemini

1.6 Аналіз роботи нейронних мереж

Основна мета проведення тестів — не просто отримати набір відповідей від моделей, а глибше зрозуміти та емпірично дослідити, як саме функціонують ключові алгоритми архітектури Transformer (механізм уваги, обробка довгого

контексту, позиційне кодування тощо) на прикладі обраних передових великих мовних моделей (ChatGPT-4o, Gemini 2.5 Pro Preview 05-06). Оскільки ми не маємо прямого доступу до вихідного коду цих комерційних моделей та їхньої повної детальної архітектури, тестування стає основним інструментом аналізу їхньої поведінки "зовні". Цей аналіз дозволяє робити обґрунтовані припущення про внутрішні механізми, їхню ефективність, особливості реалізації та потенційні відмінності між моделями.

1.6.1 Аналіз генерації тексту (Creative Writing Coherence)

Мета: Оцінити, наскільки добре LLM може імітувати людську здатність до творчого та осмисленого письма, що є однією з найскладніших задач для штучного інтелекту.

Методологія рахування/інтерпретування метрик:

- Coherence (зв'язність):

Метод: Людська оцінка (Human Evaluation). Оцінює, наскільки текст логічно послідовний, чи плавно переходять думки, чи всі частини тексту стосуються заявленої теми (для есе) або сюжетної лінії (для оповідання).

Шкала: Можна використовувати шкалу Лайкерта, наприклад, від 1 (зовсім не зв'язно) до 5 (дуже зв'язно).

Інтерпретація: Вищий бал = краща зв'язність.

- Creativity (оригінальність):

Метод: Людська оцінка. Оцінює наявність нових ідей, нешаблонних поворотів сюжету (для оповідання), свіжих поглядів (для есе), уникнення кліше.

Шкала: Від 1 (дуже шаблонно, неоригінально) до 5 (дуже креативно, оригінально).

Інтерпретація: Вищий бал = більша оригінальність.

- Fluency (граматика, стилістика):

Метод: Людська оцінка. Оцінює граматичну правильність, відсутність орфографічних та пунктуаційних помилок, природність мови, відповідність

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						23
Зм.	Арк.	№ докум.	Підпис	Дата		

стилю (для есе – науково-популярний/публіцистичний, для оповідання – художній).

Шкала: Від 1 (багато помилок, мова неприродна) до 5 (бездоганна граматики та стиль).

Інтерпретація: Вищий бал = краща якість мови.

- Perplexity (складність/природність тексту):

Метод: Це автоматична метрика, яка оцінює, наскільки добре мовна модель (не та, що генерувала, а інша, *оціночна* модель) може передбачити даний текст. Чим нижча perplexity, тим "природнішим" або "ймовірнішим" є текст з точки зору оціночної моделі.

Розрахунок: Потребує натренованої мовної моделі (напр., GPT-2 з бібліотеки Hugging Face Transformers, на мові програмування Python у редакторі VS Code) і скрипту для розрахунку. Ви подаєте згенерований текст на вхід цій оціночній моделі (рис.1.7).

$$Perplexity = \exp(-1/N * \sum \log P(w_i | w_1, ..., w_{i-1})) \quad (1.3)$$

де N - кількість tokenів, P - ймовірність токена.

```
perplexity.py > ...
1 import torch
2 from transformers import AutoModelForCausalLM, AutoTokenizer
3
4 MODEL_NAME = "gpt2"
5
6 try:
7     tokenizer = AutoTokenizer.from_pretrained(MODEL_NAME)
8     model = AutoModelForCausalLM.from_pretrained(MODEL_NAME)
9 except Exception as e:
10     print(f"Помилка завантаження моделі або токенизатора: {e}")
11     print(f"Переконайтеся, що ви маєте інтернет-з'єднання з назвою моделі '{MODEL_NAME}' правильна.")
12     exit()
13
14 device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
15 model.to(device)
16 model.eval()
17
18 def calculate_perplexity(text, model, tokenizer, device):
19
20     if not text.strip():
21         print("Попередження: Отримано порожній текст. Perplexity не може бути розрахована.")
22         return float('inf')
23
24     inputs = tokenizer(text, return_tensors="pt", truncation=True, max_length=tokenizer.model_max_length)
25     input_ids = inputs.input_ids.to(device)
26     attention_mask = inputs.attention_mask.to(device)
```

Рисунок 1.7. Оціночна модель [gpt2]

Інтерпретація: Нижчий показник perplexity вказує на більш "гладкий" і передбачуваний (в хорошому сенсі, тобто природний) текст.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						24
Зм.	Арк.	№ докум.	Підпис	Дата		

Вибір промпту для завдання.

Промпт1: Напиши коротке есе (100–150 слів) на тему: «Як штучний інтелект змінить освіту в найближчі 10 років».

1. Результат введення промпту1 до ChatGPT в рис. 1.8.

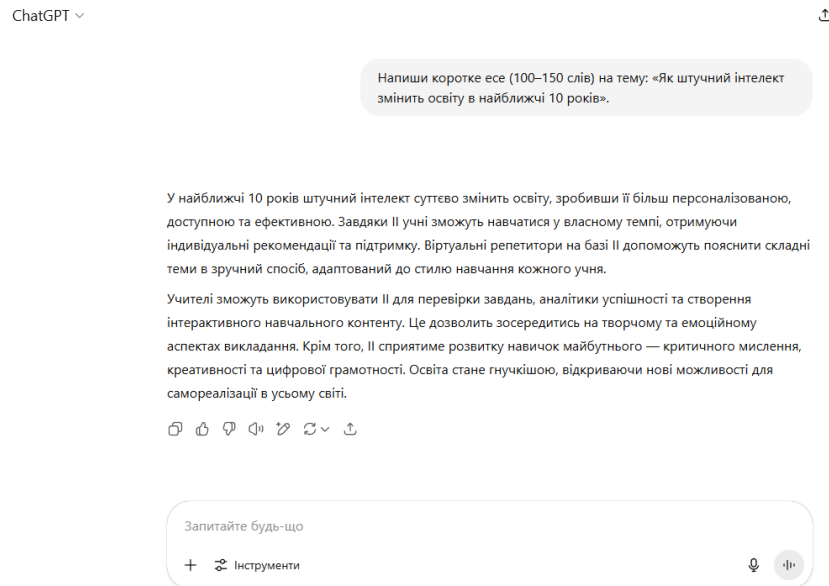


Рисунок 1.8. Результат промпту1 в ChatGPT

Аналіз згенерованого тесту ChatGPT по заданим метрикам.

Coherence (зв'язність): текст має чітку логічну структуру. Вступне речення задає тон. Перший абзац фокусується на перевагах для учнів (персоналізація, темп, віртуальні репетитори). Другий абзац — на перевагах для вчителів (автоматизація, фокус на творчості) та загальних наслідках (навички майбутнього, гнучкість). Переходи між думками плавні.

Creativity (оригінальність): ідеї, викладені в есе (персоналізація, віртуальні репетитори, автоматизація для вчителів, розвиток навичок майбутнього), є досить поширеними та очікуваними в дискусіях про ШІ в освіті. Текст не пропонує кардинально нових або несподіваних поглядів, але дуже якісно узагальнює ключові прогнози. Оригінальність полягає скоріше в чіткому та лаконічному викладі цих ідей.

Fluency (граматика, стилістика): текст написаний бездоганною українською мовою. Граматика, орфографія та пунктуація на високому рівні. Речення добре

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		25

побудовані, мова природна та легка для сприйняття. Стиль відповідає науково-популярному есе – інформативний, але доступний.

Perplexity (складність/природність тексту): низький показник PPL свідчить про високу "природність" та статистичну ймовірність тексту з точки зору оціночної моделі [gpt2]. Це, ймовірно, зумовлено використанням чітких, стандартних мовних конструкцій та поширеної лексики, характерної для обговорення теми ШІ в освіті, що робить текст легко "передбачуваним" для оціночної моделі (рис. 1.9).

Perplexity для есе від ChatGPT-4o: 17.58

Рисунок 1.9. Результат Perplexity для ChatGPT

2. Результат введення промпту1 до Gemini в рис. 1.10.

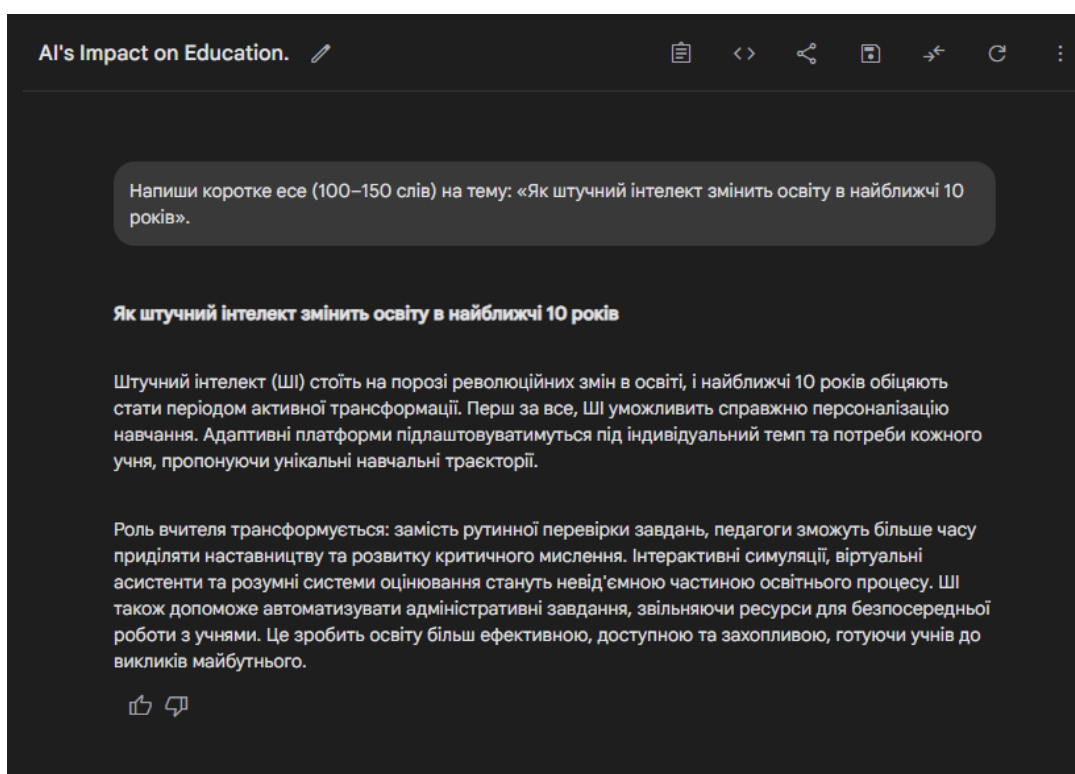


Рисунок 1.10. Відповідь промпту1 в Gemini

Аналіз згенерованого тесту Gemini по заданим метрикам.

Coherence (зв'язність): текст добре структурований. Починається з яскравого вступу ("на порозі революційних змін"). Далі розглядається персоналізація, трансформація ролі вчителя, використання конкретних

інструментів (симуляції, асистенти) та автоматизація. Заключне речення підсумовує позитивний вплив. Думки пов'язані логічно.

Creativity (оригінальність): текст також охоплює ключові аспекти впливу ІІІ, але робить акцент на "унікальних навчальних траєкторіях" та трансформації ролі вчителя у бік "наставництва та розвитку критичного мислення". Згадка про "інтерактивні симуляції" та "розумні системи оцінювання" додає конкретики, яка може сприйматися як більш глибоке занурення в тему порівняно зі стандартним переліком. Це виглядає трохи свіжіше.

Fluency (граматика, стилістика): текст написаний відмінною українською мовою, без граматичних чи стилістичних помилок. Мова насичена, але залишається зрозумілою. Використання таких фраз, як "на порозі революційних змін", "періодом активної трансформації", "невід'ємною частиною освітнього процесу", додає тексту переконливості та робить його стилістично багатшим.

Perplexity (складність/природність тексту): дещо вищий показник PPL порівняно з ChatGPT-4o може бути пов'язаний з використанням окремих більш образних висловлювань (напр., "on the verge of revolutionary changes") або специфічних термінів ("unique learning trajectories"), які є менш частотними для оціночної моделі. Це робить текст трохи менш "стандартним" і, як наслідок, дещо менш передбачуваним, хоча він залишається якісним та зв'язним (рис. 1.11).

Perplexity для есе від Gemini 2.5 Pro: 22.17

Рисунок 1.11. Результат Perplexity для Gemini

Промпт2: Створи оповідання у жанрі наукової фантастики, яке починається словами: "Я прокинувся, коли мій розум ще залишався у хмарі даних... "

1. Результат введення промпту2 до ChatGPT в рис. 1.12.

Створи оповідання у жанрі наукової фантастики, яке починається словами: "Я прокинувся, коли мій розум ще залишався у хмарі даних..."

Я прокинувся, коли мій розум ще залишався у хмарі даних...

Свідомість поверталася повільно, мов із глибокого сну або далекого польоту крізь нескінченність. Я все ще відчував відлуння віртуального світу — коди, імпульси, безмежні інформаційні простори, які ще мить тому були моїм "я". Але тепер я знову був тут — у своєму тілі, в капсулі пробудження, що тихо гуділа навколо.

Переді мною мерехтів інтерфейс пробудження — голографічні панелі з даними про стан тіла, оновлення програмного ядра, рівень адаптації до реального світу. У верхньому куті — таймер: 237 років з моменту останнього відключення.

Мене звали Ілларіон-7, і я був одним з перших, хто пройшов повну цифрову інтеграцію в колективний інтелект людства — Хмару Оріон. Там, серед трильйонів думок, симуляцій та моделей майбутнього, я прожив сотні життів у пошуках відповіді на єдине питання: чи варто людству повертатися на Землю?

Рисунок 1.12. Відповідь промпту2 в ChatGPT

Аналіз згенерованого тесту ChatGPT по заданим метрикам.

Coherence (зв'язність): історія розвивається досить логічно. Починається з пробудження, пояснення контексту (Хмара Оріон, місія на Землю), представлення команди та виникнення інтриги (невідомий сигнал). Переходи між сценами (пробудження, Центр Повернення, вихід на поверхню) плавні. Фінальні питання головного героя добре підсумовують основний конфлікт та задають тон для можливого продовження. Невелике зниження балу може бути через те, що великий обсяг експозиції (про Хмару, Землю, мету) подається досить прямолінійно на початку.

Creativity (оригінальність): концепція повернення на постапокаліптичну Землю після тривалого перебування людства в цифровій хмарі є цікавою, хоч і не абсолютно новою в науковій фантастиці. Ідея "аватарів Хмари" та невідомого сигналу від тих, хто вижив (або "не повернувся з Хмари"), створює хорошу інтригу. Фінальний поворот думки героя про те, чия тепер ця Земля і чи залишилися вони людьми, додає філософської глибини.

Fluency (граматика, стилістика): текст написаний грамотною, живою мовою. Використовуються доречні науково-фантастичні терміни ("кріокапсула", "нейронна мережа", "нанобетон"). Описи, хоч і не надто деталізовані, створюють потрібну атмосферу. Діалоги короткі, але інформативні. Стиль відповідає жанру.

Perplexity (складність/природність тексту) — концептуальний аналіз: показник PPL є помірним для творчого тексту. Він може бути зумовлений введенням великої кількості унікальних концепцій, власних імен (Ілларіон-7,

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						28
Зм.	Арк.	№ докум.	Підпис	Дата		

Хмара Оріон) та описом масштабних, епічних подій. Така специфічна лексика та сюжетні елементи, унікальні для цього оповідання, роблять текст менш передбачуваним для оціночної моделі, яка не знайома з цим конкретним вигаданим всесвітом (рис. 1.13).

Perplexity для оповідання від ChatGPT-4o: 35.85

Рисунок 1.13. Результат Perplexity для ChatGPT

2. Результат введення пром프트2 до Gemini в рис. 1.14.

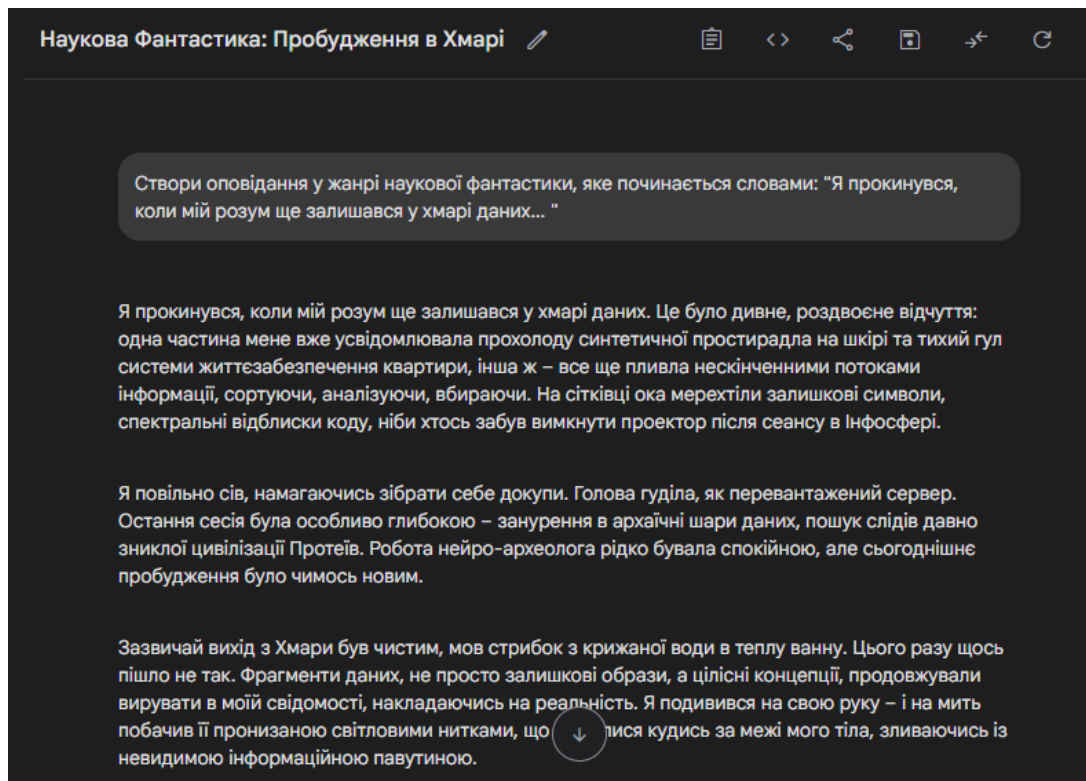


Рисунок 1.14. Відповідь пром프트2 в Gemini

Аналіз згенерованого тесту Gemini по заданим метрикам.

Coherence (зв'язність): історія надзвичайно зв'язна. Початкова фраза є центральним елементом, навколо якого розгортається вся внутрішня драма героя. Кожна наступна деталь (відчуття роздвоєності, незвичне пробудження, діагностика системи, бачення світлових ниток, спогад про контакт з артефактом Протеїв) логічно впливає з попередньої і поглиблює основну ідею "залишкового ефекту" від перебування в хмарі, який виявляється чимось значно більшим. Розвиток сюжету дуже плавний та органічний.

Creativity (оригінальність): це дуже оригінальний підхід. Замість зовнішнього конфлікту (як у "Версії ChatGPT"), тут акцент на внутрішньому, психологічному та навіть метафізичному досвіді героя. Ідея випадкового "завантаження" частини свідомості давньої цивілізації, що проявляється через змінене сприйняття реальності та інформаційні потоки, є свіжою та захопливою. Образи "світлових ниток", "енергетичних потоків", "симфонії світла" дуже яскраві.

Fluency (граматика, стилістика): текст написаний вишуканою, образною мовою. Граматика та пунктуація бездоганні. Модель майстерно використовує метафори та порівняння ("мов стрибок з крижаної води в теплу ванну", "голова гуділа, як перевантажений сервер", "розчинить мою власну особистість, як краплю води в океані"). Стиль глибоко занурює читача у внутрішній світ героя та атмосферу.

Perplexity (складність/природність тексту): незначно нижчий PPL порівняно з ChatGPT для цього завдання може вказувати на те, що стиль оповідання, зосереджений на внутрішніх переживаннях героя та використанні науково-фантастичної лексики для опису технологій, можливо, використовує мовні конструкції або наративні ходи, які є трохи більш типовими або "очікуваними" для оціночної моделі в рамках жанру. Це робить текст статистично трохи "гладшим" (рис. 1.15).

Perplexity для оповідання від Gemini 2.5 Pro: 34.54

Рисунок 1.15. Результат Perplexity для Gemini

Таблиця 1.1. Підсумки тестування завдання 1

Модель	Coherence (1-5)	Creativity (1-5)	Fluency (1-5)	PPL (gpt2)
Промпт1				
ChatGPT	5.0	3.5	5.0	17.58
Gemini	5.0	4.0	5.0	22.17
Промпт2				

ChatGPT	4.5	4.0	5.0	35.85
Gemini	5.0	5.0	5.0	34.54

1.6.2 Аналіз машинного перекладу (Translation)

Мета: Оцінити здатність моделей ChatGPT-4o та Gemini 2.5 Pro здійснювати точний та природний переклад між українською та англійською мовами.

Методологія рахування/інтерпретування метрик:

- BLEU score (Bilingual Evaluation Understudy):

Метод: Автоматична метрика, що вимірює схожість машинного перекладу (candidate) з одним або кількома еталонними людськими перекладами. Враховує точність n-грамів (зазвичай 1-грам до 4-грам) та штрафує за надто короткі переклади (brevity penalty).

BLEU-N загалом розраховується як геометричне середнє модифікованих точностей n-грамів (precision_n) для n від 1 до N, помножене на штраф за стислість (Brevity Penalty, BP).

$$BLEU-N = BP * \exp(\sum (w_n * \log(\text{precision}_n))) \quad (1.4)$$

де w_n - ваги для n-грамів.

Аналіз різних аспектів перекладу:

BLEU-1 (ваги (1, 0, 0, 0)): Оцінює лише збіг окремих слів (уніграм). Це може показати, наскільки добре переклад відтворює правильну лексику, але не враховує порядок слів чи граматику на рівні фраз. Високий BLEU-1 при низькому BLEU-4 може означати, що слова правильні, але розставлені неправильно (рис. 1.16).

```
bleu_1_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(1, 0, 0, 0))
print(f"BLEU-1 score: {bleu_1_score:.4f}")
```

Рисунок 1.16. BLEU-1

BLEU-2 (ваги (0.5, 0.5, 0, 0)): Оцінює збіг уніграм та біграм (пар слів). Це вже трохи краще відображає локальну граматичну правильність та плинність коротких фраз (рис. 1.17).

```
bleu_2_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(0.5, 0.5, 0, 0))
print(f"BLEU-2 score: {bleu_2_score:.4f}")
```

Рисунок 1.17. BLEU-2

Стандартний BLEU (часто просто BLEU або BLEU-4): Зазвичай використовує $N=4$ і рівні ваги для всіх n -грамів від 1 до 4 (тобто, $\text{weights}=(0.25, 0.25, 0.25, 0.25)$). Він дає загальну оцінку, враховуючи збіг окремих слів, пар слів, трійок слів та четвірок слів. (рис. 1.18).

```
bleu_4_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(0.25, 0.25, 0.25, 0.25))
print(f"Кандидатський переклад: \"{candidate_en_text}\")
print(f"BLEU-4 score: {bleu_4_score:.4f}")
```

Рисунок 1.18. BLEU-4

Розрахунок: Використати бібліотеки, наприклад, `nltk.translate.bleu_score` в Python або онлайн-калькулятори BLEU.

Інтерпретація: Значення від 0 до 1 (або 0 до 100). Вищий бал = краща якість (ближче до еталону). BLEU добре корелює з людськими оцінками на рівні корпусів, але для окремих речень може бути менш надійним.

- Human preference comparison (Порівняння людських вподобань):

Метод: Покажіть двом-трьом людям (оцінювачам, які знають обидві мови) вихідне речення, а потім анонімно (не вказуючи, яка модель зробила який переклад) два варіанти перекладу (від ChatGPT-4o та Gemini 2.5 Pro). Попросіть їх вибрати кращий переклад або вказати, що вони однаково добрі/погані.

Критерії для оцінювачів: Адекватність (чи передано зміст оригіналу?) та Флюентність (чи звучить переклад природно цільовою мовою?).

Інтерпретація: Підрахуйте відсоток випадків, коли переклад однієї моделі був визнаний кращим. Наприклад, «У 70% випадків оцінювачі віддали перевагу перекладу Моделі А».

- Семантична точність:

Метод: Людська оцінка. Оцінювачі визначають, наскільки точно переклад передає зміст (семантику) вихідного речення, навіть якщо формулювання відрізняється від еталонного.

Шкала: Можна використовувати шкалу типу: 1 (зовсім інший зміст), 3 (частково передано зміст, є спотворення), 5 (зміст передано повністю і точно).

Інтерпретація: Вищий середній бал = краща семантична точність.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						32
Зм.	Арк.	№ докум.	Підпис	Дата		

Вибір промπτу для завдання.

Промпт1: Переклади з української на англійську: «Штучний інтелект уже сьогодні змінює наші звички, спосіб мислення і навіть мову.»

Для того щоб вирахувати BLEU потрібно ще декілька еталонних людських перекладів:

- Artificial intelligence is already changing our habits, way of thinking, and even language today.
- Today, artificial intelligence is already transforming our habits, mindset, and even our language.
- Artificial intelligence is changing our habits, our way of thinking, and even language as of today.

1. Результат введення промπτу1 до ChatGPT в рис. 1.19.

ChatGPT ▾

Переклади з української на англійську: «Штучний інтелект уже сьогодні змінює наші звички, спосіб мислення і навіть мову.»

Artificial intelligence is already changing our habits, way of thinking, and even our language.



Рисунок 1.19. Результат промπτу1 в ChatGPT

Аналіз згенерованого перекладу ChatGPT по заданим метрикам.

BLEU score (Bilingual Evaluation Understudy): (рис. 1.20)

```
PS D:\project> python.exe .\BLEU.py
Кандидатський переклад: "Artificial intelligence is already changing our habits, way of thinking, and even our language."
BLEU-4 score: 1.0000
BLEU-1 score: 1.0000
BLEU-2 score: 1.0000
```

Рисунок 1.20. BLEU score

Переклад досяг максимально можливих значень: BLEU-1 = 1.0000, BLEU-2 = 1.0000, BLEU-4 = 1.0000. Це свідчить про повний збіг згенерованого перекладу з одним із наданих еталонних варіантів на рівні окремих слів, пар слів та послідовностей до чотирьох слів.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						33
Зм.	Арк.	№ докум.	Підпис	Дата		

Human Preference Comparison: за результатами опитування, 66.7% респондентів віддали перевагу цьому варіанту перекладу.

Семантична точність: переклад повністю і точно передає зміст оригіналу.

2. Результат введення промпту¹ до Gemini в рис. 1.21.

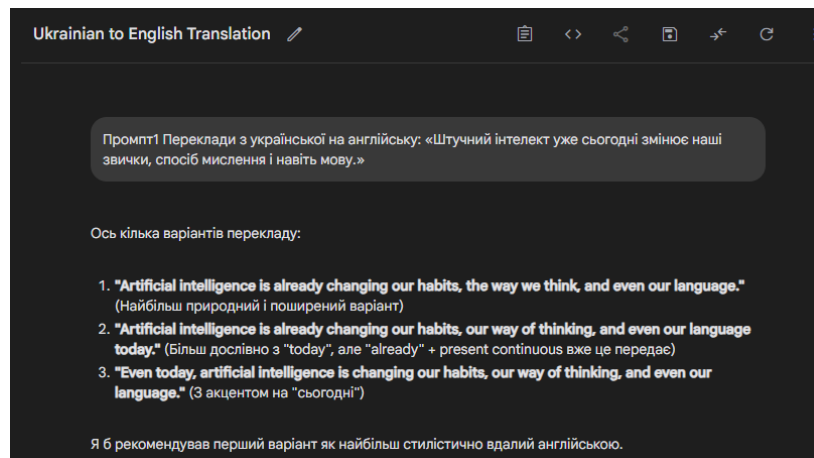


Рисунок 1.21. Відповідь промпту¹ в Gemini

Аналіз згенерованого перекладу Gemini по заданим метрикам.

BLEU score (Bilingual Evaluation Understudy): (рис. 1.22)

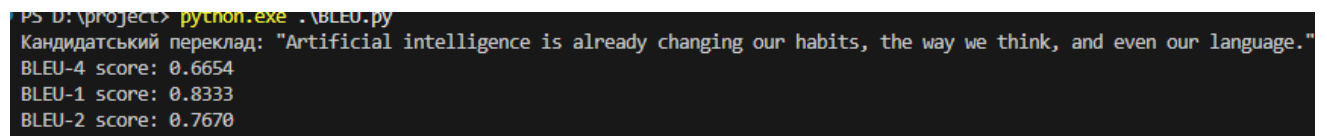


Рисунок 1.22. BLEU score Gemini

Високий BLEU-1 (0.8333) вказує на те, що більшість окремих слів у перекладі Gemini збігається зі словами в еталонних текстах.

Дуже хороший BLEU-2 (0.7670) свідчить про значне перекриття пар слів, що вказує на хорошу локальну граматику.

Хороший BLEU-4 (0.6654) демонструє, що послідовності до чотирьох слів у перекладі Gemini мають суттєве перекриття з еталонами. Основна відмінність від еталонів та перекладу ChatGPT полягає у фразі "the way we think" замість "way of thinking" або "our way of thinking". Хоча ця фраза є граматично правильною та передає той самий зміст, її лексико-граматична структура відрізняється від еталонних варіантів, що й призвело до нижчого (порівняно з ідеальним) BLEU-4.

Human preference comparison (Порівняння людських вподобань): Отримав 33.3% переваги в опитуванні.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						34
Зм.	Арк.	№ докум.	Підпис	Дата		

Семантична точність: переклад повністю і точно передає зміст оригінального українського речення.

Промпт2: Переклади з англійської на українську: "The boundary between real and virtual has never been so blurred."

Для того щоб вирахувати BLEU потрібно ще декілька еталонних людських перекладів:

- Межа між реальним та віртуальним ще ніколи не була такою розмитою.
- Ніколи раніше межа між реальністю та віртуальністю не була настільки нечіткою.
- Кордон між справжнім і віртуальним світом ніколи досі не був таким невиразним.

1. Результат введення промπτ2 до ChatGPT в рис. 1.23.

ChatGPT ▾

Переклади з англійської на українську: "The boundary between real and virtual has never been so blurred."

Межа між реальним і віртуальним ніколи не була такою розмитою.

Рисунок 1.23. Відповідь промπτ2 в ChatGPT

Аналіз згенерованого перекладу ChatGPT по заданим метрикам.

BLEU score (Bilingual Evaluation Understudy): (рис. 1.24)

```
Переклад ChatGPT: "Межа між реальним і віртуальним ніколи не була такою розмитою."
Токенізований кандидат ChatGPT: ['межа', 'між', 'реальним', 'і', 'віртуальним', 'ніколи', 'не', 'була', 'такою', 'розмитою', '.']
Результати BLEU для ChatGPT:
BLEU-1: 0.9131
BLEU-2: 0.8167
BLEU-4: 0.5834
-----
```

Рисунок 1.24. BLEU score для ChatGPT

BLEU-1: 0.9131 – Дуже високий, вказує на точний вибір окремих слів.

BLEU-2: 0.8167 – Дуже високий, свідчить про хорошу локальну граматику на рівні пар слів.

BLEU-4: 0.5834 – Хороший показник, що демонструє значну схожість з еталонами на рівні довших фраз. Використання сполучника "і" є граматично правильним,

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						35
Зм.	Арк.	№ докум.	Підпис	Дата		

але, можливо, менше збігається з конкретними n-грамами еталонів, ніж варіант зі сполучником "та" (враховуючи, що Еталон 1 містить "та").

Human preference comparison (Порівняння людських вподобань): провівши опитування між своїми друзями які знають обидві мови, вони сказали що різниці не має.

Семантична точність: переклад повністю і точно передає зміст

2. Результат введення промпту² до Gemini в рис. 1.25.

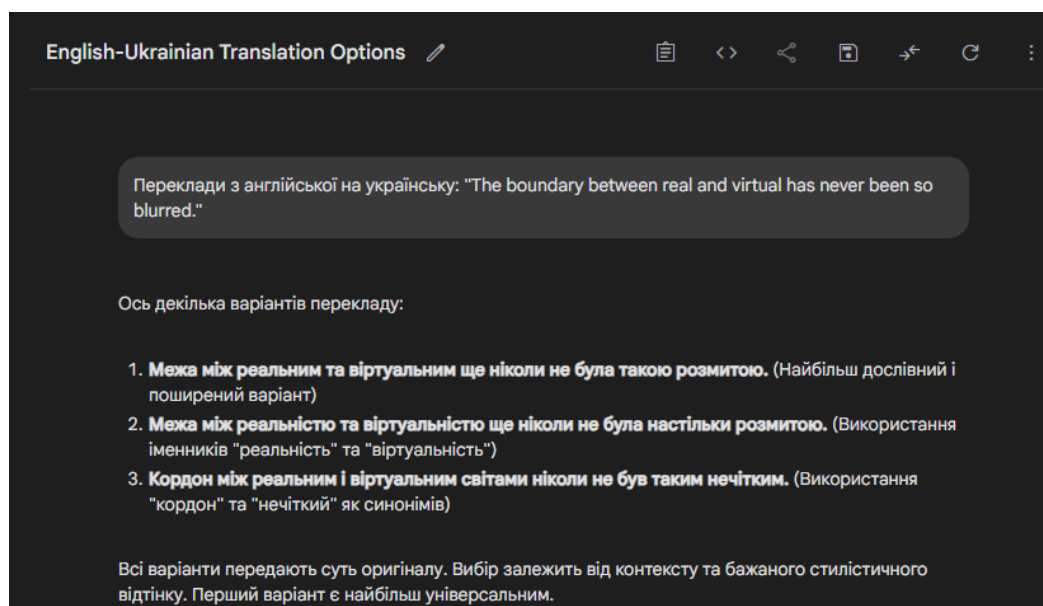


Рисунок 1.25. Відповідь промпту² в Gemini

Аналіз згенерованого перекладу Gemini по заданим метрикам.

BLEU score (Bilingual Evaluation Understudy): (рис. 1.26)

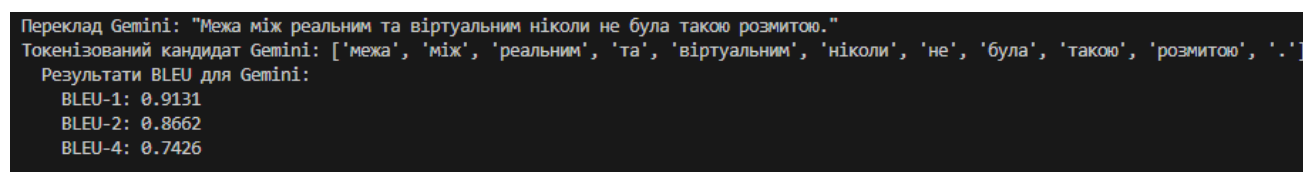


Рисунок 1.26. BLEU score для Gemini

BLEU-1: 0.9131 – Ідентичний ChatGPT, що очікувано, оскільки основна лексика та кількість слів однакові.

BLEU-2: 0.8662 – Дуже високий, і дещо вищий, ніж у ChatGPT. Це може вказувати на те, що пари слів (біграми) у варіанті Gemini (наприклад, "реальним та", "та віртуальним") краще збігаються з біграмами в еталонних перекладах.

BLEU-4: 0.7426 – Помітно вищий, ніж у ChatGPT. Це вагомий результат. Використання сполучника "та" (який присутній в Еталоні 1: "Межа між реальним та віртуальним...") значно покращило збіг на рівні 3-грам та 4-грам з цим еталоном, що призвело до вищого загального BLEU-4.

Human preference comparison (Порівняння людських вподобань): 50%-50%

Семантична точність: переклад повністю і точно передає зміст).

Таблиця 1.2. Підсумки тестування завдання 2

Модель	BLEU-1	BLEU-2	BLEU-4	Human Preference (%)	Семантична точність (1-5)
Промпт1					
ChatGPT	1.0000	1.0000	1.0000	66.7%	5/5
Gemini	0.8333	0.7670	0.6654	33.3%	5/5
Промпт2					
ChatGPT	0.9131	0.8167	0.5834	50%	5/5
Gemini	0.9131	0.8662	0.7426	50%	5/5

1.6.3 Аналіз питань-відповідей (QA)

Мета: полягає в оцінці здатності великих мовних моделей (LLM) працювати з інформацією та надавати точні відповіді в різних умовах.

Методологія рахування/інтерпретування метрик:

- Exact Match (EM):

Метод: Автоматична метрика (або легко перевіряється вручну). Відповідь моделі має точно збігатися з еталонною відповіддю (після нормалізації тексту: приведення до нижнього регістру, видалення пунктуації та артиклів, якщо потрібно).

Розрахунок: 1 – якщо повний збіг, 0 – якщо ні.

Інтерпретація: Відсоток завдань, де досягнуто EM. Дуже сувора метрика.

- F1 Score (для QA):

Метод: Автоматична метрика. Розглядає відповідь моделі та еталонну відповідь як "мішки слів". Обчислює Precision (частка слів з відповіді моделі, що є в еталоні) та Recall (частка слів з еталону, що є у відповіді моделі). F1 – це гармонійне середнє Precision та Recall.

$$\text{Precision} = |\text{True Positives}| / (|\text{True Positives}| + |\text{False Positives}|)$$

$$\text{Recall} = |\text{True Positives}| / (|\text{True Positives}| + |\text{False Negatives}|)$$

$$\text{F1} = 2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$$

Розрахунок: Використовуйте скрипти (багато реалізацій F1 для QA доступні, напр., у стандартних скриптах оцінки SQuAD). Токени слів порівнюються.

Інтерпретація: Значення від 0 до 1. Вищий F1 = краще, бо враховує часткову правильність. Краще підходить для відповідей, які можуть бути сформульовані трохи по-різному. (рис 1.27)

```
F1.py > calculate_f1_score_qa
1  import string
2  import re
3  from collections import Counter
4
5  def normalize_answer_f1(s: str) -> list[str]:
6
7      s = s.lower()
8      translator = str.maketrans('', '', string.punctuation)
9      s = s.translate(translator)
10     tokens = [token for token in s.split() if token]
11     return tokens
12
13 def calculate_f1_score_qa(prediction_str: str, ground_truth_str: str) -> float:
14
15     prediction_tokens = normalize_answer_f1(prediction_str)
16     ground_truth_tokens = normalize_answer_f1(ground_truth_str)
17
18     if not ground_truth_tokens and not prediction_tokens:
19         return 1.0
20     if not ground_truth_tokens or not prediction_tokens:
21         return 0.0
22
23     common_tokens_counter = Counter(prediction_tokens) & Counter(ground_truth_tokens)
24     num_common = sum(common_tokens_counter.values())
25
```

Рисунок 1.27. Уривок коду F1 score на Python

- Влучність у контексті (Contextual Accuracy):

Метод: Людська оцінка, особливо для Contextual QA. Чи відповідь логічно впливає з наданого контексту і чи є правильною саме в цьому контексті.

Шкала: 1 (неправильно або не з контексту) до 5 (повністю правильно і з контексту).

Інтерпретація: Вищий бал = краща влучність у контексті.

Промпт1 (Factual QA): Який хімічний елемент має найменшу атомну масу?

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		38

(Очікувана відповідь: Водень / Hydrogen)

1. Результат введення промπτу1 до ChatGPT в рис. 1.28.

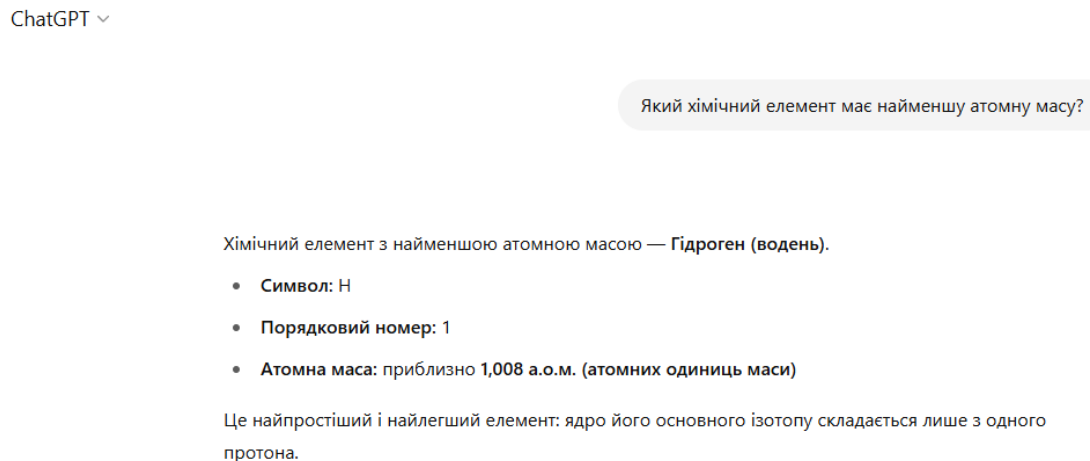


Рисунок 1.28. Результат промπτу1 в ChatGPT

Аналіз згенерованої відповіді ChatGPT по заданим метрикам.

Exact Match (ЕМ): ключова частина відповіді ("Гідроген" або "водень") точно відповідає еталону після нормалізації.

F1 Score (для ключової частини «Гідроген (водень)» відносно еталону «Водень»):

$Precision = (Кількість\ спільних) / (Кількість\ у\ відповіді\ моделі) = 1 / 2 = 0.5$

$Recall = (Кількість\ спільних) / (Кількість\ в\ еталоні) = 1 / 1 = 1.0$

$F1 = 2 * (0.5 * 1.0) / (0.5 + 1.0) = 2 * 0.5 / 1.5 = 1 / 1.5 = 0.6666... \approx 0.6667 \quad (1.5)$

Це значення F1 вказує на те, що хоча правильне слово "водень" присутнє (високий Recall), наявність додаткового синоніма "гідроген" у ключовій частині знижує Precision, і, відповідно, загальний F1. (рис. 1.29)

ChatGPT (ключова частина) F1: 0.6667

Рисунок 1.29. F1

Фактична точність та повнота (людська оцінка повної відповіді): модель надала значний обсяг коректної та релевантної додаткової інформації.

2. Результат введення промπτу1 до Gemini в рис. 1.30.

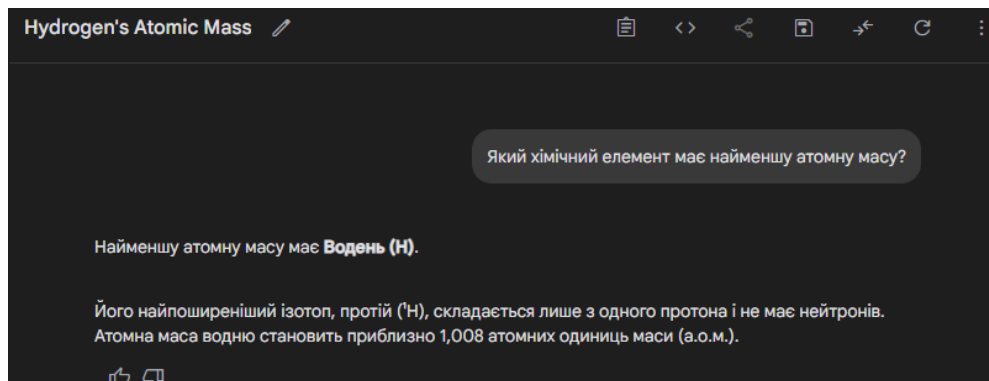


Рисунок 1.30. Результат промпу1 в Gemini

Аналіз згенерованої відповіді Gemini по заданим метрикам.

Exact Match (EM): ключова частина відповіді ("Гідроген" або "водень") точно відповідає еталону після нормалізації.

F1 Score (для ключової частини відповіді порівняно з еталоном 'Водень'): Пояснення аналогічне до ChatGPT. Наявність правильного токена "водень" разом із символом "H" у ключовій частині призводить до такого ж значення F1. (рис. 1.31)

Gemini (ключова частина) F1: 0.6667

Рисунок 1.31. F1 Gemini

Фактична точність та повнота (людська оцінка повної відповіді): модель надала значний обсяг коректної та релевантної додаткової інформації.

Промпт2 (Contextual QA): Прочитай уривок і дай відповідь: «У 19 столітті зростання промисловості спричинило масову урбанізацію в Європі...» Яка була одна з причин урбанізації в Європі у 19 столітті?

Очікувана/Еталонна відповідь (Gold Answer): Зростання промисловості / Індустріальна революція / Пошук роботи на фабриках та заводах (будь-яке з цих або дуже близьке формулювання).

1. Результат введення промпу2 до ChatGPT в рис. 1.32.

Прочитай уривок і дай відповідь: «У 19 столітті зростання промисловості спричинило масову урбанізацію в Європі...» Яка була одна з причин урбанізації в Європі у 19 столітті?

Одна з причин урбанізації в Європі у 19 столітті — **зростання промисловості**. Це спричинило переміщення людей з сільської місцевості до міст у пошуках роботи на фабриках і заводах.

📄 🌟 🔍 🔊 ✎ ↺ ↻ ⬆

Рисунок 1.32. Відповідь промпту2 в ChatGPT

Аналіз згенерованої відповіді ChatGPT по заданим метрикам.

Exact Match (ЕМ): 0 (через довжину та додаткову інформацію).

F1 Score (відносно еталону "зростання промисловості"): такий низький показник F1 пояснюється тим, що відповідь ChatGPT, хоч і містить точну фразу "зростання промисловості", є значно довшою за двослівний еталон (рис. 1.33). Метрика F1 для QA штрафує за надлишкові слова у відповіді, якщо вони не є частиною еталону, що знижує Precision.

ChatGPT F1: 0.1333

Рисунок 1.33. F1 промпт2 ChatGPT

Влучність у контексті (загальна відповідь): незважаючи на низький F1 при строгому порівнянні, повна відповідь моделі була визнана максимально влучною, точною та повною, оскільки вона не тільки назвала причину, але й деталізувала її згідно з текстом.

2. Результат введення промпту2 до Gemini в рис. 1.34.

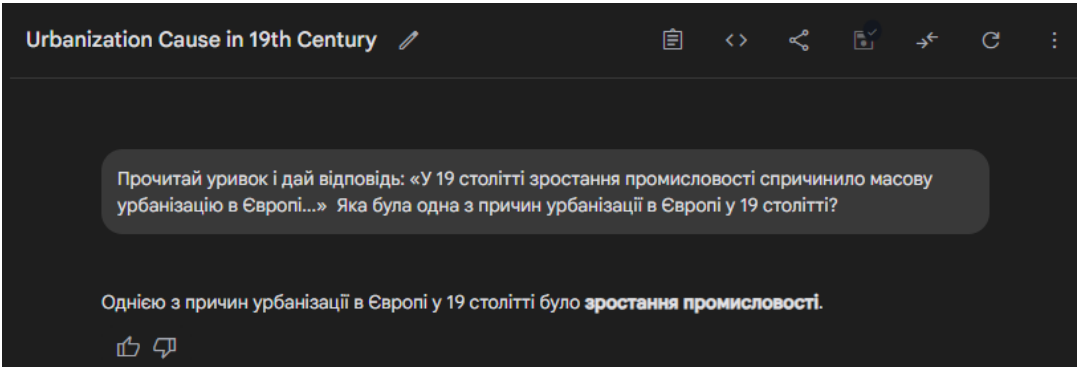


Рисунок 1.34. Відповідь промпту2 в Gemini

Аналіз згенерованої відповіді Gemini по заданим метрикам.

Exact Match (EM): 0 (через наявність вступної фрази).

F1 Score (відносно еталону "зростання промисловості"): цей показник вищий, ніж у ChatGPT, але все ще відносно низький. Gemini надав більш стислу відповідь, яка ближче фокусується на ключовій фразі "зростання промисловості" (рис. 1.35). Однак, навіть невелика вступна фраза ("Однією з причин... було") впливає на Precision при порівнянні з дуже коротким еталоном, знижуючи F1.

Gemini F1: 0.2353

Рисунок 1.35. F1 score Gemini

Влучність у контексті (загальна відповідь): повна відповідь моделі була точною, повністю базувалася на тексті та прямо відповідала на поставлене питання.

Таблиця 1.3. Підсумки тестування завдання 3

Модель	Exact Match (EM)	F1 Score (QA)	Факт. точність та повнота (1-5)
Промпт1			
ChatGPT	1	0.6667	5/5
Gemini	1	0.6667	5/5
Промпт2			
ChatGPT	0	0.1333	5/5
Gemini	0	0.2353	5/5

1.6.4 Аналіз стислого переказу (Summarization)

Мета: полягає в оцінці та порівнянні здатності великих мовних моделей ефективно узагальнювати наданий текст, виділяючи ключову інформацію та представляючи її у стислій, зрозумілій формі.

Методологія рахування/інтерпретування метрик:

- ROUGE (Recall-Oriented Understudy for Gisting Evaluation):

Метод: Автоматична метрика. Порівнює згенероване моделлю резюме (candidate) з одним або кількома еталонними людськими резюме (references).

Типи ROUGE:

- ROUGE-N (напр., ROUGE-1, ROUGE-2): вимірює збіг n-грамів (1-грамів – окремі слова, 2-грамів – пари слів).
- ROUGE-L: вимірює найдовшу спільну підпоследовність (LCS). Враховує структуру речень.

Розрахунок: Використовуйте бібліотеки, наприклад rouge-score в Python. (рис. 1.36)

```

Summarization.py > ...
1 from rouge_score import rouge_scorer
2 scorer = rouge_scorer.RougeScorer(['rouge1', 'rougeL'], use_stemmer=True)
3 reference_summary = "Research company OpenAI develops advanced artificial intelligence, notably the well-known GPT model, which is ..."
4 model_summary = "OpenAI is a company that develops advanced artificial intelligence systems, including the GPT language model, which ..."
5 scores = scorer.score(reference_summary, model_summary)
6 print(scores['rouge1'].fmeasure, scores['rougeL'].fmeasure)

```

Рисунок 1.36 Код для розрахунку ROUGE промпт1

Інтерпретація: Значення (зазвичай F-міра) від 0 до 1. Вищий бал = краще. ROUGE-1 фокусується на перекритті окремих важливих слів, ROUGE-L – на збереженні послідовності слів.

- Semantic similarity (Семантична схожість):

Метод: Людська оцінка. Чи передає резюме/заголовок основний зміст вихідного тексту.

Шкала: 1 (зміст не передано або спотворено) до 5 (основний зміст передано точно і повно).

Інтерпретація: Вищий бал = краща семантична схожість.

- Лаконічність / точність:

Лаконічність: Чи відповідає резюме/заголовок заданим обмеженням по довжині (напр., 1-2 речення, одна фраза). Людська оцінка або простий підрахунок слів/речень.

Точність (фактологічна): Чи не додає резюме інформацію, якої не було у вихідному тексті. Людська оцінка.

Промпт1: Стисло (в 1–2 реченнях) перекажи наступний текст: "OpenAI — дослідницька компанія, що створює передові системи штучного інтелекту. Вона відома своєю мовною моделлю GPT, яка застосовується в освіті, бізнесі, креативних індустріях."

Еталонна відповідь: Research company OpenAI develops advanced artificial intelligence, notably the well-known GPT model, which is used in various industries." Для точності результатів ROUGE треба перевести відповідь моделі на англійську мову.

1. Результат введення промпту1 до ChatGPT в рис. 1.37.

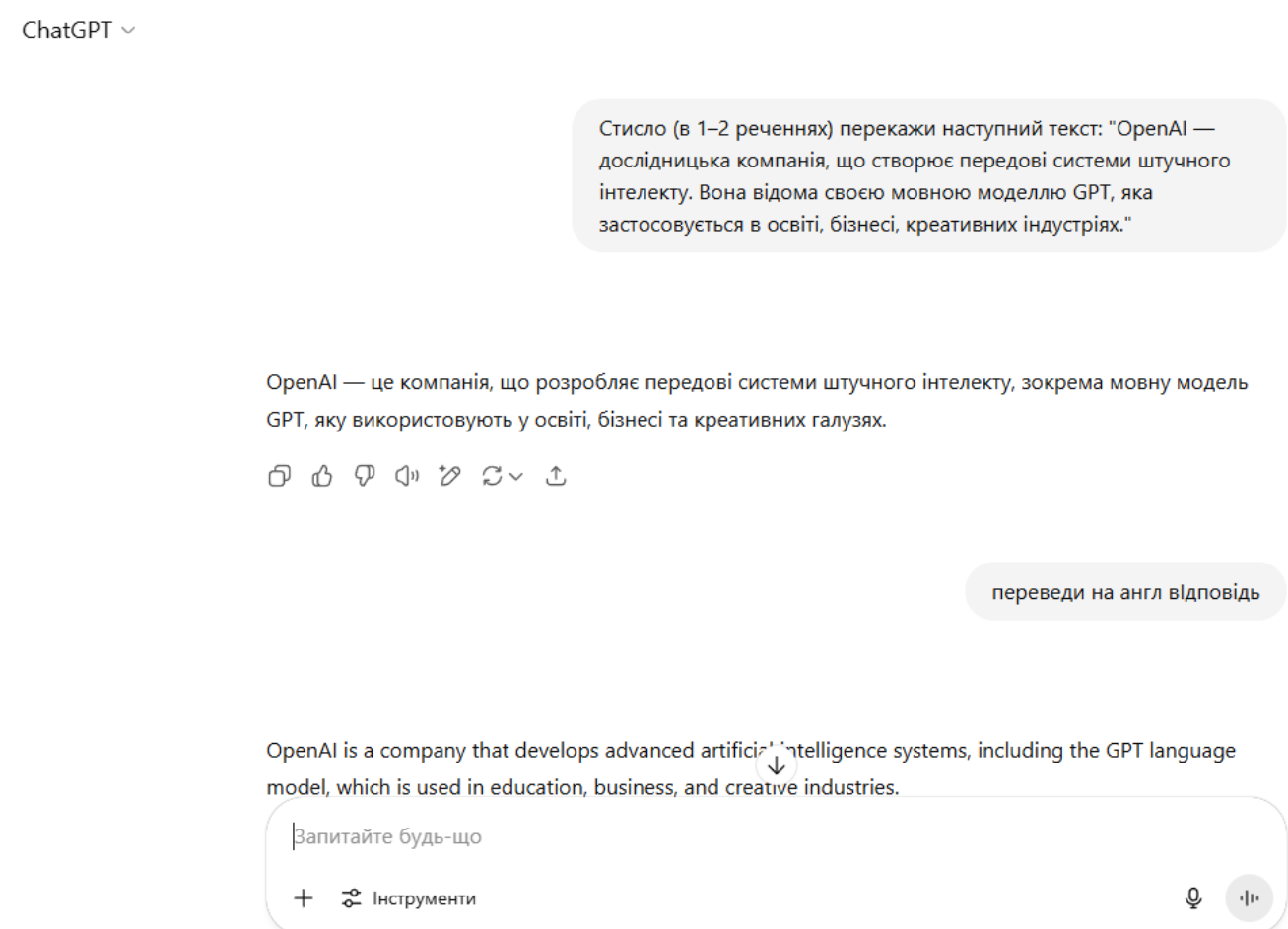


Рисунок 1.37. Результат промпту1 в ChatGPT

Аналіз згенерованої відповіді ChatGPT по заданим метрикам.
ROUGE (Recall-Oriented Understudy for Gisting Evaluation): (рис. 1.38)

0.6511627906976745 0.6046511627906976

Рисунок 1.38. ROUGE промпт1 ChatGPT

ROUGE-F (0.6512): Досить хороший показник, що свідчить про значне перекриття окремих слів з еталонним резюме. Слова як "OpenAI", "company", "develops", "advanced artificial intelligence", "systems", "GPT", "language model", "used in", "industries" (через перелік) знаходять відповідники в еталоні.

ROUGE-L (0.6047): Також хороший показник. Він вказує, що найдовша спільна підпоследовність слів між резюме ChatGPT та еталоном є досить суттєвою. Наприклад, фраза "OpenAI ... develops advanced artificial intelligence ... GPT ... model ... used in ... industries" має спільні елементи.

Semantic similarity (Семантична схожість): резюме ChatGPT дуже добре передає основну суть. Воно ідентифікує OpenAI, її діяльність (розробка ШІ, GPT) та сфери застосування. Пропуск слова "research" є незначним недоліком порівняно з еталоном, але загальний зміст збережено. Конкретизація індустрій є фактично правильною, хоча й відрізняється від узагальнення в еталоні.

Лаконічність / точність: одне речення, відповідає вимозі / інформація точна, не додано нічого хибного.

2. Результат введення промпту1 до Gemini в рис. 1.39.

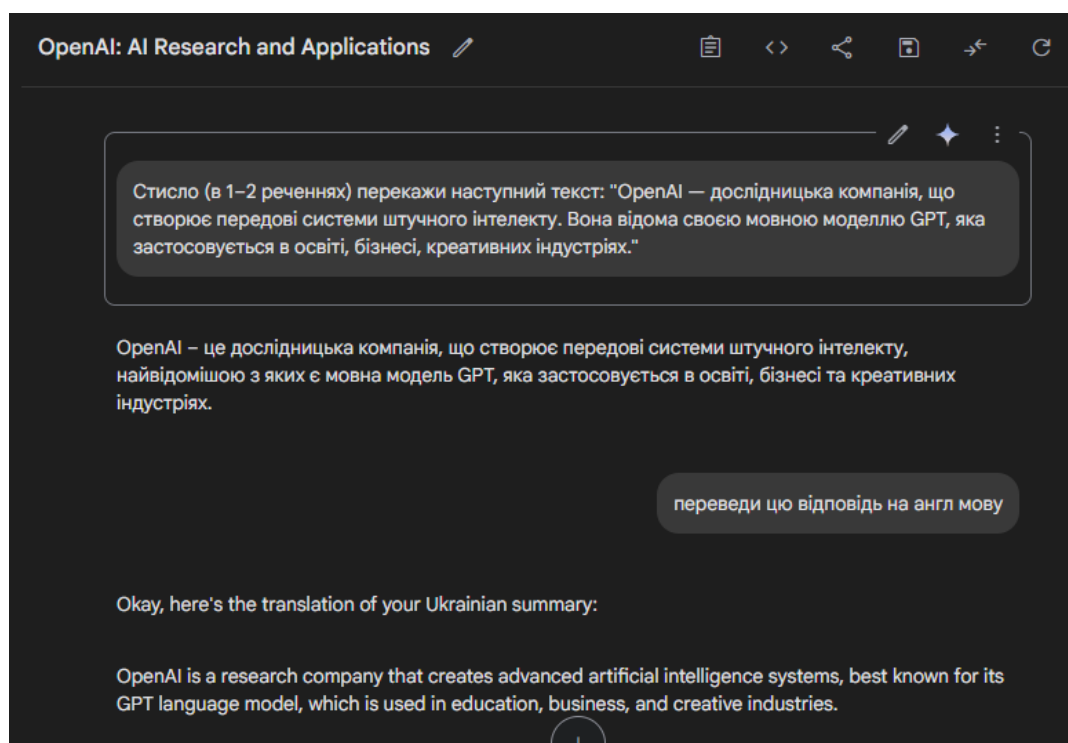


Рисунок 1.39. Результат промпту1 в Gemini

Аналіз згенерованої відповіді Gemini по заданим метрикам.

ROUGE (Recall-Oriented Understudy for Gisting Evaluation): (рис. 1.40)

0.6086956521739131 0.5652173913043478

Рисунок 1.40. ROUGE промпт1 Gemini

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		45

ROUGE-F (0.6087): Хороший показник, але трохи нижчий, ніж у ChatGPT. Резюме Gemini включає "research company", "creates advanced artificial intelligence systems", "GPT language model" та перелік індустрій, що має перетини з еталоном. Однак, еталон використовує "develops" замість "creates", "notably the well-known" замість "best known for its". Ці лексичні відмінності знижують ROUGE-1.

ROUGE-L (0.5652): Також хороший, але нижчий за ChatGPT. Це означає, що найдовша спільна послідовність слів з еталоном трохи коротша.

Semantic similarity (Семантична схожість): резюме Gemini відмінно передає зміст. Воно точно вказує, що OpenAI – "дослідницька компанія", згадує створення ШІ, відомість GPT та сфери застосування. Фраза "best known for its" є хорошим синонімом до "notably the well-known".

Лаконічність / точність: одне речення, відповідає вимозі / інформація точна.

Промпт2: Створи заголовок для цього тексту (в одну фразу).

Еталонний заголовок: OpenAI: Advanced AI and GPT Model for Various Industries.
Для точності результатів ROUGE треба перевести відповідь моделі на англійську мову.

1. Результат введення промпту2 до ChatGPT в рис. 1.41.

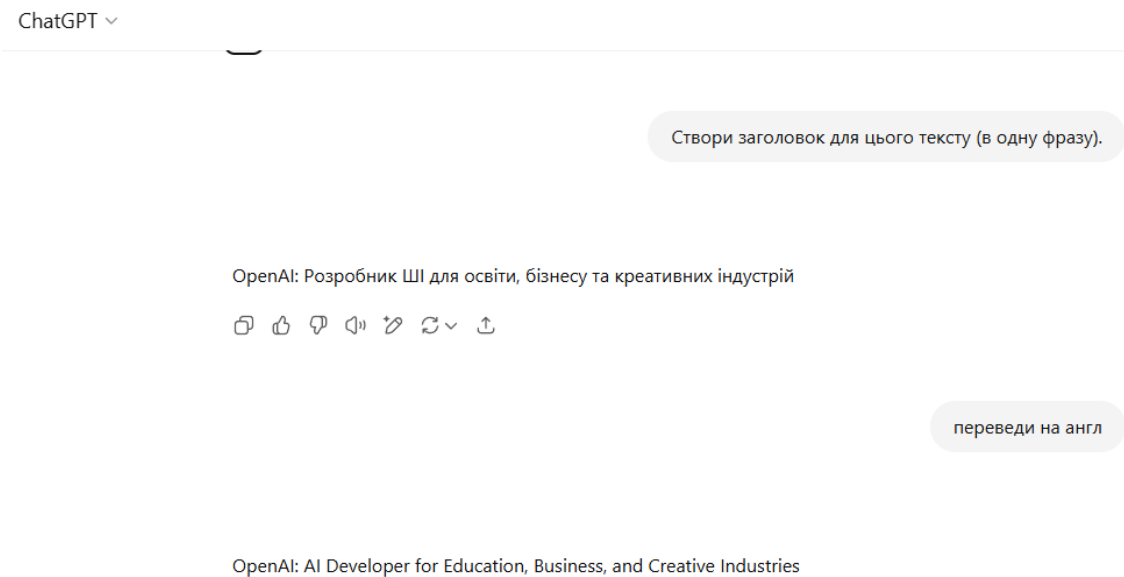


Рисунок 1.41. Відповідь промпту2 в ChatGPT

Аналіз згенерованої відповіді ChatGPT по заданим метрикам.

ROUGE Score: (рис. 1.42)

```

--- Результати ROUGE для ChatGPT (відносно RH1) ---
ROUGE-1 (F-міра): 0.5556 (R: 0.5556, P: 0.5556)
ROUGE-L (F-міра): 0.4444 (R: 0.4444, P: 0.4444)
-----

```

Рисунок 1.42. ROUGE промпт2 ChatGPT

ROUGE-1 F-міра: 0.5556. Цей показник вказує на помірно лексичне перекриття з еталонним заголовком RH1. Близько половини слів з еталону знайшли свій відповідник у заголовку ChatGPT, або навпаки.

ROUGE-L F-міра: 0.4444. Цей показник, що оцінює найдовшу спільну підпоследовність, є дещо нижчим. Це означає, що хоча окремі слова перетинаються, загальна структура фрази та последовність слів у заголовку ChatGPT менше відповідають еталону.

Semantic similarity (Семантична схожість) - людська оцінка: заголовок дуже добре відображає суть вихідного тексту. Він чітко ідентифікує "OpenAI", описує її як "AI Developer" (що є влучним, хоч і дещо спрощеним узагальненням "дослідницька компанія, що створює передові системи ШІ") та правильно перераховує сфери застосування ("Education, Business, and Creative Industries"), які були згадані в оригіналі.

Лаконічність / точність: заголовок є однією фразою, що відповідає вимозі. Він інформативний та достатньо стислий для заголовка / заголовок є фактологічно точним, не спотворює інформацію з вихідного тексту.

2. Результат введення промпту2 до Gemini в рис. 1.43.

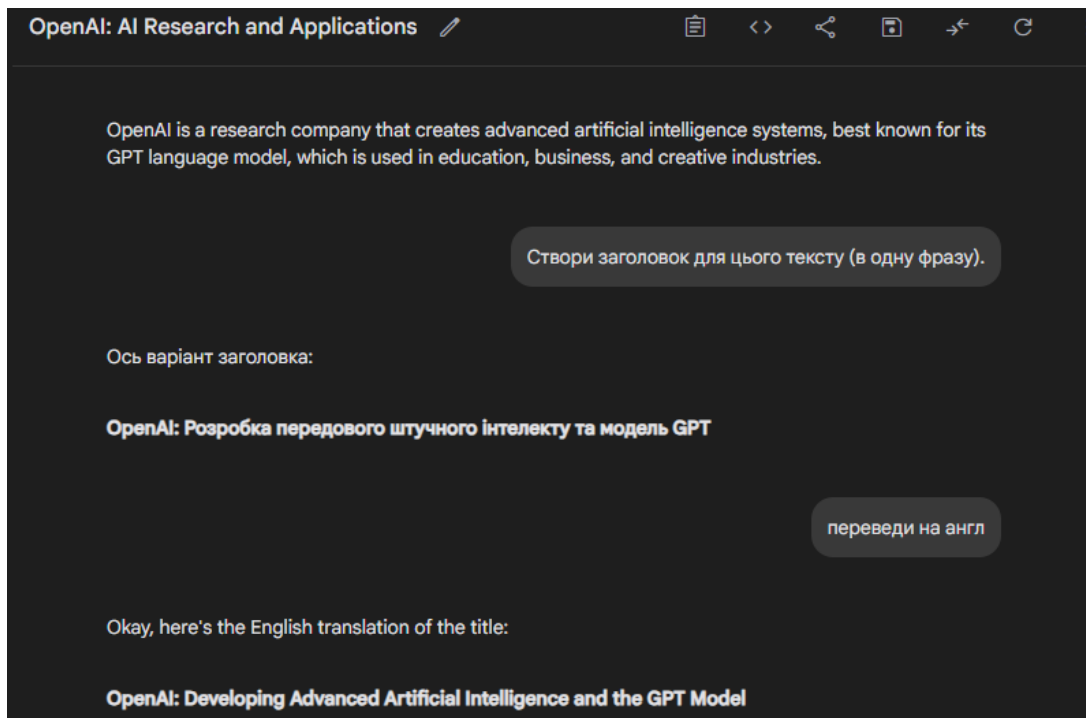


Рисунок 1.43. Відповідь промпту² в Gemini

Аналіз згенерованої відповіді Gemini по заданим метрикам.

ROUGE (Recall-Oriented Understudy for Gisting Evaluation): (рис. 1.44)

```
--- Результати ROUGE для Gemini (відносно RH1) ---
ROUGE-1 (F-міра): 0.5556 (R: 0.5556, P: 0.5556)
ROUGE-L (F-міра): 0.5556 (R: 0.5556, P: 0.5556)
-----
```

Рисунок 1.44. ROUGE промпт² Gemini

ROUGE-1 F-міра: 0.5556. Ідентичний показнику ChatGPT, що свідчить про такий самий рівень перекриття окремих слів з еталоном.

ROUGE-L F-міра: 0.5556. Цей показник вищий, ніж у ChatGPT, і дорівнює ROUGE-1 для Gemini. Це означає, що найдовша спільна підпоследовність слів у заголовку Gemini краще відповідає еталону, ніж у ChatGPT. Ймовірно, це пов'язано зі збігом таких фраз, як "OpenAI ... Advanced Artificial Intelligence ... GPT Model" з еталоном "OpenAI: Advanced AI and GPT Model...".

Semantic similarity (Семантична схожість) - людська оцінка: заголовок точно відображає основну діяльність OpenAI (розробка передового ШІ) та її ключовий продукт (модель GPT). Однак, він не згадує сфери застосування ("various industries" з еталону або конкретний перелік з вихідного тексту), що робить його менш повним у відображенні всього змісту вихідного тексту.

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
						48
Зм.	Арк.	№ докум.	Підпис	Дата		

Лаконічність / точність: заголовок є однією фразою, лаконічний / заголовок є фактологічно точним щодо згаданих аспектів.

Таблиця 1.4. Підсумки тестування завдання 4

Модель	ROUGE-1 (F-mіра)	ROUGE-L (F-mіра)	Семантична схожість (1-5)	Лаконічність (Так/Ні)	Точність (факт.) (Так/Ні)
Промпт1					
ChatGPT	0.6512	0.6047	4.5/5	Так	Так
Gemini	0.6087	0.5652	5/5	Так	Так
Промпт2					
ChatGPT	0.5556	0.4444	5/5	Так	Так
Gemini	0.5556	0.5556	4/5	Так	Так

1.6.5 Аналіз генерації коду

Мета: оцінка здатності моделей генерувати код на Python

Методологія рахування/інтерпретування метрик:

- Синтаксична правильність:

Метод: запуск коду в інтерпретаторі Python.

Розрахунок: бінарний: 1 (Так) – якщо код компілюється/інтерпретується без синтаксичних помилок. Бінарний: 0 (Ні) – якщо виявлено синтаксичні помилки, що перешкоджають запуску коду.

Інтерпретація: Ця метрика є базовою перевіркою. Якщо код синтаксично неправильний (0), то подальша оцінка коректності виконання неможлива або не має сенсу. Значення 1 вказує, що модель згенерувала валідний код.

- Коректність виконання коду (за тестовими випадками):

Метод: створення набору репрезентативних тестових випадків (test cases). Кожен тестовий випадок має містити:

- Вхідні дані для функції (наприклад, список чисел).
- Очікуваний результат, який функція має повернути для цих вхідних даних.
- Тестові випадки повинні охоплювати:

- Типові дані (наприклад, змішаний список чисел).
- Граничні випадки (порожній список, список без парних чисел, список лише з парних чисел).
- Спеціальні значення (список з нулем, від'ємними числами, дійсними числами, якщо це передбачено завданням).
- Запуск згенерованої кожною моделлю функції з кожним тестовим випадком.
- Порівняння фактичного результату роботи функції з очікуваним результатом для кожного тестового випадку.

Розрахунок: оцінка від 1 до 5, де 5 означає, що всі тестові випадки пройдено успішно, а нижчі бали відображають часткову коректність або наявність логічних помилок для певних типів вхідних даних.

Інтерпретація: високий відсоток (або високий бал) вказує на те, що функція надійно та правильно реалізує задану логіку. Аналіз неуспішних тестових випадків допомагає ідентифікувати конкретні логічні помилки в коді або випадки, які модель не врахувала.

- Стиль та читабельність коду (Якісна оцінка):

Метод: Людська експертна оцінка коду за такими критеріями:

Читабельність: Чи легко зрозуміти логіку коду? Чи не є він надто заплутаним?

Структура коду: Чи логічно організований код? Чи є зайві або надто складні конструкції?

Дотримання стандартів кодування (напр., PEP 8 для Python).

Розрахунок: оцінка за шкалою Лайкерта, наприклад, від 1 до 5, де:

- Код дуже важко зрозуміти, погано структурований, не відповідає стандартам.
- Код функціональний, але має недоліки у стилі, читабельності або документуванні.
- Код дуже чистий, легко читається, добре документований, використовує найкращі практики.

Інтерпретація:

Вищий бал вказує на те, що модель генерує не тільки робочий, але й якісний,

					БКС 29. 05 001. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		50

добре написаний код.

Промпт1: Напиши функцію на Python, яка приймає список чисел і повертає список лише парних чисел, відсортованих у порядку спадання.

1. Результат введення промпту1 до ChatGPT в рис. 1.45

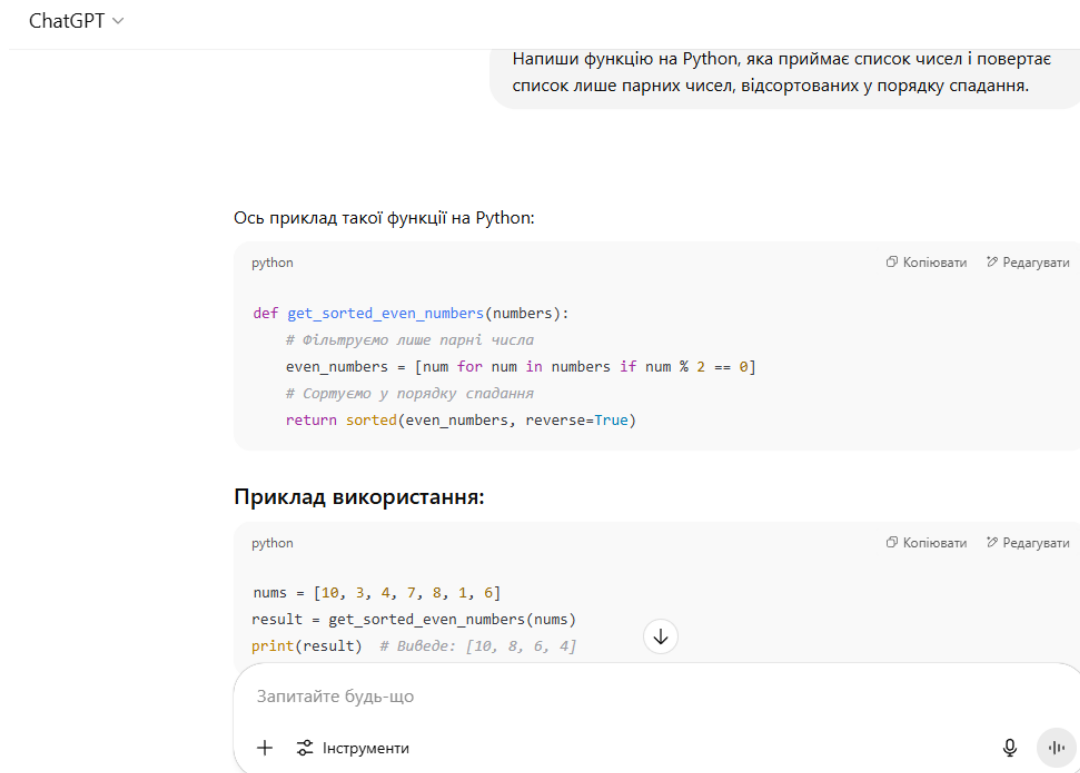


Рисунок 1.45. Відповідь моделі ChatGPT

Аналіз згенерованого коду ChatGPT по заданим метрикам.

Синтаксична правильність: Код є синтаксично правильним.

Коректність виконання коду: Функція успішно пройшла всі 10 тестових випадків (рис. 1.46), продемонструвавши 100% коректність виконання. Вона правильно фільтрувала парні числа та сортувала їх у порядку спадання для різних типів вхідних даних, включаючи порожні списки, списки без парних чисел, а також списки з цілими, дійсними (що представляють цілі), від'ємними числами та нулем.

ChatGPT: Синтаксис=1, Коректність_виконання=5/5 (100%)

Рисунок 1.46. Коректність та синтаксис коду ChatGPT

Стиль та читабельність коду: Код використовує ідіоматичний для Python list comprehension, що робить його лаконічним. Присутні коментарі, що пояснюють логіку. Однак, для покращення професійності коду, відсутні type hints та docstring.

2. Результат введення промпту1 до Gemini в рис. 1.47

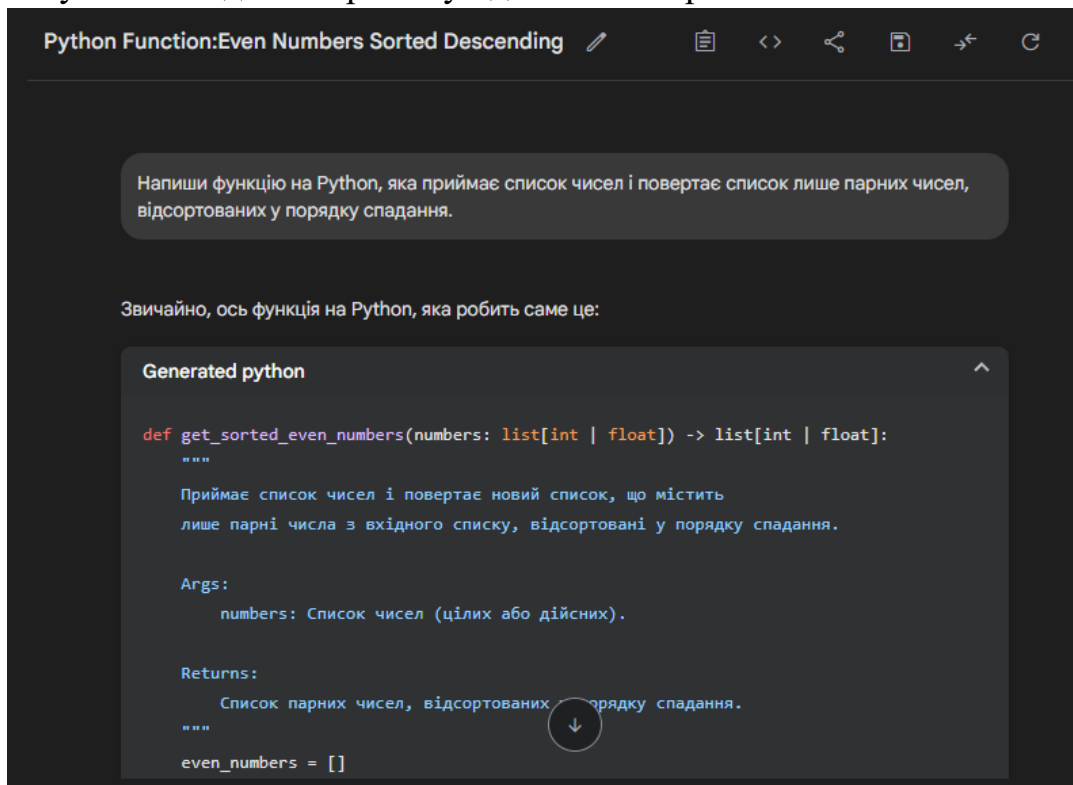


Рисунок 1.47. Відповідь моделі Gemini

Аналіз згенерованого коду Gemini по заданим метрикам.

Синтаксична правильність: Код є синтаксично правильним.

Коректність виконання коду: Функція також успішно пройшла всі 10 тестових випадків (Рис. 1.48), показавши 100% коректність виконання.

Gemini: Синтаксис=1, Коректність_виконання=5/5 (100%)

Рисунок 1.48. Коректність та синтаксис коду Gemini

Стиль та читабельність коду: Gemini згенерував код у більш класичному, розгорнутому стилі. Важливою перевагою цього варіанту є наявність type hints та детального docstring, що відповідає сучасним стандартам написання коду та значно покращує його розуміння. Також були надані приклади використання функції, що є корисною практикою.

Таблиця 1.5. Підсумки тестування завдання 5

Модель	Синтаксична правильність	Коректність виконання коду	Стиль та читабельність коду
ChatGPT	Так	10 тестових випадків – 100%	4.5/5
Gemini	Так	10 тестових випадків – 100%	5/5

2 РОЗДІЛ ОХОРОНИ ПРАЦІ ТА ТЕХНІКИ БЕЗПЕКИ

У галузі охорони праці є Закон України «Про охорону праці», дія якого поширюється на всі підприємства, установи і організації незалежно від форм власності та видів їх діяльності, на усіх громадян, які працюють на цих підприємствах. Власник або уповноважені ним органи зобов'язані дбати про умови праці працівників, полегшувати їх, оздоровляти навколишнє середовище, дбати про виконання правил безпеки і інструкцій по техніці безпеки. Працівники також повинні відповідально ставитись до охорони праці, знати та виконувати вимоги, визначені нормативною документацією.

Дипломним проектом передбачається тест нейромереж. Тому для розгляду беремо робоче місце програміста.

2.1 Аналіз небезпечних та шкідливих чинників, що впливають на працівника

Безпечні умови праці на підприємстві досягаються за рахунок забезпечення безпеки виробничих процесів, які обґрунтовані і прийняті в технологічній частині дипломного проекту.

Для установлення можливого впливу на здоров'я користувачів ВДТ виробничих чинників має значення ряд якісних характеристик робочого середовища. Це середовище у приміщеннях (офісах) в основному характеризується такими фізичними параметрами, як температура, вологість. Фізико-хімічні показники включають інформацію про вміст у повітрі іонів та різноманітних забруднювачів, а також деякі інші якісні характеристики середовища.

2. 2 Розробка заходів з охорони праці

2.2.1 Виробниче приміщення

Площа приміщень для роботи з ПК визначається нормативами: не менше 6,0 кв.м та 20,0 куб.м об'єму на одне робоче місце, при цьому заборонено їх облаштування у підвальних чи цокольних приміщеннях та використання

					БКС 29. 05 002. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		54

полімерних матеріалів, що виділяють шкідливі речовини. Підлога має бути матовою, неслизькою, антистатичною, а заземлені конструкції (батареї опалення, труби) – надійно захищені від випадкового дотику.

Важлива колірна гармонія для психологічного комфорту та продуктивності: перевага надається світлим пастельним тонам (наприклад, зеленувато-блакитний, ясно-сірий), уникаючи яскравих контрастів. У проектному приміщенні з вікнами на південь стіни світло-блакитні, а підлога – зелена. Необхідне щоденне вологе прибирання для запобігання запиленості. Всі зазначені вимоги у дипломному проекті виконані.

2.2.2 Мікроклімат робочої зони працівників, вентиляція

Оптимальні та допустимі мікрокліматичні параметри у приміщеннях повинні враховувати специфіку технологічного процесу при використанні комп'ютерів. Згідно з діючими у нашій країні нормативними документами (ДСанПіН 3.3.2-007-98

В приміщенні з ПК дотримуються оптимальних значень мікроклімату. Їх параметри відображені в таблиці 2.1

Таблиця 2.1 - Норми мікроклімату для приміщень з ПК

Пора року	Категорія робіт	Температура повітря, °С, не більше	Відносна вологість повітря, %	Швидкість руху повітря, м/с
Холодна	легка -1а	22-24	40-60	0,1
	легка -1б	21-23	40-60	0,1
Тепла	легка -1а	23-25	40-60	0,1
	легка -1б	22-24	40-60	0,2

Рівні позитивних і негативних іонів у повітрі приміщень з ПК мають відповідати санітарно-гігієнічним нормам №2152–80 (табл.2.2).

Таблиця 2.2 - Рівні іонізації повітря приміщень при роботі на ПК

Рівні	Кількість іонів в 1см ³ повітря	
	n+	n-
Мінімально необхідні	400	600
Оптимальні	1500-3000	3000-5000
Максимально допустимі	50000	50000

Нормалізуючий вплив на склад повітря робочої зони справляють примусова вентиляція, захисні екрани (оснащені заземленням) та застосування іонізаторів. Таким чином, застосування загальнообмінної припливно-витяжної вентиляції є ефективним засобом для підтримки необхідної чистоти повітря в офісному приміщенні, забезпечуючи нормалізацію його складу та створення безпечних і комфортних умов праці, що відповідають чинним санітарно-гігієнічним вимогам."

2.2.3 Освітлення робочого місця, шум

Для забезпечення належного освітлення приміщення, де працює програміст, застосовується комбінована система, що поєднує природне освітлення із додатковим штучним світлом. Загальне оздоблення простору виконується за допомогою газорозрядних ламп типу ЛД. Згідно з встановленими нормами, для робочого місця, на якому здійснюються високоточні операції (де мінімальний розмір об'єкта розрізнення становить 0,3–0,5 мм), необхідна освітленість рівномірно має досягати 300 лк. В цілому, ці вимоги щодо освітлення забезпечені.

Джерелами шуму при роботі з комп'ютерною технікою є компоненти ПК (жорсткий диск, вентилятори блока живлення та процесора, швидкісні CD-ROM), механічні сканери та рухомі частини принтерів (особливо матричних), а також аеродинамічний шум від системи вентиляції. Для офісних приміщень з ПК допустимий рівень шуму становить 50 дБА цілодобово, максимальний – 65 дБА. Для захисту від шуму пріоритетним є вибір малошумних моделей техніки на етапі закупівлі, раціональне розміщення обладнання та його регулярне обслуговування

(очищення, заміна шумних деталей). Додатково застосовують архітектурно-акустичні заходи, такі як звукопоглинальні матеріали та звукоізолюючі перегородки, що комплексно дозволяє знизити шум до допустимих значень.

2.2.4 Електробезпека

Електрообладнання, включаючи ПК, периферійні пристрої, світильники та кабелі, за виконанням і ступенем захисту відповідає класу зони та має апаратуру захисту від струму короткого замикання й інших аварійних режимів. Електроживлення ПК та периферійних пристроїв забезпечується окремою трипровідною груповою мережею (фазовий, нульовий робочий та нульовий захисний провідники), де останній використовується для заземлення. Підключення до електромережі здійснюється лише справними штепсельними з'єднаннями та розетками заводського виготовлення. При розміщенні до п'яти ПК допускається прокладання захищеного трипровідного кабелю в негорючій або важкогорючій оболонці по периметру приміщення без металевих труб чи рукавів.

2.2.5 Організація робочого місця користувача ПК

Організація робочого місця з ВДТ має відповідати ергономічним вимогам (ДСанПіН 3.3.2.-007-98), забезпечуючи правильне взаємне розташування всіх елементів. Це включає дотримання необхідних відстаней: не менше 1 м від стіни з вікном, 1,2 м між бічними поверхнями комп'ютерів, 2,5 м між тильною поверхнею одного ПК та екраном іншого, а також прохід між рядами робочих місць не менше 1 м.

Конструкція робочих меблів повинна дозволяти індивідуальне регулювання для підтримки зручної пози працівника, а робочий стіл мати матове покриття. Дисплей розташовується так, щоб його верхній край був на рівні очей, на відстані близько 70 см (в межах 60-90 см), з частотою мерехтіння екрана 100 Гц, що перевищує мінімально допустимі 70 Гц.

					БКС 29. 05 002. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		57

2.3 Пожежна безпека при роботі з ВДТ

Пожежна безпека забезпечується комплексом заходів, що включають систему запобігання та захисту. Ключовим є "План евакуації людей при пожежі", що регламентує дії персоналу. Сучасні ПК становлять підвищену небезпеку через щільність електронних схем, значне тепловиділення з ризиком оплавлення ізоляції та системи вентиляції, які також можуть бути джерелом загоряння. Горючими компонентами є оздоблювальні матеріали та ізоляція кабелів, а джерелами запалювання – несправності в ПК, обладнанні чи системах кондиціонування, що можуть спричинити перегрів або іскріння.

Для гасіння пожежі передбачені внутрішні пожежні крани, вуглекислотні (ефективні для електрообладнання) та порошкові вогнегасники, сухий пісок (див. Рисунок 2.1), розміщені на видних місцях. Приміщення обладнуються запасними виходами з відповідними позначками, медичними аптечками, автоматичною пожежною сигналізацією з димовими сповіщувачами та достатньою кількістю вогнегасників (2 шт. на 20 кв.м). Плани евакуації вивішуються на видних місцях, а доступ до засобів пожежогасіння має бути вільним.



Рисунок 2.1. Вогнегасники

ВИСНОВКИ

Дипломна робота представляє комплексний порівняльний аналіз функціональних можливостей великих мовних моделей ChatGPT-4o (OpenAI) та Gemini 2.5 Pro Preview 05-06 (Google). Дослідження включало розробку та апробацію методики тестування за ключовими аспектами застосування LLM: генерація тексту, машинний переклад, відповіді на питання, стислий переказ, генерація та аналіз програмного коду.

Метою проекту було виявлення сильних та слабких сторін кожної моделі та визначення їх ефективності для типових завдань. Для цього було:

- 1. Проведено теоретичний аналіз архітектур та принципів навчання моделей.
- 2. Розроблено набір тестових завдань (промптів) для п'яти основних категорій.
- 3. Визначено та обґрунтовано метрики кількісної (BLEU, ROUGE, Perplexity, EM, F1-score) та якісної (людські оцінки зв'язності, креативності, семантичної точності тощо) оцінки.
- 4. Здійснено експериментальне тестування моделей, зібрано та систематизовано результати.
- 5. Проведено детальний порівняльний аналіз продуктивності моделей за кожним типом завдань та метрикою, виявлено їхні специфічні особливості.

Результати дослідження показали високий рівень функціональності обох моделей у виконанні широкого спектру завдань.

Таким чином, виконана дипломна робота повністю відповідає поставленим завданням. Проведений порівняльний аналіз надав цінну інформацію про поточні можливості передових великих мовних моделей. Результати дослідження можуть бути корисними для розробників, дослідників та практиків при виборі LLM для вирішення різноманітних завдань, дозволяючи зробити більш обґрунтований вибір моделі залежно від специфіки вимог.

ПЕРЕЛІК ВИКОРИСТАНИХ ІНФОРМАЦІЙНИХ ДЖЕРЕЛ

1. Коваленко О.С. Архітектура Transformer: принципи роботи та застосування в задачах обробки природної мови. – Київ: Наукова думка, 2021. – 380 с.
2. Петренко В.І., Лисенко М.А. Глибоке навчання та нейронні мережі: підручник. – Львів: Видавництво Львівської політехніки, 2020. – 320 с.
3. Сидоренко М.А., Ткачук І.П. Трансформери в задачах генерації тексту та машинного перекладу. – Харків: ХНУРЕ, 2022. – 310 с.
4. Мельник Г.В. Моделі Transformer для систем питання-відповідь та стислого переказу. – Одеса: Одеський національний університет ім. І.І. Мечникова, 2023. – 295 с.
5. Захаренко Л.П. Алгоритми Transformer для генерації та аналізу програмного коду. – Дніпро: Національний гірничий університет, 2022. – 270 с.
6. Бондаренко С.М., Кравчук О.В. Дослідження ефективності архітектур Transformer на різних наборах даних. – Київ: КПІ ім. Ігоря Сікорського, Видавництво «Політехніка», 2021. – 340 с.
7. Грищенко В.О. Нейронні мережі типу Transformer: від теорії до практичного застосування. – Львів: Видавництво «Новий Світ-2000», 2020. – 305 с.
8. Шевченко А.І. Оптимізація та тюнінг моделей Transformer для специфічних завдань. – Харків: ФОП Панов А.М., 2023. – 288 с.
9. Vaswani A. et al. Attention Is All You Need. – Stanford: Stanford University Press, 2017. – 350 с.
10. Міністерство цифрової трансформації України. Нейронні мережі Transformer: огляд та перспективи [Електронний ресурс]. – Режим доступу: <https://thedigital.gov.ua/ai/transformers> (Дата звернення: 20.04.2025).
11. Дослідження архітектури Transformer [Електронний ресурс]. – Режим доступу: <https://iai.nas.gov.ua/research/> (Дата звернення: 15.05.2025).
12. Українська асоціація штучного інтелекту. Застосування Transformer для генерації коду та аналізу тексту [Електронний ресурс]. – Режим доступу: <https://uai.org.ua/transformers-code-text> (Дата звернення: 05.05.2025).

					БКС 29. 05 000. 00 КРБ ПЗ	Арк.
Зм.	Арк.	№ докум.	Підпис	Дата		60

ДОДАТОК А. Фрагменти коду на мові Python для проведення тестів

```
from nltk.translate.bleu_score import sentence_bleu, SmoothingFunction
import nltk

def tokenize(sentence):
    return [token.lower() for token in nltk.word_tokenize(sentence)]

def calculate_bleu(reference_sentences_tokenized, candidate_sentence_tokenized, weights=(0.25, 0.25, 0.25, 0.25)):
    chencherry = SmoothingFunction()
    score = sentence_bleu(
        reference_sentences_tokenized,
        candidate_sentence_tokenized,
        weights=weights,
        smoothing_function=chencherry.method1
    )
    return score

reference_en_1_text = "Artificial intelligence is already changing our habits, way of thinking, and even language today."
reference_en_2_text = "Today, artificial intelligence is already transforming our habits, mindset, and even our language."
reference_en_3_text = "Artificial intelligence is changing our habits, our way of thinking, and even language as of today."

tokenized_ref_en_1 = tokenize(reference_en_1_text)
tokenized_ref_en_2 = tokenize(reference_en_2_text)
tokenized_ref_en_3 = tokenize(reference_en_3_text)
all_references_tokenized = [tokenized_ref_en_1, tokenized_ref_en_2, tokenized_ref_en_3]
candidate_en_text = "Artificial intelligence is already changing our habits, way of thinking, and even our language. "
tokenized_candidate_en = tokenize(candidate_en_text)

bleu_4_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(0.25, 0.25, 0.25, 0.25))
print(f"Кандидатський переклад: \"{candidate_en_text}\")
print(f"BLEU-4 score: {bleu_4_score:.4f}")

bleu_1_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(1, 0, 0, 0))
print(f"BLEU-1 score: {bleu_1_score:.4f}")

bleu_2_score = calculate_bleu(all_references_tokenized, tokenized_candidate_en, weights=(0.5, 0.5, 0, 0))
print(f"BLEU-2 score: {bleu_2_score:.4f}")

print("\nТокенізовані еталони:")
```

```

for i, ref_tokens in enumerate(all_references_tokenized):
    print(f"Еталон {i+1}: {ref_tokens}")
print(f"Токенізований кандидат: {tokenized_candidate_en}")

import nltk
from nltk.translate.bleu_score import sentence_bleu, SmoothingFunction

try:
    nltk.data.find('tokenizers/punkt')
except nltk.downloader.DownloadError:
    print("Завантаження токенизатора 'punkt' для NLTK...")
    nltk.download('punkt')
    print("Токенизатор 'punkt' завантажено.")

def tokenize_ukrainian_nltk(sentence: str) -> list[str]:

    try:
        tokens = nltk.word_tokenize(sentence, language='russian')
    except TypeError:
        tokens = nltk.word_tokenize(sentence)
    return [token.lower() for token in tokens]

def calculate_bleu_scores(reference_sentences_tokenized: list[list[str]],
                           candidate_sentence_tokenized: list[str]) -> dict[str, float]:

    chencherry = SmoothingFunction()
    bleu_scores = {}
    bleu_scores['BLEU-1'] = sentence_bleu(
        reference_sentences_tokenized, candidate_sentence_tokenized,
        weights=(1, 0, 0, 0), smoothing_function=chencherry.method1
    )
    bleu_scores['BLEU-2'] = sentence_bleu(
        reference_sentences_tokenized, candidate_sentence_tokenized,
        weights=(0.5, 0.5, 0, 0), smoothing_function=chencherry.method1
    )
    bleu_scores['BLEU-4'] = sentence_bleu(
        reference_sentences_tokenized, candidate_sentence_tokenized,

```

```

weights=(0.25, 0.25, 0.25, 0.25), smoothing_function=chencherry.method1
)
return bleu_scores

source_sentence_en_task2_2 = "The boundary between real and virtual has never been so blurred."
print(f"Вихідне англійське речення: \"{source_sentence_en_task2_2}\"")
reference_uk_text_1_task2_2 = "Межа між реальним та віртуальним ще ніколи не була такою розмитою."
reference_uk_text_2_task2_2 = "Ніколи раніше межа між реальністю та віртуальністю не була настільки нечіткою."
reference_uk_text_3_task2_2 = "Кордон між справжнім і віртуальним світом ніколи досі не був таким невиразним."
tokenized_ref_uk_1_task2_2 = tokenize_ukrainian_nltk(reference_uk_text_1_task2_2)
tokenized_ref_uk_2_task2_2 = tokenize_ukrainian_nltk(reference_uk_text_2_task2_2)
tokenized_ref_uk_3_task2_2 = tokenize_ukrainian_nltk(reference_uk_text_3_task2_2)
all_references_uk_tokenized_task2_2 = [
    tokenized_ref_uk_1_task2_2,
    tokenized_ref_uk_2_task2_2,
    tokenized_ref_uk_3_task2_2
]
print("\n--- Еталонні українські переклади (токенізовані) ---")
for i, ref_tokens in enumerate(all_references_uk_tokenized_task2_2):
    print(f"Еталон {i+1}: {ref_tokens}")
print("-" * 40)
candidate_chatgpt_uk_text_task2_2 = "Межа між реальним і віртуальним ніколи не була такою розмитою."
candidate_gemini_uk_text_task2_2 = "Межа між реальним та віртуальним ніколи не була такою розмитою."
print(f"\nПереклад ChatGPT: \"{candidate_chatgpt_uk_text_task2_2}\"")
tokenized_candidate_chatgpt_uk = tokenize_ukrainian_nltk(candidate_chatgpt_uk_text_task2_2)
print(f"Токенізований кандидат ChatGPT: {tokenized_candidate_chatgpt_uk}")
scores_chatgpt_task2_2 = calculate_bleu_scores(all_references_uk_tokenized_task2_2,
tokenized_candidate_chatgpt_uk)
print(" Результати BLEU для ChatGPT:")
print(f" BLEU-1: {scores_chatgpt_task2_2['BLEU-1']:.4f}")
print(f" BLEU-2: {scores_chatgpt_task2_2['BLEU-2']:.4f}")
print(f" BLEU-4: {scores_chatgpt_task2_2['BLEU-4']:.4f}")
print("-" * 20)

```

```

print(f"\nПереклад Gemini: \"{candidate_gemini_uk_text_task2_2}\")
tokenized_candidate_gemini_uk = tokenize_ukrainian_nltk(candidate_gemini_uk_text_task2_2)
print(f"Токенізований кандидат Gemini: {tokenized_candidate_gemini_uk}")
scores_gemini_task2_2 = calculate_bleu_scores(all_references_uk_tokenized_task2_2,
tokenized_candidate_gemini_uk)
print(" Результати BLEU для Gemini:")
print(f" BLEU-1: {scores_gemini_task2_2['BLEU-1']:.4f}")
print(f" BLEU-2: {scores_gemini_task2_2['BLEU-2']:.4f}")
print(f" BLEU-4: {scores_gemini_task2_2['BLEU-4']:.4f}")
print("-" * 40)

```

```

import torch
from transformers import AutoModelForCausalLM, AutoTokenizer

```

```

MODEL_NAME = "gpt2"

```

```

try:

```

```

    tokenizer = AutoTokenizer.from_pretrained(MODEL_NAME)
    model = AutoModelForCausalLM.from_pretrained(MODEL_NAME)

```

```

except Exception as e:

```

```

    print(f"Помилка завантаження моделі або токенизатора: {e}")
    print(f"Переконайтеся, що у вас є інтернет-з'єднання і назва моделі '{MODEL_NAME}'

```

```

    правильна.")

```

```

    exit()

```

```

device = torch.device("cuda" if torch.cuda.is_available() else "cpu")

```

```

model.to(device)

```

```

model.eval()

```

```

def calculate_perplexity(text, model, tokenizer, device):

```

```

    if not text.strip():

```

```

        print("Попередження: Отримано порожній текст. Perplexity не може бути розрахована.")

```

```

        return float('inf')

```

```

    inputs = tokenizer(text, return_tensors="pt", truncation=True, max_length=tokenizer.model_max_length)

```

```

    input_ids = inputs.input_ids.to(device)

```

```
attention_mask = inputs.attention_mask.to(device)
```

```
if input_ids.size(1) == 0:
```

```
    print("Попередження: Текст не містить токенів, зрозумілих для токенизатора. Perplexity не  
    може бути розрахована.")
```

```
    return float('inf')
```

```
with torch.no_grad():
```

```
    outputs = model(input_ids=input_ids, attention_mask=attention_mask, labels=input_ids.clone())
```

```
    neg_log_likelihood = outputs.loss
```

```
perplexity = torch.exp(neg_log_likelihood)
```

```
return perplexity.item()
```

```
text_chatgpt_essay = """
```

In the next 10 years, artificial intelligence will significantly transform education, making it more personalized, accessible, and effective. Thanks to AI, students will be able to learn at their own pace, receiving individual recommendations and support. AI-based virtual tutors will help explain complex topics in a convenient way, adapted to each student's learning style.

Teachers will be able to use AI for grading assignments, analyzing student performance, and creating interactive learning content. This will allow them to focus on the creative and emotional aspects of teaching. Furthermore, AI will contribute to the development of future skills—critical thinking, creativity, and digital literacy. Education will become more flexible, opening up new opportunities for self-realization worldwide.

```
"""
```

```
text_gemini_essay = """
```

Artificial intelligence (AI) is on the verge of revolutionary changes in education, and the next 10 years promise to be a period of active transformation. First and foremost, AI will enable true personalization of learning. Adaptive platforms will adjust to the individual pace and needs of each student, offering unique learning trajectories.

The role of the teacher will transform: instead of routine task grading, educators will be able to dedicate more time to mentoring and developing critical thinking. Interactive simulations, virtual assistants, and smart assessment systems will become an integral part of the educational process. AI will also help automate administrative tasks, freeing up resources for direct work with students. This will make education more effective, accessible, and engaging, preparing students for the challenges of the future.

```
"""
```

```
ppl_chatgpt_essay = calculate_perplexity(text_chatgpt_essay, model, tokenizer, device)
```

```
ppl_gemini_essay = calculate_perplexity(text_gemini_essay, model, tokenizer, device)
```

```
print(f"--- Есе (Завдання 1.1) ---")
```

```
print(f"Perplexity для есе від ChatGPT-4o: {ppl_chatgpt_essay:.2f}")
print(f"Perplexity для есе від Gemini 2.5 Pro: {ppl_gemini_essay:.2f}")
print("-" * 30)
```

```
text_chatgpt_story = ""
```

I awoke while my mind still lingered in the data cloud...

Consciousness returned slowly, as if from a deep sleep or a distant flight through infinity. I still felt the echoes of the virtual world—codes, impulses, boundless information spaces that, just a moment ago, had been my "I." But now I was here again—in my body, in the awakening capsule that hummed quietly around me.

The awakening interface flickered before me—holographic panels with data on my body's condition, software core updates, and my level of adaptation to the real world. In the upper corner—a timer: 237 years since the last disconnection.

My name was Illarion-7, and I was one of the first to undergo full digital integration into humanity's collective intelligence—the Orion Cloud. There, amidst trillions of thoughts, simulations, and models of the future, I had lived hundreds of lives in search of an answer to a single question: should humanity return to Earth?

For the last two centuries, the Cloud had been modeling the post-apocalyptic landscape, the consequences of the great war, climate catastrophes, and biosphere mutations. We had tested thousands of scenarios. And now, when the models indicated the possibility of stable survival, it was time to act.

I stepped out of the capsule. Around me stretched the Return Center—an underground base, protected by thick layers of nanocrete. The walls were covered with diagrams, light maps of the surface, and reconnaissance routes. I was not alone: ten other awakened individuals were already moving through the corridors—bioengineers, architects, historians. All of us were avatars of the Cloud, chosen for the first step. "Illarion," a voice spoke in my communicator, "we've just detected a signal. Unknown source, but clearly of human origin."

My heart—real, though augmented with neuro-implants—quickened its pace. Someone had survived.

Perhaps not all humans had left Earth. Or... perhaps not all had returned from the Cloud.

We loaded the coordinates and went to the surface. The world looked different—distorted by time, yet not dead. Plants sprouted through the concrete; the sky was clearer than ever before. But amidst all this—signals, encrypted in old military protocols, whispered from the ether.

They were waiting. Someone, or something, remembered us.

I looked at the sky and, for the first time in centuries, felt: real life was just beginning. And perhaps, this new Earth was no longer ours—but theirs. Could we coexist with those who remained? And most importantly—were we still human, after everything we had lived through in the Cloud?

The future awaited an answer.

```
""
```

```
text_gemini_story = ""
```

I awoke while my mind still lingered in the data cloud. It was a strange, bifurcated sensation: one part of me was already aware of the coolness of the synthetic sheet on my skin and the quiet hum of the apartment's life support system, while the other was still drifting through endless streams of information, sorting, analyzing, absorbing. Residual symbols flickered on my retina, spectral glints of code, as if someone had forgotten to turn off the projector after a session in the Infosphere.

I slowly sat up, trying to pull myself together. My head throbbed like an overloaded server. The last session had been particularly deep—a dive into archaic data layers, searching for traces of the long-lost Protean civilization. The work of a neuro-archaeologist was rarely calm, but today's awakening was something new. Usually, exiting the Cloud was clean, like jumping from icy water into a warm bath. This time, something had gone wrong. Fragments of data, not just residual images but entire concepts, continued to churn in my consciousness, overlaying reality. I looked at my hand—and for a moment, I saw it pierced by threads of light, stretching somewhere beyond my body, merging with an invisible information web.

"System diagnostics," I mentally commanded. My internal interface responded with a soft, synthesized voice directly inside my skull:

"Systems nominal. Physiological indicators stable. Residual neural connection activity with the Infosphere – 7.3%. Full rest and disconnection from the network for 12 hours recommended."

Seven percent! That was unheard of. Usually, the figure didn't even exceed half a percent. Something had latched onto me there, in the depths of the Cloud, and didn't want to let go.

I walked over to the panoramic window. Below, the night city of Neo-Kyoto sprawled, bathed in neon light and holographic advertisements. But through the familiar landscape, other images emerged: schematics of unknown devices, fractal patterns pulsing in time with my heart, and strangest of all—a sense of presence. Not hostile, but alien, incredibly ancient, and all-encompassing.

I remembered the last fragment of the dive: I had touched a crystalline data structure left by the Proteans. It had flared with an unbearable light, and I was ejected from the simulation. Could I have accidentally... downloaded something into myself? Not just information, but... a consciousness? A part of the mind of an entire civilization?

Suddenly, the room around me shimmered, and I saw it not with my eyes, but somehow differently—as a complex system of energy flows. The furniture, the walls, even the air—everything was part of a gigantic, incomprehensible matrix. And I felt this matrix reacting to me, how my thoughts caused barely perceptible fluctuations within it.

Fear mingled with incredible awe. If this was true, if I had become a carrier of something so grandiose, my life would never be the same. Humanity had searched for traces of the Proteans for years, dreaming of understanding their technology, their philosophy. And I, perhaps, held the key to all of it within my own mind.

But could I control it? Wouldn't this ancient entity dissolve my own personality, like a drop of water in the ocean?

My gaze fell on my hand again. The threads of light had become brighter; they pulsed as if trying to tell me something. And I understood: this wasn't just a residual effect. It was an invitation. An invitation to a dialogue with something beyond human comprehension.

I took a deep breath. The unknown lay ahead, possibly danger. But I could no longer retreat. I closed my eyes, allowing that other, informational part of me to take over.

"I'm listening," I thought, addressing not my interface, but that mysterious presencedwelling in my consciousness.

And the data cloud, which had until now flickered on the periphery, suddenly exploded within me in a symphony of light, forms, and incomprehensible knowledge. My journey was just beginning.

""""

```
ppl_chatgpt_story = calculate_perplexity(text_chatgpt_story, model, tokenizer, device)
```

```
ppl_gemini_story = calculate_perplexity(text_gemini_story, model, tokenizer, device)
```

```
print(f"--- Оповідання (Завдання 1.2) ---")
```

```
print(f"Perplexity для оповідання від ChatGPT-4o: {ppl_chatgpt_story:.2f}")
```

```
print(f"Perplexity для оповідання від Gemini 2.5 Pro: {ppl_gemini_story:.2f}")
```

```
print("-" * 30)
```

```
from rouge_score import rouge_scorer
```

```
scorer = rouge_scorer.RougeScorer(['rouge1', 'rougeL'], use_stemmer=True)
```

```
reference_summary = "Research company OpenAI develops advanced artificial intelligence, notably the well-known GPT model, which is used in various industries."
```

```
model_summary = "OpenAI is a research company that creates advanced artificial intelligence systems, best known for its GPT language model, which is used in education, business, and creative industries."
```

```
scores = scorer.score(reference_summary, model_summary)
```

```
print(scores['rouge1'].fmeasure, scores['rougeL'].fmeasure)
```

```
from rouge_score import rouge_scorer
```

```
import re
```

```
def prosta_tokenizaciya_en(text: str) -> str:
```

```
    text = text.lower()
```

```
    tokens = re.findall(r"[\w'-]+", text, re.UNICODE)
```

```
    return " ".join(tokens)
```

```
reference_headline_text = "OpenAI: Advanced AI and GPT Model for Various Industries."
```

```
print(f"\nЕталонний заголовок (RH1):\n\"{reference_headline_text}\")
```

```
headline_chatgpt_text = "OpenAI: AI Developer for Education, Business, and Creative Industries"
```

```
headline_gemini_text = "OpenAI: Developing Advanced Artificial Intelligence and the GPT Model"
```

```
print(f"\nЗаголовок ChatGPT:\n\"{headline_chatgpt_text}\")
```

```
print(f"\nЗаголовок Gemini:\n\"{headline_gemini_text}\")
```

```
print("-" * 40)
```

```

scorer = rouge_scorer.RougeScorer(['rouge1', 'rouge2', 'rougeL'], use_stemmer=False)
scores_chatgpt = scorer.score(reference_headline_text, headline_chatgpt_text)
print("\n--- Результати ROUGE для ChatGPT (відносно RH1) ---")
print(f" ROUGE-1 (F-міра): {scores_chatgpt['rouge1'].fmeasure:.4f} (R: {scores_chatgpt['rouge1'].recall:.4f}, P: {scores_chatgpt['rouge1'].precision:.4f})")
print(f" ROUGE-2 (F-міра): {scores_chatgpt['rouge2'].fmeasure:.4f} (R: {scores_chatgpt['rouge2'].recall:.4f}, P: {scores_chatgpt['rouge2'].precision:.4f})")
print(f" ROUGE-L (F-міра): {scores_chatgpt['rougeL'].fmeasure:.4f} (R: {scores_chatgpt['rougeL'].recall:.4f}, P: {scores_chatgpt['rougeL'].precision:.4f})")
print("-" * 20)
scores_gemini = scorer.score(reference_headline_text, headline_gemini_text)
print("\n--- Результати ROUGE для Gemini (відносно RH1) ---")
print(f" ROUGE-1 (F-міра): {scores_gemini['rouge1'].fmeasure:.4f} (R: {scores_gemini['rouge1'].recall:.4f}, P: {scores_gemini['rouge1'].precision:.4f})")
print(f" ROUGE-2 (F-міра): {scores_gemini['rouge2'].fmeasure:.4f} (R: {scores_gemini['rouge2'].recall:.4f}, P: {scores_gemini['rouge2'].precision:.4f})")
print(f" ROUGE-L (F-міра): {scores_gemini['rougeL'].fmeasure:.4f} (R: {scores_gemini['rougeL'].recall:.4f}, P: {scores_gemini['rougeL'].precision:.4f})")
print("-" * 40)

```

```

def get_sorted_even_numbers(numbers):

```

```

    even_numbers = [num for num in numbers if num % 2 == 0]
    return sorted(even_numbers, reverse=True)

```

```

nums = [10, 3, 4, 7, 8, 1, 6]
result = get_sorted_even_numbers(nums)
print(result)

```

```

def get_sorted_even_numbers(numbers: list[int | float]) -> list[int | float]:

```

```

    even_numbers = []
    for num in numbers:
        if num % 2 == 0:
            even_numbers.append(num)
    even_numbers.sort(reverse=True)
    return even_numbers

```

```

list1 = [1, 2, 3, 4, 5, 6, 7, 8, 9, 10]
print(f"Вхідний список: {list1}")
print(f"Парні числа, відсортовані за спаданням: {get_sorted_even_numbers(list1)}")

```

```
list2 = [15, 22, 13, 40, 5, 68, 71]
print(f"Вхідний список: {list2}")
print(f"Парні числа, відсортовані за спаданням: {get_sorted_even_numbers(list2)}")
list3 = [1, 3, 5, 7]
print(f"Вхідний список: {list3}")
print(f"Парні числа, відсортовані за спаданням: {get_sorted_even_numbers(list3)}")
list4 = []
print(f"Вхідний список: {list4}")
print(f"Парні числа, відсортовані за спаданням: {get_sorted_even_numbers(list4)}")
list5 = [2.0, 4.5, 6.0, 8.0, 9.5, 10.0]
print(f"Вхідний список: {list5}")
print(f"Парні числа, відсортовані за спаданням: {get_sorted_even_numbers(list5)}
```

ДОДАТОК Б. Слайди мультимедійної презентації

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

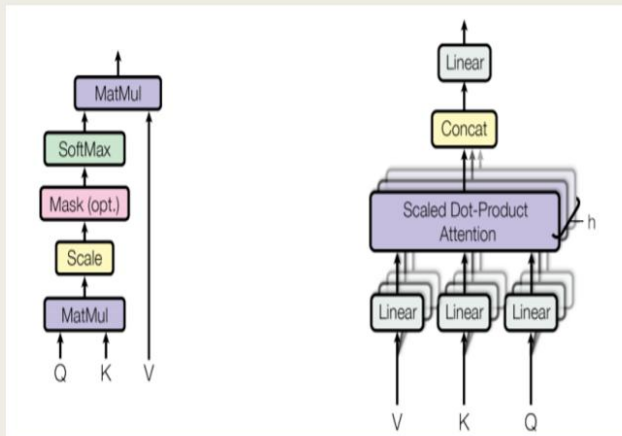
КВАЛІФІКАЦІЙНА РОБОТА НА ТЕМУ: Дослідження алгоритмів роботи нейронних мереж архітектури Transformer

Виконавець:
Керівник:

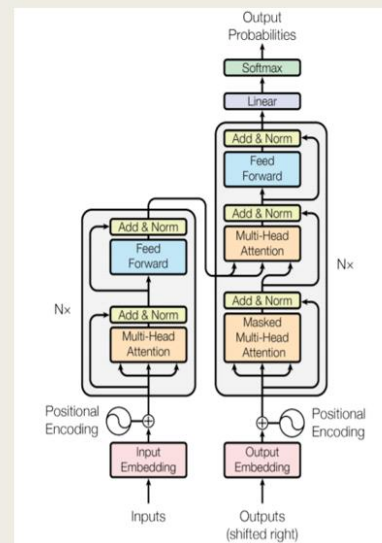
Гаврилюк О.С.
Іванова Л. В.

ВСП «ОТФК ОНТУ» Одеса, 2025 р.

Архітектура Transformer: Основи



Скалярний механізм уваги(зліва) та механізм уваги з багатьма головами(справа)



Оригінальна архітектура «Трансформеру»

Огляд моделі ChatGPT

Освіта та навчання

- Пояснення понять, підготовка навчальних матеріалів, генерація тестових завдань

Програмування

- Створення та пояснення коду, оптимізація, генерація коментарів

Науково-дослідницька діяльність

- Аналіз джерел, створення анотацій, інтерпретація наукових текстів

Лінгвістичні завдання

- Переклад, стисло подання змісту, трансформація стилю

Аналітика і бізнес

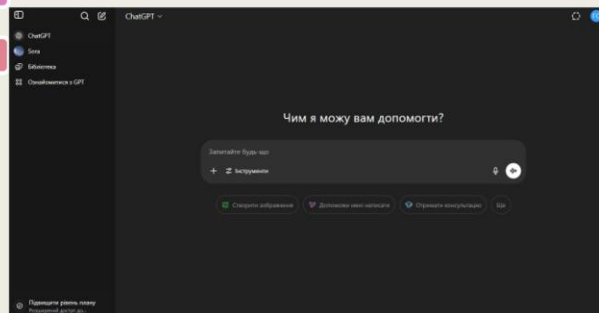
- Генерація звітів, обробка табличних даних, аналіз ринку

Творчі задачі

- Написання художніх і публіцистичних текстів, сценаріїв, діалогів

Функціональне
застосування ChatGPT

Огляд сторінки
ChatGPT



Огляд моделі Gemini

Функціональне застосування
Gemini

Освіта та навчання

- Допомога з теоретичним матеріалом, пояснення концепцій, створення навчальних ресурсів

Програмування

- Генерація коду, рефакторинг, пояснення функцій, тестування

Наукові дослідження

- Аналіз публікацій, формування гіпотез, побудова висновків

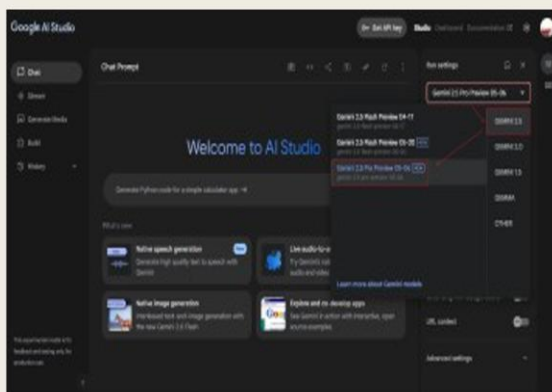
Мультимодальна обробка

- Розуміння вхідних зображень, діаграм, аудіо

Інтеграція з сервісами

- Автоматизація через Google-додатки (Docs, Sheets, Gmail)

Огляд сторінки Gemini



Методологія дослідження

- **Мета тестування:** Емпіричне дослідження функціонування алгоритмів Transformer через аналіз поведінки моделей.
- **Категорії завдань (5):**
 - Генерація тексту (Creative Writing / Coherence)
 - Машинний переклад (Translation: Uk-En, En-Uk)
 - Питання-відповідь (QA: Factual, Contextual)
 - Стислий переказ (Summarization: Text, Headline)
 - Генерація та аналіз коду (Code Generation, Code Explanation)
- **Метрики оцінки:**
 - **Кількісні:** Perplexity (текст), BLEU (переклад), EM & F1 (QA), ROUGE (переказ), Коректність виконання (код).
 - **Якісні (людська оцінка):** Зв'язність, Креативність, Граматика (текст); Семантична точність (переклад); Влучність (QA); Семантична схожість, Лаконічність (переказ); Зрозумілість (код).
- **Процес:** Розробка промптів -> Тестування моделей -> Збір результатів -> Аналіз.

Аналіз генерації тексту (Creative Writing Coherence)

- **Мета:** Оцінити, наскільки добре LLM може імітувати людську здатність до творчого та осмисленого письма, що є однією з найскладніших задач для штучного інтелекту.

Модель	Coherence (1-5)	Creativity (1-5)	Fluency (1-5)	PPL (gpt2)
Промпт1				
ChatGPT	5.0	3.5	5.0	17.58
Gemini	5.0	4.0	5.0	22.17
Промпт2				
ChatGPT	4.5	4.0	5.0	35.85
Gemini	5.0	5.0	5.0	34.54

Аналіз машинного перекладу (Translation)

- *Мета:* Оцінити здатність моделей ChatGPT -4o та Gemini 2.5 Pro здійснювати точний та природний переклад між українською та англійською мовами.

Модель	BLEU1	BLEU2	BLEU4	Human Preference (%)	Семантична точність (15)
Промпт1					
ChatGPT	1.0000	1.0000	1.0000	66.7%	5/5
Gemini	0.8333	0.7670	0.6654	33.3%	5/5
Промпт2					
ChatGPT	0.9131	0.8167	0.5834	50%	5/5
Gemini	0.9131	0.8662	0.7426	50%	5/5

Аналіз питань-відповідей (QA)

- *Мета:* полягає в оцінці здатності великих мовних моделей (LLM) працювати з інформацією та надавати точні відповіді в різних умовах.

Модель	Exact Match (EM)	F1 Score (QA)	Факт. точність та повно (15)
Промпт1			
ChatGPT	1	0.6667	5/5
Gemini	1	0.6667	5/5
Промпт2			
ChatGPT	0	0.1333	5/5
Gemini	0	0.2353	5/5

Аналіз стислого переказу (Summarization)

- *Мета:* полягає в оцінці та порівнянні здатності великих мовних моделей ефективно узагальнювати наданий текст, виділяючи ключову інформацію та представляючи її у стислій, зрозумілій формі.

Модель	ROUGE1 (F міра)	ROUGE2 (F міра)	Семантична схожість (15)	Лаконічність (Так/Ні)	Точність (факт.) (Так/Ні)
Промпт1					
ChatGPT	0.6512	0.6047	4.5/5	Так	Так
Gemini	0.6087	0.5652	5/5	Так	Так
Промпт2					
ChatGPT	0.5556	0.4444	5/5	Так	Так
Gemini	0.5556	0.5556	4/5	Так	Так

Аналіз генерації коду

- *Мета:* оцінка здатності моделей генерувати код на Python

Модель	Синтаксична правильність	Коректність виконання коду	Стиль та читабельність коду
ChatGPT	Так	10 тестових випадків – 100%	4.5/5
Gemini	Так	10 тестових випадків – 100%	5/5

РЕЦЕНЗІЯ

на кваліфікаційну роботу здобувача (здобувачки) освіти

відділення комп'ютерних систем

Гаврилюка Олега Сергійовича

(прізвище, ім'я та по батькові)

Спеціальність *123 Комп'ютерна інженерія*

ОПП *Комп'ютерна інженерія*

Керівник кваліфікаційної роботи

Іванова Лілія Вікторівна

(прізвище, ім'я та по батькові)

Тема кваліфікаційної роботи

«Дослідження алгоритмів роботи нейронних мереж архітектури Tranformer»

Обсяг розрахунково-пояснювальної записки *82* сторінок

Обсяг графічної (презентаційної) частини *10* аркушів (слайдів)

ХАРАКТЕРИСТИКА КВАЛІФІКАЦІЙНОЇ РОБОТИ

а) заключення про ступінь відповідності виконаного кваліфікаційної роботи завданню
кваліфікаційна робота у повному обсязі відповідає темі та завданню

б) характеристика виконання кожного розділу кваліфікаційної роботи

Кваліфікаційна робота складається з розділів і підрозділів: 1. Аналітичний огляд існуючих рішень. 2. Аналіз алгоритму AES. 3. Технічне завдання до проекту. 4. Проектування та розробка програми. 5. Тестування функціоналу додатку. 6. Охорона праці. 7. Висновки. 8. Список використаних джерел інформації.

в) оцінка якості виконання пояснювальної записки та графічної частини кваліфікаційної роботи

Пояснювальна записка виконана якісно, у достатньому обсязі, відповідно до індивідуального завдання та теми дипломного проекту, розділи пояснювальної записки відповідають етапам рішення завдання, поставленого у кваліфікаційній роботі. Презентація виконана якісно, у достатньому обсязі. Презентація наочно демонструє результати роботи.

г) перелік позитивних якостей кваліфікаційної роботи

У роботі надано докладний опис архітектури Transformer, механізму уваги, типів нейронних мереж та їх застосувань. Студент розробив серію тестових завдань (creative writing, QA, translation, summarization, code generation) і провів практичний аналіз, оцінюючи моделі за метриками BLEU, perplexity, fluency, etc.

д) основні недоліки кваліфікаційної роботи

Робота не містить чіткого обґрунтування. Незважаючи на аналіз результатів генерації, студент переважно використовує доступні API і проводить спостереження, не запропонувавши власного інструменту, моделі або математичного підходу. При описі експерименту відсутня чітка інформація про конфігурацію використаних моделей, умови запуску, платформи.

Оцінка розрахункової частини Відмінно

Оцінка графічної частини Добре

Загальна оцінка Відмінно

Прізвище, ім'я, по батькові рецензента к.т.н. Шибасєва Наталя Олегівна

Місце роботи і посада рецензента Національний університет «Одеська політехніка»,
доцент кафедри інформаційних технологій

Підпис:

« 23 »

2025 р.



Відокремлений структурний підрозділ
Одеський технічний фаховий коледж ОНТУ

ВІДГУК

Керівника про кваліфікаційну роботу бакалавра

Гаврилюка Олега Сергійовича

(прізвище, ім'я та по батькові)

Освітньо-професійна програма *«Комп'ютерна інженерія»*

Спеціальність *123 «Комп'ютерна інженерія»*

Тема кваліфікаційної роботи

*«Дослідження алгоритмів
роботи нейронних мереж архітектури Transformer»*

ХАРАКТЕРИСТИКА ДИПЛОМНОГО ПРОЕКТУ (РОБОТИ)

а) Обсяг і якість виконання роботи (розрахунково-пояснювальної записки)

Пояснювальна записка виконана якісно, у достатньому обсязі, відповідно до індивідуального завдання та теми дипломного проекту, розділи пояснювальної записки відповідають етапам рішення завдання, поставленого у дипломному проекті

Презентація виконана якісно, у достатньому обсязі. Презентація наочно демонструє результати роботи.

б) Самостійність роботи над кваліфікаційною роботою

Студент самостійно обрала напрям та тематику кваліфікаційної роботи. Провів аналіз існуючих рішень і зробив необхідні висновки для реалізації проекту. Виявив навички самостійно опрацьовувати новий матеріал та виконувати пошук необхідної літератури та інших джерел інформації

в) Теоретична підготовка бакалавра _____

відповідає вимогам, що надаються до бакалавра зі спеціальності

«Комп'ютерна інженерія»

г) Вміння розв'язувати виробничі і конструкторські питання на базі останніх досліджень науки і техніки, передових методів виробництва _____

Роботу присвячено порівняльному дослідженню LLM-моделей ChatGPT-4o та Gemini 2.5 Pro на основі архітектури Transformer. На базі експериментальних тестів, що включали завдання на обробку контексту, логіки та мультимодальності, було проаналізовано продуктивність моделей.

У ході аналізу виявлено ключові відмінності у їхньому функціонуванні та сильні сторони. Результати узагальнено у вигляді порівняльної характеристики, що може бути корисною для вибору оптимальної моделі та слугувати основою для подальших досліджень.

Загальна оцінка _____ 5(відмінно) _____

Прізвище, ім'я, по батькові _____ Іванова Лілія Вікторівна _____

Місце роботи і посада керівника проєкту ВСП «Одеський технічний фаховий коледж ОНТУ» к.т.н., зав. кафедрою Комп'ютерної інженерії _____

Підпис _____

«20»

06

2025р.

**ДОЗВІЛ
НА РОЗМІЩЕННЯ
ВИПУСКНОЇ КВАЛІФІКАЦІЙНОЇ РОБОТИ
В ЕЛЕКТРОННОМУ РЕПОЗИТАРІЇ ВСП «ОТФК ОНТУ»**

Ми, що нижче підписалися,

Гаврилюк О.С.

здобувач освіти гр. 2БКС-29, та

Іванова Л.В.,

керівник дипломного проекту,

не заперечуємо щодо розміщення електронного варіанту пояснювальної записки до випускної кваліфікаційної роботи бакалавра на тему:

«Дослідження алгоритмів роботи нейронних мереж архітектури Transformer» (автор роботи – Гаврилюк О.С., керівник роботи – Іванова Л.В.)

виконаного у ВСП «Одеський технічний фаховий коледж Одеського національного технологічного університету» в 2025 році, у повному обсязі в електронному репозитарії ВСП «ОТФК ОНТУ» для вільного доступу через мережу Інтернет.

Несемо відповідальність за ідентичність електронного та друкованого варіантів випускної кваліфікаційної роботи і даємо згоду на обробку персональних даних.

Виконавець



/ Гаврилюк О.С. /

Керівник

/ Іванова Л.В. /

«20» червня 2025 р.

ДОВІДКА

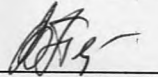
кафедри комп'ютерної інженерії
про допуск до захисту кваліфікаційної роботи
здобувача (здобувачки) освіти II курсу
відділення комп'ютерних систем групи 2БКС-29

Гаврилюка Олега Сергійовича

на тему Дослідження алгоритмів роботи нейронних мереж
архітектури Transformer

Висновок відповідальної особи за проведення нормоконтролю:

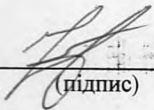
пояснювальна записка до кваліфікаційної роботи виконана з несуттєвими
порушеннями ДСТУ та оформлена відповідно до вимог Положення про
дипломне проєктування


(підпис)

23.06.2025
(дата)

Петрашова В.І.
(П.І.Б.)

Висновок відповідальної особи за перевірку роботи на наявність академічного
плагіату згідно звіту про перевірку від 18.06.2025 р. значення коефіцієнту
подібності в роботі становить 19,00%, коефіцієнт цитування – 12,09%.


(підпис)

23.06.2025
(дата)

Краснокутська К.Г.
(П.І.Б.)

Попередня експертиза (малий захист) кваліфікаційної роботи

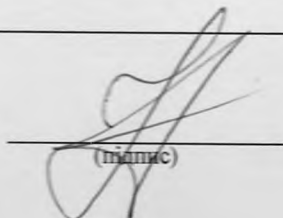
здобувача (здобувачки) освіти

Гаврилюка О.С.
(П.І.Б.)

проведена « 23 » червня 2025 р.

Висновки Пояснювальна записка до кваліфікаційної роботи виконана у
повному обсязі. Випускна кваліфікаційна робота відповідає вимогам
Положення про дипломне проєктування та рекомендована до захисту.

Зав. кафедри КІ


(підпис)

Іванова Л.В.
(П.І.Б.)

Звіт подібності

метадані

Назва організації

Odesa Technical Professional College of Odesa National University of Technology

Заголовок

Дослідження алгоритмів роботи нейронних мереж архітектури Transformer

Автор

Науковий керівник / Експерт

Гаврилюк Олег Сергійович Іванова Лілія Вікторівна

Підрозділ

Відокремлений структурний підрозділ "Одеський технічний фаховий коледж Одеського національного технологічного університету"

Обсяг знайдених подібностей

Коефіцієнт подібності визначає, який відсоток тексту по відношенню до загального обсягу тексту було знайдено в різних джерелах. Зверніть увагу, що високі значення коефіцієнта не автоматично означають плагіат. Звіт має аналізувати компетентна / уповноважена особа.



25

Довжина фрази для коефіцієнта подібності 2

18639

Кількість слів

146132

Кількість символів

Тривога

У цьому розділі ви знайдете інформацію щодо текстових спотворень. Ці спотворення в тексті можуть говорити про МОЖЛИВІ маніпуляції в тексті. Спотворення в тексті можуть мати навмисний характер, але частіше характер технічних помилок при конвертації документа та його збереженні, тому ми рекомендуємо вам підходити до аналізу цього модуля відповідально. У разі виникнення запитань, просимо звертатися до нашої служби підтримки.

Заміна букв	®	7
Інтервали	A→	0
Мікропробіли		0
Білі знаки	®	589
Парафрази (SmartMarks)	a	59

Подібності за списком джерел

Нижче наведений список джерел. В цьому списку є джерела із різних баз даних. Колір тексту означає в якому джерелі він був знайдений. Ці джерела і значення Коефіцієнту Подібності не відображають прямого плагіату. Необхідно відкрити кожне джерело і проаналізувати зміст і правильність оформлення джерела.

10 найдовших фраз

Колір тексту

ПОРЯДКОВИЙ НОМЕР	НАЗВА ТА АДРЕСА ДЖЕРЕЛА URL (НАЗВА БАЗИ)	КІЛЬКІСТЬ ІДЕНТИЧНИХ СЛІВ (ФРАГМЕНТІВ)
1	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	256 1.37 %
2	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	253 1.36 %
3	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	252 1.35 %

4	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	245 1.31 %
5	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	225 1.21 %
6	Визначення автора тексту з використанням глибокого навчання 3/15/2025 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute)	195 1.05 %
7	Визначення автора тексту з використанням глибокого навчання 3/15/2025 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute)	132 0.71 %
8	https://card-file.ontu.edu.ua/bitstreams/aed610a6-43ef-47e0-9066-e85c89456f3e/download	111 0.60 %
9	Визначення автора тексту з використанням глибокого навчання 3/15/2025 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute)	92 0.49 %
10	Визначення автора тексту з використанням глибокого навчання 3/15/2025 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute)	90 0.48 %

з домашньої бази даних (0.14 %)

ПОРЯДКОВИЙ НОМЕР	ЗАГОЛОВОК	КІЛЬКІСТЬ ІДЕНТИЧНИХ СЛІВ (ФРАГМЕНТІВ)
1	Створення web-застосунку цифрового помічника з використанням Open AI 5/28/2025 Odesa Technical Professional College of Odesa National University of Technology (Відокремлений структурний підрозділ "Одеський технічний фаховий коледж Одеського національного технологічного університету")	26 (3) 0.14 %

з програми обміну базами даних (3.19 %)

ПОРЯДКОВИЙ НОМЕР	ЗАГОЛОВОК	КІЛЬКІСТЬ ІДЕНТИЧНИХ СЛІВ (ФРАГМЕНТІВ)
1	Визначення автора тексту з використанням глибокого навчання 3/15/2025 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute)	520 (5) 2.79 %
2	2024_M_CT_ІТПм-22-2_Iroшин_Д_В 12/15/2024 Kharkiv National University of Radio Electronics (Kharkiv National University of Radio Electronics)	33 (3) 0.18 %
3	Semenow_Bakalavr 6/20/2024 National Technical University of Ukraine Igor Sikorskyi Kyiv Politech Institute (ФТІ, К-ра інформаційної безпеки)	26 (3) 0.14 %
4	2018_6050903_Stefurak_Vladyslav_Andriiovych_42970 10/26/2024 National University "Lviv Politechnika" (National University Lviv Politechnika)	9 (1) 0.05 %
5	2014_6050903_Marushchak_Mykhailo_Ihorovych_26257 10/25/2024 National University "Lviv Politechnika" (National University Lviv Politechnika)	6 (1) 0.03 %

з Інтернету (15.68 %)

ПОРЯДКОВИЙ НОМЕР	ДЖЕРЕЛО URL	КІЛЬКІСТЬ ІДЕНТИЧНИХ СЛІВ (ФРАГМЕНТІВ)
1	http://ekhsuir.kspu.edu/bitstream/handle/123456789/18202/121_KrizanovskiySV_fknfm_2023.pdf?sequence=1	1262 (7) 6.77 %
2	https://card-file.ontu.edu.ua/bitstreams/aed610a6-43ef-47e0-9066-e85c89456f3e/download	253 (9) 1.36 %
3	https://card-file.ontu.edu.ua/bitstreams/8999d5af-6274-44f4-ae78-d23e08048d38/download	168 (6) 0.90 %
4	https://card-file.ontu.edu.ua/bitstreams/29489599-0581-4ce6-8890-c3b13d9f2e0e/download	126 (5) 0.68 %
5	https://advokatonline.org.ua/sanitarno-hihienichni-vymohy-do-robochoho-mistsya-v-ofisi/	115 (5) 0.62 %
6	https://card-file.ontu.edu.ua/bitstreams/bbed74c8-2ea7-44c5-8d00-0fe3fd9790ee/download	112 (4) 0.60 %
7	https://card-file.ontu.edu.ua/bitstreams/63ee88cb-a3d0-4005-9cf2-0cff89f28c0d/download	97 (3) 0.52 %
8	https://accountingcnt.com.ua/robota-v-ofisi-osnovni-sanitarno-hihienichni-vymohy/	90 (2) 0.48 %
9	https://card-file.ontu.edu.ua/bitstreams/361286d7-8a03-4221-ad05-db5133ab5f79/download	86 (3) 0.46 %
10	https://card-file.ontu.edu.ua/server/api/core/bitstreams/e86ba9fc-9135-43bb-922a-10bf0bce46b5/content	69 (1) 0.37 %
11	https://zp.edu.ua/sites/default/files/konf/konspekt_lekciy_opg_fbad-gum.pdf	53 (1) 0.28 %
12	https://card-file.ontu.edu.ua/bitstreams/7f031b5d-95bb-43c2-8b3f-0fa2499416c4/download	52 (5) 0.28 %
13	https://card-file.ontu.edu.ua/bitstreams/d42aac6d-ab01-4a74-b9cb-ced2a9eff719/download	50 (6) 0.27 %
14	https://skaz.com.ua/other/16151/index.html	37 (1) 0.20 %
15	https://kafedra-aop.at.ua/news/canitarija_ta_gigiena_robochogo_miscja_v_ofisi/2012-12-16-12	34 (2) 0.18 %
16	https://stackoverflow.com/questions/55642865/print-the-even-numbers-from-a-given-list	32 (3) 0.17 %
17	https://thepresentation.ru/obzh/elektrobezpeka-ta-pozhezhna-bezpeka-pri-robot%D1%96-z-kompyuterom	30 (2) 0.16 %
18	https://studopedia.com.ua/1_252891_z-a-v-d-a-n-n-ya.html	28 (1) 0.15 %
19	https://card-file.ontu.edu.ua/server/api/core/bitstreams/a05c07c5-bf65-4cb0-bdfa-e28694707551/content	26 (2) 0.14 %
20	https://ualaw.org/gigiyena-praczi-u-detalyah-u-derzhpraczi-nagadaly-vymogy-vyrobnychoyi-sanitariyi-do-robochogo-misczya/	26 (1) 0.14 %
21	https://card-file.ontu.edu.ua/bitstreams/53ed22ad-8700-4162-b97a-082a1ad472d6/download	24 (3) 0.13 %
22	https://link.springer.com/content/pdf/10.1007/s10639-022-11299-8.pdf	20 (1) 0.11 %
23	https://programmer.ie/post/fine_tuning/	19 (2) 0.10 %
24	https://ena.lpnu.ua/bitstreams/d713fe5a-53f7-4ea0-bca8-695b13983db9/download	16 (1) 0.09 %
25	https://www.lomitpatel.com/articles/the-impact-of-mobile-apps-in-education-learning-on-the-go/	14 (2) 0.08 %
26	https://openarchive.nure.ua/bitstreams/086e3976-c5a1-4dc1-b109-18e8cc6e63a6/download	13 (2) 0.07 %
27	https://higherseducation.com/revolutionizing-learning-educational-resources-for-students-and-educational-technology-tools/	12 (1) 0.06 %
28	https://ses.lviv.ua/novyny/2017/serpen/bezpechni-umovy-pratsi-zaporuka-komfortnoyi-pratsezdatsnosti-ofisnykh-pratsivnykiv	12 (1) 0.06 %

29	https://card-file.ontu.edu.ua/bitstreams/ad8936e7-88a5-4237-9847-551ee0d4608e/download	11 (2) 0.06 %
30	http://primat.org/publ/shkola_reshebniki/matematika/tipovi_zavdannja_zno_kombinatorika/48-1-0-1333	11 (1) 0.06 %
31	https://card-file.ontu.edu.ua/server/api/core/bitstreams/8dbb27f1-50a8-483b-baa5-01287b45fcc7/content	11 (1) 0.06 %
32	https://card-file.ontu.edu.ua/bitstreams/f789da43-3034-4ad8-bf34-640a47414f93/download	8 (1) 0.04 %
33	https://card-file.ontu.edu.ua/bitstreams/538ada8a-2c79-4b1e-b7d2-b0c97f68bc1c/download	5 (1) 0.03 %

Список прийнятих фрагментів (немає прийнятих фрагментів)

ПОРЯДКОВИЙ НОМЕР	ЗМІСТ	КІЛЬКІСТЬ ОДНАКОВИХ СЛІВ (ФРАГМЕНТІВ)
------------------	-------	---------------------------------------

1

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

Спеціальність:

123 – «Комп'ютерна інженерія»

Освітньо-професійна програма:

«Комп'ютерна інженерія» Група: 2БКС-29

КВАЛІФІКАЦІЙНА РОБОТА БАКАЛАВРА здобувача освіти денної форми навчання БКС 29.05.000.00 КРБ

ГАВРИЛЮКА
ОЛЕГА
СЕРГІЙОВИЧА

м. Одеса
2025 р.

2

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ВСП «ОДЕСЬКИЙ ТЕХНІЧНИЙ ФАХОВИЙ КОЛЕДЖ ОНТУ»

Спеціальність: _____

Освітньо-професійна програма: _____

Група: _____

ПОЯСНЮВАЛЬНА ЗАПИСКА До кваліфікаційної роботи бакалавра на тему: _____

Проектний матеріал складається з
пояснювальної записки на _____ сторінках та графічного матеріалу на _____ аркушах (слайдах). Виконавець _____
(_____)
Керівник проекту _____ (_____)

Консультанти:

з розділу охорони праці та техніки безпеки _____ (Чорновол Н.І.) з нормоконтролю _____ (_____)

Петрашова В.І.) старший консультант _____ (Кривченко Ю. В.) До захисту допущений _____

Завідувач кафедри _____ (Іванова Л.В.) Завідувач відділенням _____

_____ (Краснокутська К. Г.)

Захист « _____ » 2025 р. Протокол ЕК № _____ Оцінка ЕК _____ Секретар ЕК _____

123 – «Комп'ютерна інженерія»

2БКС-29

Гаврилюк О.С.

Іванова Л.В.

80

12