

ПОЛТАВСЬКИЙ ДЕРЖАВНИЙ АГРАРНИЙ УНІВЕРСИТЕТ
Навчально-науковий інститут економіки, управління, права та
інформаційних технологій
Кафедра інформаційних систем та технологій

КВАЛІФІКАЦІЙНА РОБОТА

на здобуття ступеня вищої освіти магістр

на тему: **«Нейроконсультант інтернет-магазину спортивного одягу»**

Виконав: здобувач вищої освіти
за освітньою програмою
Інформаційні управляючі системи та
технології
спеціальності 126 Інформаційні
системи та технології
ступеня вищої освіти магістр
групи 126ІСТ_мд_2023
Влох Тарас Сергійович
Керівник: Слюсар Вадим Іванович
Рецензент: Муравльов Володимир
Вячеславович

Полтава – 2024 року

ВСТУП

Актуальність теми кваліфікаційної роботи підтверджується необхідністю покращення користувацького досвіду в інтернет-магазинах через впровадження технологій штучного інтелекту. Інтеграція нейронних мереж сприяє оптимізації процесів пошуку та підбору товарів, персоналізації рекомендацій і автоматизації підтримки клієнтів. Це дозволяє значно підвищити ефективність роботи онлайн-магазинів, водночас покращуючи зручність взаємодії для користувачів. Застосування нейронних мереж у сфері електронної комерції залишається недостатньо розвиненим, що робить тему роботи актуальною для дослідження.

Зв'язок роботи з науковими програмами, темами. Робота відповідає дослідженням в межах науково-дослідної ініціативної тематики «Організаційно-методологічні аспекти впровадження інформаційно-комунікаційних систем і технологій в управлінні діяльністю сучасних організацій та підприємств за умов переходу до цифрової економіки» (ДРН 0123U105060, 2023-2028 рр.), що реалізується на кафедрі інформаційних систем та технологій, тематиці досліджень навчально-дослідної лабораторії інтелектуальних систем, комп'ютерних мереж та інтернет речей кафедри інформаційних систем та технологій Полтавського державного аграрного університету.

Метою кваліфікаційної роботи є підвищення ефективності взаємодії клієнтів з інтернет-магазином та збільшення конверсії за рахунок впровадження штучного інтелекту.

Завданнями кваліфікаційної роботи є:

- аналіз особливостей інтеграції штучного інтелекту та інтернет-комерції;
- дослідження API нейронних мереж;
- оцінка синтезованої моделі нейронної мережі для реалізації нейроконсультанта;
- розробка рекомендацій щодо інтеграції нейроконсультанта в інтернет-магазин для оптимізації обслуговування клієнтів.

Об'єктом дослідження є процес взаємодії клієнтів із інтернет-магазином.

Предметом дослідження є нейронні мережі, що використовуються для створення інтелектуальних консультантів.

Методами дослідження є аналітичний, інформаційно-пошуковий, методи синтезу та тестування GPT-моделей, робота з бібліотекою OpenAI API, розробка та інтеграція нейроконсультанта на сайт.

Інформаційна база кваліфікаційної роботи базується на ресурсах про трансформерні моделі, що використовуються в OpenAI API, принципи роботи великих мовних моделей, а також приклади інтеграції нейронних мереж у системи електронної комерції.

Елементи наукової новизни роботи полягають в розробці моделі нейроконсультанта на основі генеративного штучного інтелекту.

Практична значущість роботи полягає в розробці рекомендацій щодо застосування штучних нейронних мереж для розробки нейроконсультантів – можуть бути використані для подальших досліджень за даною тематикою та при проєктуванні хмарних сервісів.

Апробація результатів відбувалася в рамках V Міжнародної студентської наукової конференції «Міждисциплінарні наукові дослідження та перспективи їх розвитку» (жовтень 2024 р., м. Черкаси), VII Міжнародної студентської наукової конференції «Розвиток суспільства та науки в умовах цифрової трансформації» (листопад 2024 р., м. Тернопіль).

За результатами досліджень здійснено 2 публікації тез доповідей.

Структура кваліфікаційної роботи логічно пов'язана з завданнями досліджень і містить вступ, три розділи основної частини, висновки, список використаних джерел, додатки. Загальний обсяг пояснювальної записки кваліфікаційної роботи складає 71 сторінку формату А4. Вона містить 33 рисунки та 1 таблицю.

РОЗДІЛ 1

АНАЛІЗ ОСОБЛИВОСТІ ІНТЕГРАЦІЇ ШТУЧНОГО ІНТЕЛЕКТУ ТА ІНТЕРНЕТ-КОМЕРЦІЇ

1.1 Роль штучного інтелекту в сучасних технологіях

Штучний інтелект (Artificial Intelligence, AI) на сьогодні є одним із найбільш перспективних напрямків у розвитку сучасних технологій. Здатність AI аналізувати великі обсяги даних, навчатися на їх основі та адаптуватися до нових умов значно розширює можливості у багатьох галузях. Основна мета AI – створити розумні системи, здатні виконувати складні завдання, що вимагали людських інтелектуальних зусиль. Поява AI створює потенціал для розвитку розумних рішень, що суттєво впливають на бізнес, економіку, повсякденне життя та навіть мистецтво [1].

AI сприяє технологічному прогресу в багатьох аспектах, від автоматизації рутинних процесів до розробки інноваційних рішень для складних завдань, які вимагають творчого підходу. Завдяки методам, як-от машинне навчання, глибинне навчання та обробка природної мови, AI здатен не лише обробляти дані, але й прогнозувати події, оцінювати ризики та знаходити нові шляхи вирішення проблем. У результаті AI стає важливим елементом бізнесу та технологічних процесів [2].

Машинне навчання (Machine Learning, ML) дає змогу системам автоматично вчитися та вдосконалюватися на основі досвіду без потреби в ручному програмуванні. Воно включає широкий спектр підходів і методів, що дозволяють створювати адаптивні системи, здатні знаходити приховані закономірності у даних і робити точні прогнози. ML активно використовується в різних галузях, зокрема у фінансах, для прогнозування ринкових тенденцій, і в медицині, для діагностики захворювань. Наприклад, алгоритми класифікації допомагають у виявленні підозрілих операцій, що підвищує рівень захищеності фінансових транзакцій [3].

Глибоке навчання (Deep Learning, DL) – це підвид ML, що використовує багатошарові нейронні мережі. Завдяки здатності DL обробляти великий обсяг неструктурованих даних, його застосування стає необхідним у таких складних завданнях, як розпізнавання обличчя, аналіз медичних зображень, обробка мовлення. Наприклад, у сфері охорони здоров'я DL дозволяє лікарям діагностувати захворювання на ранніх стадіях, що значно підвищує ефективність лікування. Сьогодні DL застосовується також у сфері автономного водіння, де AI забезпечує транспортні засоби здатністю ідентифікувати об'єкти, приймати рішення в реальному часі та передбачати потенційні загрози [4].

Обробка природної мови (Natural Language Processing, NLP) дозволяє AI сприймати, аналізувати та розуміти людську мову, що робить його особливо цінним у сфері комунікацій та клієнтського обслуговування. NLP використовується для розробки чат-ботів, систем перекладу, аналізу текстових даних у соціальних мережах. Такі системи, як чат-боти та голосові помічники, здатні значно спростити процес взаємодії з клієнтами. NLP також допомагає автоматизувати процес обробки великих обсягів текстових даних, зокрема аналізуючи настрої користувачів у соціальних мережах для створення персоналізованих рекламних кампаній [5].

Сфера інтернет-комерції є однією з тих, де AI допомагає вдосконалити досвід покупців завдяки персоналізації. Аналізуючи попередні дії користувачів, AI здатен рекомендувати товари, що можуть зацікавити клієнтів, і таким чином підвищити їхню задоволеність від процесу покупок. Крім того, за допомогою AI відбувається оптимізація управління запасами, що забезпечує ефективне планування поставок та зменшення витрат на зберігання.

В охороні здоров'я AI використовується для обробки медичних зображень, виявлення захворювань, прогнозування результатів лікування. AI сприяє скороченню часу на постановку діагнозу, що є особливо важливим у випадках, коли своєчасне втручання може врятувати життя. Такі технології, як роботизовані хірургічні системи, вже зараз допомагають хірургам під час операцій, підвищуючи точність роботи та знижуючи ризик ускладнень [6].

Фінансовий сектор активно використовує AI для запобігання шахрайству, управління інвестиціями та аналізу ризиків. Завдяки автоматизації рутинних процесів, як-от обробка кредитних заявок або перевірка трансакцій, AI значно спрощує діяльність фінансових установ і підвищує їхню ефективність. Здатність AI аналізувати величезні масиви інформації дозволяє фінансовим аналітикам швидко виявляти ринкові тренди та приймати обґрунтовані рішення [7].

Зі стрімким зростанням обсягів даних, які генеруються людьми та пристроями, сучасний світ стикається з викликом ефективної обробки інформації для прийняття оптимальних рішень. Великі масиви даних перевищують людські можливості обробки, що підсилює потребу у використанні AI. AI допомагає витягти цінну інформацію з даних, що має критичне значення для розуміння трендів і підтримки обґрунтованих рішень, зокрема у процесах ML.

Активне впровадження AI обумовлене стрімким розвитком сучасних обчислювальних технологій та доступністю великих обсягів даних, які слугують основою для тренування нейронних мереж. Важливу роль відіграють також інновації у сфері хмарних обчислень, які спрощують інтеграцію інтелектуальних рішень у різні бізнес-процеси. Крім того, компанії все частіше звертають увагу на нові джерела даних, включаючи поведінкові та контекстуальні дані, які раніше залишалися поза увагою.

AI надає суттєві конкурентні переваги для сучасного бізнесу, а також підвищує його адаптивність. Використовуючи алгоритми ML і DL, AI дозволяє аналізувати великі обсяги даних у реальному часі, що сприяє швидкому реагуванню на зміни ринкових умов. Наприклад, за допомогою AI компанії можуть пропонувати клієнтам більш точні й персоналізовані рекомендації щодо продуктів, що позитивно впливає на лояльність клієнтів і рівень конверсії.

Багато компаній визнають значення AI у створенні інноваційних підходів до бізнесу, що сприяє вдосконаленню бізнес-процесів і стратегій адаптації. У результаті AI може бути перспективним інструментом для забезпечення

довгострокових конкурентних переваг, дозволяючи бізнесу розвиватися в умовах швидких технологічних змін.

AI також значно покращує обслуговування клієнтів за допомогою чат-ботів та віртуальних асистентів, які здатні швидко реагувати на запити в будь-який час доби, що підвищує задоволеність користувачів і зменшує навантаження на працівників. Крім того, AI допомагає оптимізувати ланцюги постачання, аналізуючи попит, прогножуючи потреби в продуктах і коригуючи запаси. Такий підхід допомагає компаніям швидко адаптуватися до ринкових змін і мінімізувати втрати. У сфері безпеки AI активно використовується для моніторингу аномалій та захисту даних, що стає критично важливим в умовах зростання кіберзагроз. Водночас AI сприяє розвитку нових бізнес-моделей, як-от підписка на основі споживання або індивідуалізовані послуги, що відкриває нові джерела доходів для компаній [8].

Завдяки автоматизації на базі AI, організації можуть звільнити співробітників від рутинних завдань, що дозволяє їм зосередитися на стратегічно важливих питаннях. Такий підхід сприяє оптимізації внутрішніх процесів, підвищенню ефективності роботи та зменшенню кількості людських помилок. AI також відкриває нові можливості для персоналізації клієнтського досвіду, забезпечуючи індивідуальний підхід до кожного користувача.

AI відкриває нові можливості у сфері прогнозної аналітики, що дозволяє виявляти тенденції та закономірності, які можуть бути непомітні для людини. Такі технології широко застосовуються, наприклад, у фінансових послугах для ідентифікації потенційних можливостей та ризиків на ринку, що дозволяє постачальникам послуг забезпечити клієнтів більш точною інформацією для прийняття зважених рішень. Таким чином, AI стає важливим елементом для ефективного управління у різних галузях.

AI дозволяє надавати клієнтам персоналізовані рекомендації, що особливо важливо для галузей електронної комерції, розваг і цифрового маркетингу. Компанії на кшталт Netflix і Amazon використовують AI для формування

рекомендацій на основі попередніх переглядів та покупок користувачів, що підвищує рівень їхнього задоволення та сприяє залученню нових клієнтів [9].

1.2 Перспективні архітектури нейронних мереж

У сучасному світі стрімкого розвитку технологій AI займає провідне місце, однією з найбільш прогресивних технологій у сфері NLP є моделі генеративного AI, які здатні не лише аналізувати великі обсяги текстових даних, але й створювати їх із високою точністю та зв'язністю.

Особливу увагу привертає генеративний попередньо тренований трансформер (Generative Pre-trained Transformer, GPT) – технологія, що являє собою універсальну систему для обробки й прогнозування мови. Це спеціалізований алгоритм, здатний аналізувати текст, витягувати з нього ключову інформацію, узагальнювати її та генерувати осмислений і структурований контент [10]. На рис. 1.1 наведено приклад трансформерної архітектури GPT-1.

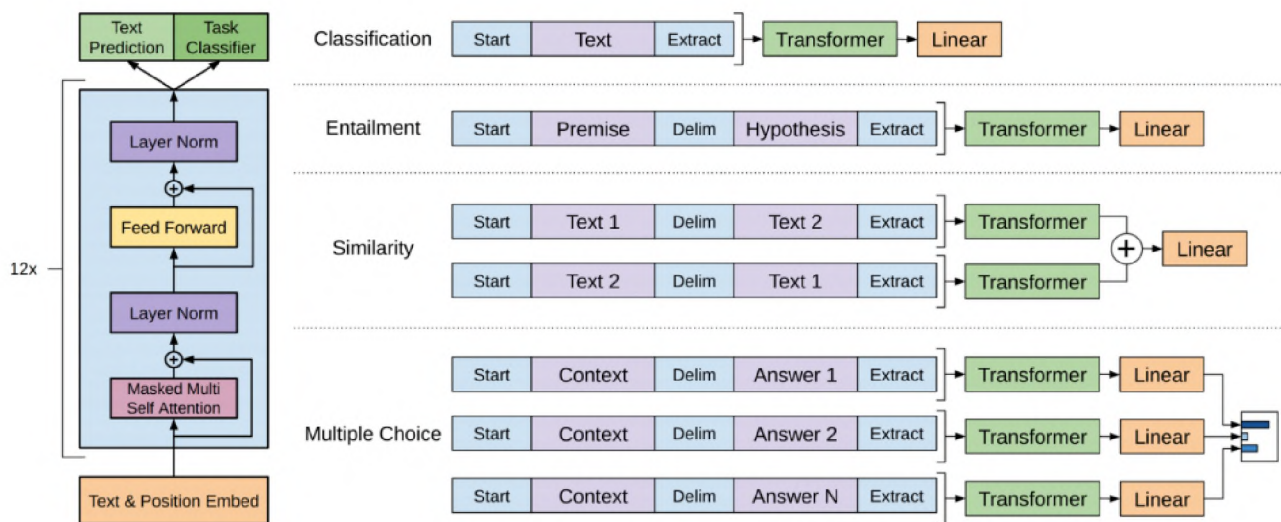


Рисунок 1.1 – Трансформерна архітектура GPT-1 [11]

В умовах зростання вимог до автоматизації процесів і підвищення ефективності, GPT стає незамінним інструментом, який не лише оптимізує

рутинні завдання, а й відкриває нові горизонти для інновацій. У цьому контексті важливо детально розглянути суть GPT, зрозуміти, що сприяє її популярності, та дослідити ключові можливості, які забезпечують її широке використання у світі.

GPT, хоч і відноситься до сфери AI, більш точно визначається як нейромережева модель обробки та генерації тексту, побудована на архітектурі Transformer. Ці моделі аналізують введення у вигляді NLP та створюють найбільш імовірну відповідь, спираючись на глибоке розуміння мовних структур.

Основою роботи GPT є знання, отримані під час навчання на величезних обсягах даних, що включають сотні мільярдів параметрів. Завдяки цьому GPT здатна враховувати контекст введеного тексту та динамічно адаптувати відповіді, генеруючи не лише окремі слова, а й складні осмислені тексти. Наприклад, якщо GPT отримує запит створити вірш у стилі відомого поета, вона аналізує контекст і відтворює унікальний текст, зберігаючи стиль та структуру оригіналу. На відміну від інших типів нейронних мереж, таких як рекурентні або згорткові, GPT базується на трансформерній архітектурі. Головна перевага цієї архітектури полягає у використанні механізму самоуваги, який дозволяє моделі приділяти більше уваги ключовим частинам тексту під час обробки даних. Це забезпечує глибше розуміння контексту та підвищує ефективність у задачах NLP. Трансформерна архітектура складається з двох модулів: кодера та декодера, які спеціалізуються на обробці вхідних текстів і побудові послідовних відповідей, що робить GPT одним із найпотужніших інструментів у сучасному AI [12].

Кодер моделі GPT перетворює вхідний текст у векторне представлення, яке є математичною інтерпретацією кожного слова. У цьому векторному просторі слова, які мають схожий сенс, розташовані ближче одне до одного. Під час обробки тексту кодер враховує контекст кожного слова у послідовності, що дозволяє створити більш точне розуміння значення тексту. Коли текст надходить на вхід, блок кодера розбиває його на частини й присвоює ваги кожному слову. Ці ваги вказують на важливість конкретних слів у контексті всього речення. Завдяки цьому модель може визначати, які слова мають найбільший вплив на зміст тексту. Додатково позиційні кодувальники допомагають уникнути

плутанини у значенні тексту, навіть якщо структура речення змінюється [12]. Наприклад, позиційне кодування дозволяє моделі розрізняти значення між такими реченнями: Пташка летить за метеликом. Метелик летить за пташкою.

Кодер формує векторне представлення тексту, яке включає всі особливості контексту, і передає цю інформацію до наступного компонента – декодера.

Декодер отримує сформоване кодером векторне представлення тексту та використовує його для створення прогнозованого результату. Він застосовує механізми самоконтролю (self-attention), щоб зосередитися на релевантних частинах вхідного тексту та передбачити найбільш відповідну відповідь. У ході роботи декодер оцінює кілька можливих варіантів і обирає найкращий з них [12].

Цей підхід забезпечує високу точність і адаптивність при створенні текстів, відповідаючи на запити різної складності. На відміну від рекурентних нейронних мереж, трансформатори, на яких побудована GPT, обробляють текст паралельно, а не послідовно. Це значно підвищує продуктивність і робить моделі GPT більш ефективними для розв’язання задач NLP. На рис. 1.2 наведено приклад навчання архітектури GPT-3.5.

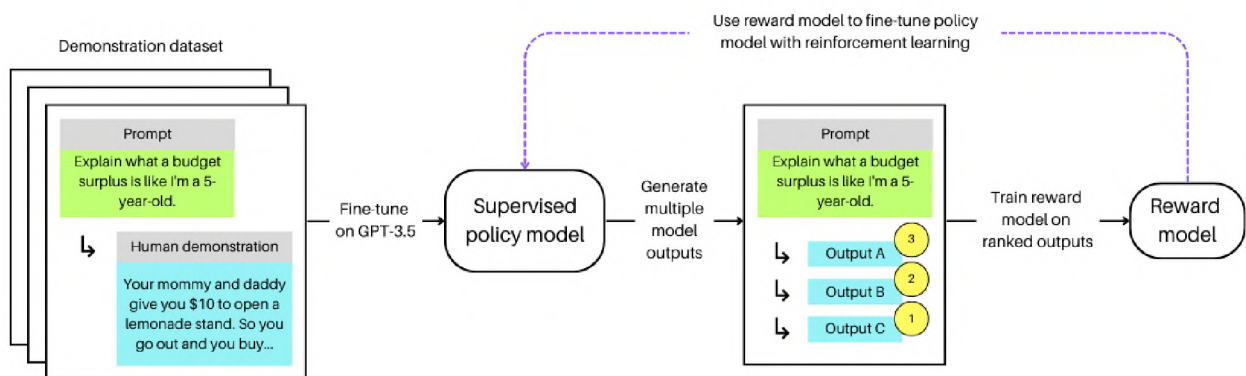


Рисунок 1.2 – Приклад навчання GPT-3.5 [13]

GPT-моделі є потужним інструментом для автоматизації процесів, надаючи гнучкі та ефективні рішення для багатьох завдань. Вони можуть використовуватись у різних сферах, таких як створення персоналізованого контенту, автоматизація обробки текстів, написання кодів або оптимізація роботи

чат-ботів. Висока точність і адаптивність моделей дозволяє впроваджувати їх у таких напрямках, як автоматизація юридичного аналізу, підтримка клієнтів, аналіз великих даних та кібербезпека. Однією з головних переваг GPT є можливість їхньої інтеграції в бізнес-процеси та автоматизація рутинних задач, що економить час і ресурси. GPT здатні створювати текст високої якості та відповідати на складні запити, що дозволяє використовувати їх у науці, освіті, розвагах і навіть творчих індустріях [14].

Попри свої численні переваги, GPT-моделі мають певні обмеження. Одним із ключових є те, що їхнє навчання базується на історичних наборах даних, через що вони не враховують подій чи даних, які з'явилися після завершення навчання. Наприклад, GPT не зможе надати інформацію про нещодавні події чи технічні новинки, які стали актуальними після дати завершення навчання. Ще одним аспектом є залежність GPT від текстових форматів. Моделі цього типу найкраще працюють із текстовими даними, що може обмежувати їхнє застосування для аналізу мультимедійного контенту. В окремих випадках мультимодальні моделі можуть бути більш доцільними для вирішення задач, які потребують аналізу зображень, відео чи звуку [14].

GPT, як мовна модель, навчається на величезних масивах текстових даних, доступних у відкритих джерелах, включаючи інтернет. Через це існує ризик, що модель може отримати доступ до шкідливого, токсичного або некоректного контенту. Для забезпечення безпеки і відповідності людським цінностям, GPT піддається ретельному процесу адаптації. Цей процес спрямований на вирівнювання моделі з етичними нормами та очікуваннями користувачів. Мета полягає в тому, щоб GPT забезпечував відповіді, які відповідають людським намірам і не суперечать моральним чи соціальним принципам. Значну роль у цьому відіграє метод навчання з підкріпленням і зворотним зв'язком від людини.

Під час цього етапу створюються спеціальні демонстраційні набори даних, які показують бажану поведінку моделі у відповідь на різноманітні запити. Відповіді розміщуються у певному порядку, що дозволяє GPT розпізнавати, які

відповіді є найбільш доречними в певному контексті. Після цього модель доопрацьовується, щоб мінімізувати ризик небажаних відповідей [15].

Однак, слід зазначити, що жодна система AI не є абсолютно досконалою. GPT працює відповідно до заданих обмежень і правил, але може бути вразливим до спроб обійти ці обмеження. Незважаючи на це, модель демонструє високий рівень безпеки та корисності, якщо її використовують відповідально.

1.3 Особливості та фактори зростання популярності інтернет-комерції

Сучасний розвиток нейронних мереж привертає все більше уваги з боку бізнес-лідерів, які починають усвідомлювати їхній величезний потенціал і численні можливості. У міру того, як обсяги онлайн-покупок зростають, точний збір і ретельна організація даних стають надзвичайно важливими. Нейронні мережі вже зарекомендували себе як ефективний інструмент для покращення взаємодії компаній з клієнтами, допомагаючи краще розуміти їхні потреби і вчасно реагувати на зміни в їхній поведінці [16].

Пошукові запити в електронній комерції часто не дають бажаних результатів через обмеження традиційних алгоритмів, які зосереджуються лише на співставленні ключових слів. Нейронні мережі з NLP дозволяють додати «людський» елемент, інтерпретуючи запити на природній мові, що допомагає користувачам швидше знаходити потрібний товар. Наприклад, система може розпізнати й відобразити потрібні результати на основі специфічних характеристик, як-от стиль, матеріал чи колір одягу. Інновації в цій сфері дозволяють здійснювати пошук із врахуванням текстових і візуальних елементів, що значно підвищує точність результатів.

Нейронні мережі збирають і аналізують величезні обсяги даних про поведінку споживачів, їхні вподобання та історію покупок, дозволяючи компаніям формувати персоналізовані рекомендації та таргетовану рекламу. Завдяки сегментації клієнтів на основі їхніх індивідуальних потреб, бізнес може

ефективніше просувати продукти, що відповідають інтересам кожного покупця. Наприклад, такі платформи, як Amazon, використовують нейромережі для рекомендацій, що ґрунтуються на поведінкових паттернах подібних клієнтів, що підвищує ймовірність купівлі та рівень задоволеності клієнтів.

Згідно з прогнозами, у період з 2024 по 2027 рік обсяги продажів продовжать зростати зі середнім річним темпом 7,8%, досягнувши \$ 8 трлн (рис. 1.3). Це перевищуватиме обсяги виторгу традиційних магазинів більш ніж удвічі, що свідчить про те, що електронна комерція стає все більш привабливим варіантом для ритейлерів у глобальному масштабі [17].

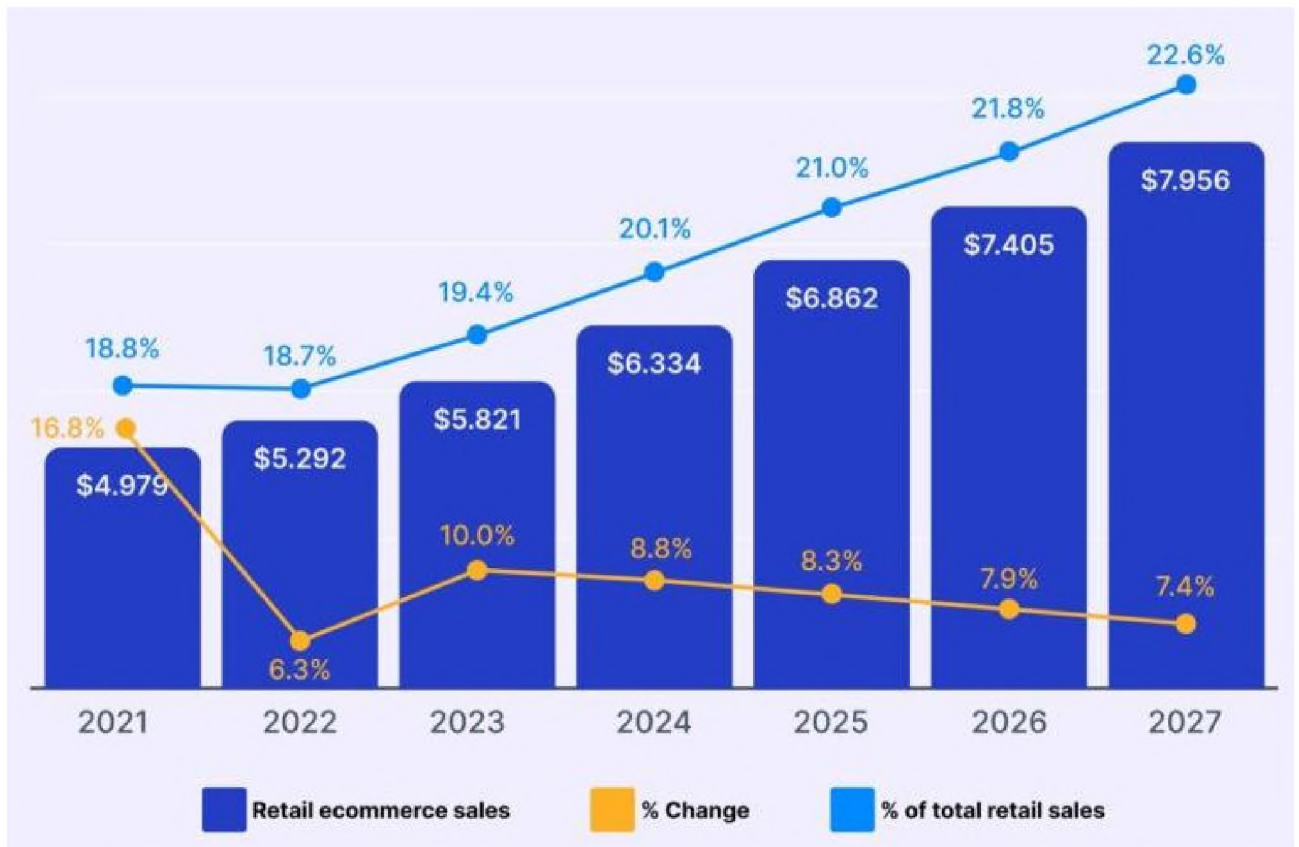


Рисунок 1.3 – Графік обсягів продажів інтернет-комерції за 2021-2027 роки [17]

Традиційні методи прогнозування продажів часто базуються на історичних даних, але не враховують швидко змінювані переваги споживачів. Нейронні мережі можуть аналізувати величезні обсяги інформації з різних джерел, відстежуючи зміни в попиті та споживчих настроях. Це дозволяє створювати більш точні прогнози та гнучко реагувати на ринкові тренди. Використання

нейронних мереж для аналізу цих даних може допомогти компаніям заздалегідь підготуватися до сезонних коливань і нових трендів, що забезпечує їм стратегічну перевагу на конкурентному ринку.

Крім здатності аналізувати попит і прогнозувати сезонні коливання, нейронні мережі також забезпечують компаніям можливість впроваджувати персоналізовані стратегії для задоволення різних споживчих груп. Завдяки глибокому аналізу вподобань, поведінкових патернів і демографічних характеристик, нейронні мережі можуть сприяти не тільки продажам, але й розробці продуктів, які точно відповідають очікуванням клієнтів. Бренди отримують шанс швидко адаптувати маркетингові та рекламні кампанії до нових тенденцій, використовуючи прогнози, що враховують як історичні, так і поточні дані в реальному часі [18].

Нейропомічники, такі як чат-боти та віртуальні асистенти, кардинально змінюють взаємодію клієнтів із брендами, додаючи автоматизацію, персоналізацію та ефективність на кожному етапі продажів та підтримки. Їх впровадження стало можливим завдяки нейронним мережам, що здатні аналізувати велику кількість даних, розпізнавати шаблони поведінки клієнтів і вчитися на основі попередньої взаємодії. Це робить нейропомічників надзвичайно корисними як для споживачів, так і для компаній.

Один з основних аспектів, як нейропомічники покращують клієнтську підтримку, полягає в забезпеченні миттєвого зворотного зв'язку. Замість того щоб чекати на оператора, клієнти можуть швидко отримати відповіді на свої запитання або розв'язати типові проблеми, звертаючись до чат-бота чи віртуального асистента. Нейронні мережі дозволяють цим помічникам обробляти запити NLP, розуміючи навіть складні питання, що ставлять клієнти. Це створює більш природний діалог, наближений до розмови з реальною людиною, і значно підвищує зручність та швидкість обслуговування [19].

Завдяки можливостям аналізу великих даних, нейропомічники здатні пропонувати персоналізовані рішення на основі історії покупок, уподобань та поведінки клієнтів на сайті. Наприклад, віртуальний асистент може

запропонувати клієнту продукти, які потенційно можуть зацікавити його, спираючись на попередні покупки або історію пошукових запитів. Це додає індивідуальний підхід у взаємодію, що сприяє задоволеності клієнтів і підвищує їхню лояльність до бренду. У результаті персоналізація сприяє зростанню продажів, оскільки клієнти відчують, що їхні потреби враховані.

Нейропомічники відіграють важливу роль на кожному етапі клієнтського шляху: від першого контакту з брендом до етапу післяпродажного обслуговування. Вони допомагають клієнтам обирати товари, розповідаючи про особливості продуктів, надають підтримку під час здійснення покупки, а також залишаються доступними після продажу для вирішення питань щодо гарантійного обслуговування або повернення. Наприклад, якщо клієнт звертається з питанням щодо статусу свого замовлення, чат-бот може надати цю інформацію автоматично, що економить час клієнта та знижує навантаження на операторів підтримки.

Автоматизація обслуговування клієнтів через нейропомічників дозволяє компаніям значно скоротити витрати. Віртуальні асистенти можуть обробляти тисячі запитів одночасно, що робить їх ефективнішими за команду операторів. Це особливо корисно для компаній з великим обсягом клієнтських звернень, оскільки нейропомічники можуть займатися типовими питаннями, залишаючи складніші завдання для кваліфікованих фахівців. Це не лише скорочує витрати, але й підвищує ефективність роботи всієї служби підтримки.

На відміну від традиційної служби підтримки, нейропомічники можуть навчатися та вдосконалюватися з кожною новою взаємодією. Нейронні мережі дозволяють їм аналізувати, які запити зустрічаються частіше, як краще відповідати на певні питання і що можна вдосконалити у відповідях. Це означає, що з часом нейропомічники стають все більш ефективними та адаптивними, підвищуючи загальну якість обслуговування клієнтів.

Чат-боти та віртуальні асистенти на базі нейронних мереж можуть стати важливими складовими іміджу бренду. Швидка, коректна, і водночас персоналізована підтримка створює враження високого рівня обслуговування.

Клієнти схильні залишатися з брендом, який пропонує якісне обслуговування і зручну комунікацію. У результаті нейропомічники допомагають зміцнити бренд і утримувати клієнтів, що робить їх потужним інструментом для розвитку довгострокової лояльності.

Вебсайти, особливо у сфері електронної комерції, активно використовують алгоритми на основі AI для покращення користувацького досвіду. Сучасні пошукові системи, що працюють на нейронних мережах, здатні обробляти запити NLP, що дозволяє користувачам отримувати більш релевантні результати пошуку. Розширене індексування та сегментація контенту на вебсайтах дозволяють краще структурувати інформацію, полегшуючи її обробку та забезпечуючи точнішу відповідь на запит користувача. Вебсайти використовують пошукові алгоритми зі AI, що допомагають користувачам швидко знаходити потрібний контент, обробляючи запити NLP для інтуїтивного пошуку. Сегментація тексту в таких системах дозволяє розділяти великі обсяги інформації на логічні блоки, полегшуючи аналіз і прискорюючи пошук релевантних даних. Ключовим прикладом використання нейромереж на вебсайтах є поява нейропомічників або чат-ботів на основі AI. Ці віртуальні помічники призначені для обробки запитів користувачів у режимі реального часу, пропонуючи допомогу, відповідаючи на запитання та проводячи користувачів через різні процеси на сайті. Нейронні мережі дозволяють цим чат-ботам вчитися на попередніх взаємодіях, вдосконалюючись з часом, коли вони отримують більше даних. Наприклад, сучасні вебсайти для обслуговування клієнтів, онлайн-банкінгу та електронної комерції значною мірою покладаються на нейропомічників для обробки типових запитів клієнтів, надання технічної підтримки, персоналізації послуг і навіть проведення безпечних транзакцій.

Нейронні мережі кардинально змінюють роботу вебсайтів, покращуючи персоналізацію, автоматизуючи взаємодію з користувачами за допомогою нейропомічників та посилюючи безпеку. З розвитком технологій роль нейронних мереж у веброзробці буде тільки зростати, пропонуючи ще більш досконалі інструменти для покращення користувацького досвіду та ефективності. Завдяки

своїй здатності до навчання та адаптації нейронні мережі стають незамінними у створенні більш розумних та адаптивних вебсайтів [20].

Висновки до розділу 1

У цьому розділі було досліджено вплив AI та нейронних мереж на електронну комерцію, акцентуючи увагу на їхніх ключових функціях та потенціалі. AI відіграє вирішальну роль у трансформації бізнесу, забезпечуючи інноваційні підходи до обробки даних, автоматизації процесів та підвищення точності в прогнозуванні споживчої поведінки. Нейронні мережі, в свою чергу, стають невід'ємною складовою сучасної електронної комерції, надаючи компаніям можливість краще розуміти потреби клієнтів, забезпечувати персоналізацію послуг, аналізувати дані та створювати нейропомічників, які оптимізують процеси продажів і підтримки.

Інтеграція таких технологій дозволяє бізнесу досягати нових рівнів ефективності, знижувати витрати на обслуговування клієнтів, утримувати споживачів та збільшувати прибуток. Особлива увага приділяється потенціалу нейропомічників у продажах і підтримці клієнтів, де завдяки ML та NLP, чат-боти й віртуальні асистенти значно покращують взаємодію з брендами та підвищують задоволеність користувачів.

Таким чином, нейронні мережі та AI стають потужними інструментами для розвитку електронної комерції, забезпечуючи вигідні конкурентні переваги й сприяючи глибшій інтеграції сучасних технологій у бізнес-процеси.

РОЗДІЛ 2

СИНТЕЗ АРХІТЕКТУРИ ТА ВИБІР ТЕХНОЛОГІЇ ДЛЯ НЕЙРОКОНСУЛЬТАНТА

2.1 Сучасні API нейронних мереж для інтеграції у вебдодатки

Інтерфейс програмування додатків (Application Programming Interface, API) ML пропонують стандартизовані інтерфейси до потужних попередньо навчених моделей та алгоритмів. Така стандартизація значно спрощує процес інтеграції та знижує зусилля незалежно від типу моделі чи сфери застосування. Додатки можуть встановлювати єдину комунікацію з різними частинами програмних систем, що сприяє стабільній роботі.

Зазвичай ці API працюють через інтернет: розробники надсилають запити на сервер із даними для обробки, і API, використовуючи власні моделі ML, обробляє інформацію та повертає відповідь з результатами. Це значно спрощує роботу з ML, дозволяючи розробникам використовувати AI без глибокого занурення у процеси навчання моделі або підготовку даних.

Наприклад, у задачах розпізнавання зображень розробник може відправити зображення API для аналізу. API використовує попередньо навчені моделі, щоб ідентифікувати об'єкти чи риси на зображенні, і повертає результати у структурованому форматі, наприклад, у вигляді запису об'єктів JavaScript (JavaScript Object Notation, JSON) або розширюваній мови розмітки (Extensible Markup Language, XML). Так само, для аналізу настроїв, можна надіслати текстовий фрагмент на API, що оцінить емоційне забарвлення тексту – позитивне, негативне чи нейтральне. Ці API беруть на себе основну обробку даних і процес виводу моделі, дозволяючи розробникам зосередитися на інтеграції результатів у застосунок [21].

Крім попередньо навчених моделей, багато API ML пропонують опції для налаштування. Це дозволяє розробникам адаптувати моделі під конкретні

потреби чи специфічні набори даних, надаючи API додаткову інформацію, яка допомагає оптимізувати параметри моделі для кращих результатів.

Використання API ML надає легкий доступ до функціоналу ML без необхідності створювати, навчати або розгортати моделі самостійно. Це дозволяє заощадити ресурси на збирання даних, навчання моделей та їх оптимізацію. Стандартизовані інтерфейси забезпечують легку інтеграцію, а повторне використання моделей значно скорочує час і зусилля, необхідні для створення AI-рішень.

Ці API зазвичай працюють на хмарних платформах, що дозволяє легко масштабувати додатки під навантаження без втрати продуктивності. Крім того, API ML забезпечують доступ до передових моделей та методик, які постійно оновлюються і підтримуються їхніми розробниками. Це означає, що додатки можуть використовувати новітні досягнення в AI без необхідності в глибоких знаннях або ручних оновленнях.

Як зазначалося раніше, багато API ML підтримують налаштування моделей, що дозволяє адаптувати їх до специфічних потреб чи наборів даних.

API ML забезпечують універсальні кінцеві точки, які можуть взаємодіяти з моделями ML. Це значно спрощує інтеграцію функцій AI у вже існуючі системи, дозволяючи легко організувати взаємодію між незалежними компонентами і моделями. Така модульність відкриває можливість для різноманітних застосувань та користувачів ефективно користуватися AI-моделями.

Класифікація API нейронних мереж за типами застосувань у вебдодатках включає кілька основних категорій, які допомагають розробникам обрати відповідні інструменти для інтеграції AI у свої продукти. Основні категорії можна поділити на такі напрямки, як обробка тексту, аналіз візуальних даних, персоналізація, обробка аудіо, аналітика та прогнозування.

1. Обробка тексту є ключовою функціональністю для багатьох вебдодатків, особливо тих, які потребують обробки великих обсягів текстової інформації. API для обробки тексту використовуються у таких завданнях, як аналіз настроїв, модерація контенту, розпізнавання мовлення та машинний

переклад. Наприклад, аналіз настроїв дозволяє виявити емоційне забарвлення тексту, що може бути корисним для підтримки клієнтів або аналізу відгуків у соціальних мережах. Інші задачі, як модерація контенту, допомагають забезпечити безпеку та відповідність вимогам платформи, автоматично перевіряючи текст на наявність небажаних елементів.

2. Візуальні дані – це ще одна важлива категорія API, яка охоплює обробку зображень і відео. Вона є корисною для таких функцій, як розпізнавання об'єктів, класифікація зображень, оптичне розпізнавання символів (Optical Character Recognition, OCR) та аналіз відео. Наприклад, розпізнавання об'єктів застосовується у таких сферах, як електронна комерція, де можна автоматично ідентифікувати продукти на зображеннях, або в безпеці для розпізнавання обличчя. OCR дозволяє автоматично зчитувати текст із зображень, що дуже корисно для оцифрування документів.

3. Персоналізація є незамінною у сучасних вебдодатках, оскільки дозволяє створювати індивідуальний досвід для кожного користувача. API для персоналізації допомагають розробникам створювати рекомендаційні системи, що підбирають контент чи товари залежно від уподобань та поведінки користувачів. Також вони надають інструменти для адаптації інтерфейсу, що дозволяє зробити додаток максимально зручним для різних категорій користувачів, враховуючи їхні потреби.

Також існують спеціалізовані API для аналітики, прогнозування та обробки аудіо. Аналітика та прогнозування дозволяють вебдодаткам обробляти великі обсяги даних для створення корисних інсайтів та прогнозів. Наприклад, аналіз поведінки користувачів може допомогти оптимізувати продукт та маркетингові стратегії. Водночас обробка аудіо охоплює такі завдання, як розпізнавання мовлення, аналіз емоцій у голосі та синтез мови. Це робить API для обробки аудіо корисними для додатків з голосовими інтерфейсами, таких як віртуальні помічники чи системи аудіокниг.

Інтеграція AI-технологій для NLP через API стала ключовим елементом для вебдодатків, які прагнуть забезпечити більш інтуїтивну і зручну взаємодію з

користувачами. NLP дозволяє обробляти текстові та голосові дані, інтерпретувати запити природною мовою, надавати автоматизовані відповіді та навіть адаптувати комунікацію залежно від контексту (рис. 2.1). Це відкриває широкі можливості для автоматизації та персоналізації обслуговування клієнтів, оптимізації процесів і підвищення рівня задоволеності користувачів. NLP також сприяє створенню інтерактивних систем, які здатні розпізнавати емоції клієнтів і відповідно коригувати стиль спілкування.

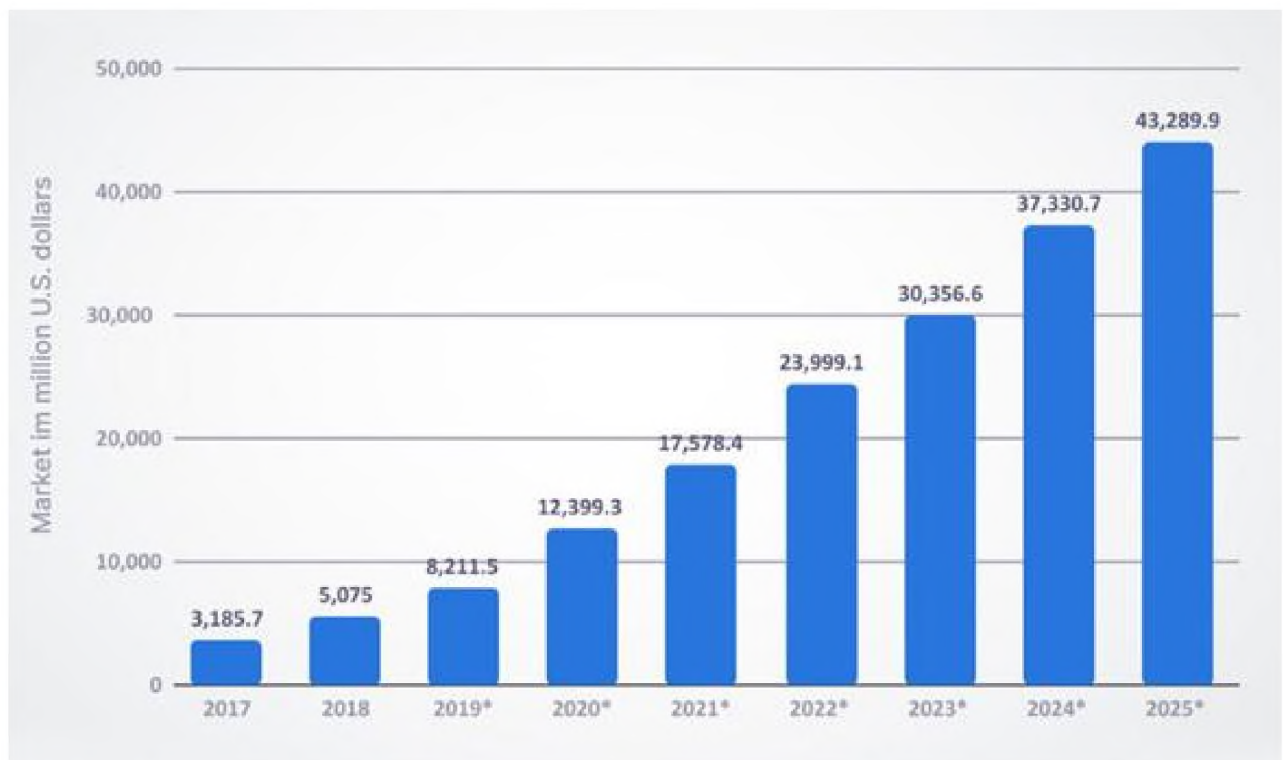


Рисунок 2.1 – Доходи від ринку NLP у всьому світі з 2017 по 2025 рік (у мільйонах доларів США) [22]

Одним із основних завдань NLP API є розпізнавання та розуміння тексту, введеного користувачем. Це включає такі аспекти, як виявлення ключових слів, розпізнавання сутностей (імена, дати, географічні назви) і визначення основного змісту повідомлення. Завдяки цій функції, вебдодатки можуть розуміти зміст запитів і автоматично класифікувати їх для подальшої обробки. Наприклад, у сфері підтримки клієнтів це дозволяє автоматично спрямовувати запити до відповідних відділів або рекомендувати релевантну інформацію. Багато NLP API

пропонують функцію аналізу настроїв, яка допомагає визначити емоційне забарвлення тексту – позитивне, негативне чи нейтральне. Це особливо корисно для моніторингу відгуків у соціальних мережах, аналізу відгуків клієнтів та загального контролю за репутацією бренду. Наприклад, якщо користувач залишає негативний коментар або повідомлення, система може автоматично виділяти його для пріоритетного обслуговування, щоб швидко реагувати на незадоволення і покращити враження клієнта.

За допомогою NLP API можна створити систему автоматизованих відповідей, яка реагуватиме на запити користувачів у режимі реального часу. Це дає можливість створювати чат-ботів та віртуальних помічників, які можуть відповідати на часті запитання, допомагати з навігацією по вебсайту чи підтримувати користувачів під час покупки. Така автоматизація значно знижує навантаження на відділи підтримки, даючи можливість фокусуватися на складніших завданнях. NLP API також пропонують можливості машинного перекладу, що дозволяє автоматично переводити тексти з однієї мови на іншу. Це корисно для багатомовних вебдодатків, які обслуговують клієнтів із різних країн. Завдяки високоточним перекладам, API дають змогу підтримувати кілька мов одночасно без необхідності залучення додаткових ресурсів для перекладу.

Одна з важливих функцій NLP API – це здатність інтерпретувати складні запити, сформульовані природною мовою. Такі API можуть розуміти питання, які потребують деталізованої відповіді, а також враховувати контекст. Наприклад, коли користувач запитує: «Який товар зараз у вас найпопулярніший?», система може дати точну відповідь, спираючись на актуальні дані про продажі, забезпечуючи інтерактивну та персоналізовану взаємодію. Багато API для NLP також підтримують функцію розпізнавання мовлення, що дозволяє користувачам взаємодіяти з додатками за допомогою голосових команд. Це особливо важливо для мобільних додатків та інтерфейсів, де голосове керування може бути більш зручним, ніж текстове введення. Така можливість забезпечує ще вищий рівень зручності і дозволяє охопити більшу аудиторію, включаючи користувачів з обмеженими можливостями.

Використання API для нейронних мереж надає численні переваги, основною з яких є швидкість розгортання. API забезпечують готові інструменти для інтеграції з додатками, дозволяючи бізнесам впроваджувати функціональність AI без необхідності у створенні моделей з нуля. Це економить час і ресурси, особливо коли йдеться про обмежений бюджет або короткі терміни. До того ж, API від провідних розробників регулярно оновлюються, включаючи нові функції і вдосконалення моделей, що дозволяє користувачам завжди мати доступ до найновіших технологій без додаткових витрат.

Однак, використання API також має свої недоліки. Головний мінус полягає в обмеженій можливості налаштування, оскільки API, як правило, надають стандартизовані моделі, які не завжди можуть врахувати специфічні потреби бізнесу. У випадках, коли потрібна унікальна модель для обробки певних видів даних або для досягнення високої точності, кастомні нейромережні рішення можуть бути більш ефективними. Створення власної нейронної мережі дозволяє мати повний контроль над архітектурою і гіперпараметрами моделі, адаптуючи її до конкретних завдань, що може значно покращити якість результатів.

Ще одним мінусом використання API є залежність від постачальників і можливі обмеження на обробку даних. Використовуючи API, компанії часто зобов'язані слідувати політикам обробки даних і конфіденційності постачальника, що може створити додаткові виклики у випадках, коли йдеться про конфіденційні або персоналізовані дані. У той же час, при створенні кастомного рішення всі дані можуть зберігатися і оброблятися в межах внутрішньої інфраструктури компанії, забезпечуючи кращий контроль над їхньою безпекою і конфіденційністю.

Загалом, вибір між API та кастомними рішеннями залежить від потреб бізнесу та ресурсів, які він може виділити на впровадження нейронних мереж. API – це ефективний інструмент для швидкого впровадження базової функціональності, але кастомні рішення можуть виявитися корисними для компаній, які потребують точності або унікальних можливостей обробки. Кастомні нейромережі дозволяють враховувати специфіку галузі, інтегрувати

додаткові алгоритми чи джерела даних, а також забезпечувати вищу безпеку обробки конфіденційної інформації.

Інтеграція API нейронних мереж у мобільні вебдодатки дозволяє значно розширити їх функціональні можливості, додаючи такі інструменти, як розпізнавання зображень, NLP, аналіз настроїв тощо. Включення таких технологій у мобільні додатки відкриває нові можливості для користувацької взаємодії та автоматизації, надаючи можливість створювати інтуїтивно зрозумілі та персоналізовані інтерфейси. Наприклад, мобільний додаток, інтегрований з сучасним NLP API, може автоматично обробляти текстові повідомлення від користувачів, розпізнавати їхній настрій або наміри і пропонувати релевантні відповіді, персоналізовані пропозиції чи рекомендації.

Перевагою інтеграції нейромережних API у мобільні додатки є їхня масштабованість та гнучкість. API, як правило, працюють на хмарних сервісах, що дозволяє обробляти великі обсяги даних без значного навантаження на мобільні пристрої. Це також забезпечує швидке розгортання нових функцій та оновлень, оскільки все, що потрібно, – це виклик до хмарного API. Крім того, оскільки API надаються провідними технологічними компаніями, користувачі можуть бути впевнені у високій точності, надійності та стабільності моделей, які постійно вдосконалюються і оновлюються.

Інтеграція нейромережних API у мобільні додатки також дозволяє створювати омніканальні стратегії взаємодії з користувачами. Використовуючи одні й ті ж API на різних платформах, компанії можуть забезпечити послідовний досвід користування незалежно від того, де саме клієнт взаємодіє з додатком – на мобільному телефоні, у браузері чи через смарт-пристрої. Це сприяє побудові лояльності та підвищенню задоволеності користувачів, оскільки вони отримують однаковий рівень обслуговування на всіх каналах.

Крім того, такі технології дозволяють компаніям швидше адаптуватися до змін у поведінці користувачів та реагувати на їхні потреби в режимі реального часу. Нейромережні API спрощують впровадження персоналізованих функцій,

таких як індивідуальні рекомендації чи автоматичний аналіз запитів, а також забезпечують більш точну обробку великих обсягів даних.

2.2 Принципи вибору API для створення нейроконсультанта

Визначення вимог до функціональності нейроконсультанта є критичним етапом у процесі його створення, оскільки ці вимоги формують основу для вибору відповідного API та визначають, які можливості має мати система для ефективної роботи. Нейроконсультант – це інтелектуальний інструмент, призначений для взаємодії з клієнтами, надання консультацій та рекомендацій, що мають бути максимально персоналізованими та відповідати конкретним потребам кожного користувача.

Однією з ключових функцій нейроконсультанта є NLP. Ця функція дозволяє йому розуміти текстові або голосові запити користувачів, інтерпретувати їхні запитання, іноді навіть з огляду на емоційний стан, тон та контекст. NLP включає такі підфункції, як розпізнавання намірів, виявлення ключових слів, розбір синтаксису, а також аналіз семантики та контексту запиту. Це дозволяє нейроконсультанту спілкуватися з клієнтами так, ніби він «розуміє» їхні потреби, що значно покращує якість взаємодії з користувачами. Наступна важлива функція – розуміння запитів користувачів. Нейроконсультант повинен мати здатність не просто аналізувати вхідні дані, а й розпізнавати приховані потреби та надавати коректні й обґрунтовані відповіді. Це означає, що система має розпізнавати навіть непрямі або складні запитання, коли користувачі не висловлюють свої потреби безпосередньо. Наприклад, якщо клієнт задає питання на кшталт «Що ви можете порадити для зимових подорожей?», нейроконсультант повинен зрозуміти, що користувач шукає товари для холодного сезону, і запропонувати відповідні продукти. Ще один важливий аспект – персоналізація рекомендацій. Нейроконсультант має надавати поради та пропозиції, які враховують історію попередніх запитів, уподобання та

поведінкові патерни кожного окремого користувача. Це вимагає від API доступу до функцій ML, що дозволяють аналізувати дані та формувати персоналізовані пропозиції на основі профілю клієнта. Така персоналізація сприяє більшому залученню користувачів і підвищує ймовірність, що клієнт скористається пропозицією або придбає продукт.

Крім того, нейроконсультант повинен підтримувати контекстну обізнаність. Це означає, що він має пам'ятати попередні запити та відповіді в межах одного сеансу, щоб продовжувати розмову на основі попередньої інформації. Така контекстна обізнаність дозволяє створити більш природну й безперервну взаємодію з користувачем, уникаючи повторення запитів або непослідовних відповідей.

1. Нейроконсультант повинен мати можливість підлаштовувати стиль і тон спілкування під різні аудиторії, наприклад, бути більш формальним або неформальним залежно від контексту.

2. Нейроконсультант має вміти навчатися на основі нових даних та оновлювати свої знання, щоб надавати більш точні рекомендації та краще розуміти потреби клієнтів.

3. Це дозволяє нейроконсультанту використовувати додаткові дані про користувачів, що покращує якість відповідей і дозволяє враховувати більше параметрів для персоналізації.

Отже, визначення функціональних вимог для нейроконсультанта дозволяє точно окреслити, якими можливостями має володіти система для надання ефективної допомоги користувачам.

У сучасному ринку нейронних мереж AI існує багато API, але не всі вони забезпечують однаковий рівень точності чи якості відповідей. Тому під час вибору API для створення нейроконсультанта потрібно провести ретельний аналіз характеристик кожної доступної моделі, оцінити її здатність точно обробляти інформацію і давати адекватні відповіді.

Оцінка якості моделей API базується на кількох ключових критеріях.

По-перше, важливим показником є точність моделі, яка вимірюється як здатність правильно інтерпретувати запити користувачів і відповідати на них у відповідності до очікувань. Цей аспект включає тестування моделі на широкому наборі запитів, щоб визначити, наскільки вона здатна розпізнавати контекст, ключові слова та інші аспекти NLP. Наприклад, у випадку запитань із неточно сформульованими потребами або тих, що містять сленг чи скорочення, модель повинна зберігати адекватність відповіді.

Другий важливий аспект – це рівень узгодженості відповіді з початковим запитом. Це означає, що відповіді мають бути не лише правильними, а й відповідати інтонації та стилю, які очікуються від нейроконсультанта. Оцінка узгодженості також включає перевірку на послідовність відповідей у рамках одного діалогу, тобто, чи здатен нейроконсультант «пам'ятати» попередні запити й відповідати у відповідному контексті. Непослідовність у відповідях може призвести до непорозумінь із користувачами, знизити довіру до нейроконсультанта та вплинути на загальну якість обслуговування.

Третім критерієм є швидкість обробки запитів і надання відповідей. Для ефективної роботи нейроконсультанта важливо, щоб обробка запитів була швидкою та не створювала затримок, особливо у випадках, коли користувачі очікують миттєву реакцію. Тому під час вибору API слід звернути увагу на те, наскільки оперативно модель здатна обробляти великі обсяги даних, чи підтримує вона швидкодію при одночасній обробці запитів від багатьох користувачів, а також її масштабованість у різних умовах.

Четвертий аспект – це гнучкість і можливості кастомізації моделей. Хороший API має пропонувати можливість донавчання чи налаштування моделей, щоб нейроконсультант міг краще адаптуватися до специфічних потреб бізнесу або аудиторії. Якщо модель підтримує кастомізацію, це дозволяє розширити її функціонал і точність відповідно до конкретних сценаріїв використання. Наприклад, для нейроконсультанта в спортивному інтернет-магазині можливість донавчання допоможе більш точно рекомендувати товари на основі попередніх уподобань клієнтів або нових тенденцій.

Останнім, але не менш важливим, є здатність моделі до підтримки різноманітних мов та культурних контекстів. Якщо нейроконсультант повинен працювати для багатомовної аудиторії, то вибраний API має забезпечувати високу точність обробки різних мов, діалектів і культурних особливостей. Це дозволяє забезпечити зручну та якісну взаємодію для користувачів із різних регіонів і уникнути непорозумінь, пов'язаних із лінгвістичними або культурними бар'єрами. Здатність налаштовувати API відповідно до специфічних вимог дозволяє створити індивідуальні рішення, які максимально відповідають потребам користувачів і забезпечують високий рівень ефективності. Це особливо актуально для нейроконсультантів, яким часто потрібно розуміти контекст запитів клієнтів, підтримувати персоналізовану взаємодію та інтегрувати рекомендації на основі специфіки певної галузі чи компанії.

Однією з ключових можливостей кастомізації є донавчання моделі на специфічних даних. API, що дозволяють донавчання, дозволяють завантажувати додаткові дані, релевантні для конкретного бізнесу, щоб підвищити точність і релевантність відповідей нейроконсультанта. Наприклад, якщо нейроконсультант використовується в онлайн-магазині спортивного одягу, додаткові дані про продукти, відгуки покупців або статистику продажів допоможуть моделі краще розуміти потреби клієнтів і рекомендувати відповідні товари. Цей підхід дає можливість розширити знання моделі за межі загальних даних, які використовуються для початкового навчання.

Іншим важливим аспектом кастомізації є здатність налаштувати тон і стиль відповіді, що дозволяє адаптувати комунікацію під бренд або специфічні очікування аудиторії. Деякі API підтримують налаштування параметрів, які контролюють тональність і формальність відповіді, що особливо важливо для нейроконсультантів, орієнтованих на споживачів різного віку чи культурного фону. Наприклад, для преміум-бренду може бути доречнішим більш офіційний і стриманий стиль відповідей, тоді як для молодіжного бренду можна обрати неформальний, дружній тон.

API, що підтримують кастомізацію, також часто мають можливість інтеграції специфічних бізнес-логік і умов. Це може включати налаштування параметрів моделі для обробки складних логічних запитів або створення унікальних сценаріїв взаємодії. Наприклад, для нейроконсультанта, що працює в банківській сфері, можна налаштувати правила, які допомагають розрізняти типи фінансових запитів або надавати спеціалізовані відповіді залежно від категорії клієнта. Така інтеграція бізнес-логіки допомагає забезпечити високу точність і релевантність відповідей у складних сценаріях.

Ще однією перевагою кастомізації є налаштування мовної та культурної підтримки. Якщо бізнес обслуговує клієнтів з різних країн, важливо, щоб нейроконсультант міг точно відповідати на запити різними мовами, враховуючи культурні особливості кожного ринку. Деякі API дозволяють адаптувати мовні моделі, що робить відповіді нейроконсультанта зрозумілішими та ближчими до специфіки конкретної аудиторії. Це сприяє кращій взаємодії з клієнтами і підвищує їх задоволеність сервісом.

Основні можливості кастомізації API, необхідні для ефективного створення нейроконсультанта, включають:

1. Додавання на спеціалізованих даних, що підвищує точність і адаптованість моделі до бізнес-завдань.
2. Налаштування тону та стилю відповідей для відповідності корпоративній культурі та очікуванням цільової аудиторії.
3. Інтеграція бізнес-логік для забезпечення коректної обробки специфічних запитів і сценаріїв взаємодії.
4. Підтримка багатомовності для зручності клієнтів з різних регіонів та культурних груп.

Сумісність API з іншими компонентами вебдодатку охоплює легкість інтеграції обраного API з наявною інфраструктурою, а також взаємодію з іншими сервісами та компонентами, що вже використовуються в додатку. Основне завдання тут полягає у забезпеченні плавної роботи всіх компонентів, щоб уникнути конфліктів і мінімізувати складність впровадження.

Одним із важливих аспектів сумісності є підтримка стандартних протоколів взаємодії, таких як REST або GraphQL. Більшість сучасних API нейронних мереж працюють через ці протоколи, що значно полегшує їх інтеграцію у вебдодатки, розроблені на основі цих стандартів. Крім того, використання загальноприйнятих форматів передачі даних, як-от JSON або XML, спрощує обмін інформацією між різними компонентами, що може включати бази даних, інтерфейси користувача, або сторонні сервіси. Це дозволяє значно зменшити потребу в додаткових конвертаціях даних та сприяє швидкій інтеграції API в систему.

Ще одним важливим фактором є підтримка різних мов програмування і фреймворків, адже вебдодатки можуть бути побудовані на різних технологічних платформах. Якщо API підтримує популярні мови, такі як JavaScript, Python, Java або PHP, це забезпечує гнучкість для розробників у виборі інструментів. Наявність бібліотек або SDK для конкретних мов також сприяє швидшому впровадженню, оскільки вони містять готові функції для взаємодії з API. Це дозволяє команді заощадити час на розробку власних інтерфейсів та використовувати перевірені методи для інтеграції.

Взаємодія з базами даних є ще одним важливим аспектом сумісності API з іншими компонентами вебдодатку. Для додатків, що активно використовують бази даних для зберігання і обробки інформації, важливо, щоб API міг легко обробляти та зберігати дані в існуючих структурах баз даних. Наприклад, у випадку з нейроконсультантом, результати взаємодії можуть зберігатися в базах даних для подальшого аналізу. Інтеграція з популярними базами даних, такими як MySQL, PostgreSQL, MongoDB, або іншими, забезпечує збереження даних у потрібному форматі та дозволяє уникнути проблем з сумісністю, що виникають при передачі інформації між різними системами.

Також варто врахувати можливість інтеграції з зовнішніми сервісами та API. У багатьох випадках, вебдодаток використовує декілька сторонніх сервісів, таких як платіжні системи, сервіси аналітики, або інструменти для відстеження поведінки користувачів. API повинен бути сумісним з цими сервісами, щоб

забезпечити ефективне управління даними та координацію дій. Наприклад, при роботі з нейроконсультантом, інтеграція з CRM-системами може бути корисною для збереження історії спілкування з клієнтами, а з інструментами аналітики – для відстеження ефективності роботи консультанта.

Ще один важливий аспект – це масштабованість API, тобто його здатність адаптуватися до зростаючих обсягів роботи та нових функцій. Якщо вебдодаток передбачає можливість розширення функціоналу або зростання користувацької бази, обраний API повинен бути готовим до підвищених навантажень. Сумісність з хмарними платформами, такими як AWS, Google Cloud або Azure, може надати необхідну масштабованість та забезпечити стабільну роботу системи при збільшенні обсягів даних. Це дозволяє гнучко адаптувати інфраструктуру під постійні зміни, що важливо для підтримки високої продуктивності вебдодатку.

2.3 Особливості застосування OpenAI API

OpenAI API – це інтерфейс програмування, який надає розробникам доступ до потужних моделей AI, зокрема серії GPT. Цей API дозволяє інтегрувати можливості NLP, генерації текстів, перекладу, аналізу емоцій, кодування та інших функцій у веб та мобільні додатки. OpenAI API працює через REST-запити, що дозволяє легко та ефективно використовувати його в різних платформах і мовах програмування.

OpenAI API пропонує інтернет-магазинам сучасні інструменти для автоматизації та персоналізації взаємодії з клієнтами. Цей API дозволяє створювати персоналізовані рекомендації товарів на основі індивідуальних уподобань користувачів, завдяки чому покупці швидше знаходять необхідні товари. Іншим важливим напрямком використання API є автоматизація обслуговування клієнтів – завдяки чат-ботам на базі GPT моделей компанії можуть обробляти сотні запитів одночасно, швидко та якісно.

Зазначений API також допомагає автоматизувати створення контенту: це стосується генерації описів товарів, SEO-оптимізованих текстів і публікацій для залучення клієнтів. Цей підхід дозволяє значно заощадити час на ведення контенту, зберігаючи при цьому його високу якість. OpenAI API корисний і для аналітики: він дозволяє виявляти тренди та прогнозувати попит, що допомагає краще планувати запаси та маркетингові кампанії [23].

OpenAI API надає розширені можливості налаштування параметрів, що дозволяє адаптувати відповіді моделей для конкретних бізнес-завдань. На рис. 2.2 наведено приклад коду, який генерує текстову відповідь та опис коду з налаштуванням основних параметрів OpenAI API.

```
import OpenAI from "openai";

const openai = new OpenAI({
  apiKey: process.env.REACT_APP_OPENAI_API_KEY,
});

const response = openai.Completion.create({
  engine="text-davinci-003",
  prompt="Напиши інформацію про спортивний одяг",
  temperature=0.7,
  max_tokens=150,
  top_p=0.9,
  frequency_penalty=0.5,
  presence_penalty=0.6,
  stop=["\n", "Одяг"]
})

console.log(response.choices[0].text);
```

Рисунок 2.2 – Приклад коду з налаштуванням основних параметрів

Import OpenAI from «openai» – імпортує бібліотеку OpenAI, яка дозволяє взаємодіяти з API. Ця бібліотека містить методи для налаштування й виконання запитів до моделей та використовується для інтеграції моделей OpenAI у різних середовищах програмування.

Const openai – створює змінну, в якій необхідно вказати значення ключа apiKey від API. Цей ключ необхідно створити в профілі користувача OpenAI та

вставити як значення до ключа `apiKey`. Це потрібно для того, щоб забезпечити аутентифікацію користувача при надсиланні запитів до сервера OpenAI. Без цього ключа сервер не зможе ідентифікувати запити та надавати доступ до функціоналу API. Ключ також контролює використання ресурсів та запобігає несанкціонованому доступу, забезпечуючи безпеку даних і доступу до моделей.

`Const response` – змінна, де вказуються різні параметри для генерації відповіді та зберігається відповідь від API. У цьому рядку запит до API здійснюється через метод `openai.Completion.create()`, а результат зберігається в змінну `response` для подальшої обробки.

`Engine=«text-davinci-003»` – визначає модель, яку буде використано для відповіді на запит. В цьому випадку, модель `text-davinci-003` призначена для генерації розгорнутих і контекстуально багатих відповідей.

`Prompt=«Напиши інформацію про спортивний одяг»` – текст запиту, або «промпт», який передається в API. Це конкретне завдання, на яке модель швидко дає відповідь. `Prompt` можна адаптувати залежно від вимог користувача, уточнюючи контекст або очікуваний формат відповіді.

`Temperature` – визначає ступінь креативності відповіді. Діапазон значень від 0 до 1. Низькі значення роблять відповіді більш передбачуваними і логічними, а високі – більш варіативними та креативними.

`Max-tokens` – максимальна кількість слів чи символів у відповіді. Дозволяє обмежувати довжину відповіді, що корисно для коротких відповідей або структурованих діалогів.

`Top-p` – це параметр, який застосовується для контролю генеративної поведінки моделі як альтернативний підхід до `temperature`. Його значення варіюється від 0 до 1 і визначає, які варіанти відповіді модель розглядатиме. Цей метод сприяє збереженню балансу між креативністю та релевантністю: високі значення `top-p` дозволяють отримати більш варіативні відповіді, а нижчі – обмежують модель до більш визначених і прямолінійних відповідей [24].

`Frequency-penalty` – контролює частоту повторення слів у відповіді. Значення може варіюватися від 0 до 2, де 0 означає відсутність контролю, а

значення ближче до 2 знижують імовірність повторення тих самих слів, що важливо для унікальних відповідей.

Presence-penalty – стимулює включення нових тем і запобігає замкненому колу відповіді. Діапазон – від -2 до 2. Значення ближче до 2 заохочують модель вводити нові ідеї або слова, що допомагає у творчих завданнях.

Stop – задає слова або символи, які сигналізують завершення відповіді. Наприклад, встановлення `stop=[«\n», «Одяг»]` дозволить завершити відповідь після виявлення цих символів.

`console.log(response.choices[0].text)` – код, який повертає згенерований текст у консоль.

OpenAI API пропонує передові можливості для роботи з текстом, які базуються на сучасних алгоритмах NLP. Функціональність API охоплює генерацію тексту, редагування текстів і їх аналіз, що дозволяє ефективно автоматизувати бізнес-процеси, пов'язані зі створенням контенту, комунікацією та обробкою даних.

Генерація тексту – одна з основних функцій OpenAI API, яка дозволяє створювати контент у відповідь на заданий запит. Ця можливість ґрунтується на використанні мовних моделей, таких як GPT-4, які аналізують контекст запиту та створюють осмислений текст. Наприклад, в електронній комерції API може автоматично створювати опис товарів, які враховують його унікальні характеристики, або формувати персоналізовані маркетингові матеріали. Генерація тексту також корисна для автоматизації клієнтського обслуговування: чат-боти, побудовані на основі OpenAI API, можуть формулювати відповіді на запитання клієнтів із високою точністю та релевантністю. Модель дозволяє контролювати тональність тексту, стиль і навіть обсяг відповіді, що робить її адаптивною до різних завдань, таких як складання листів, написання статей чи створення сценаріїв [25].

Функція редагування, яку надає OpenAI API через endpoint `edits`, дозволяє виправляти, уточнювати або переформулювати текст на основі заданих інструкцій. Ця можливість особливо корисна в задачах, пов'язаних із

поліпшенням граматики, стилістики чи структури тексту. Наприклад, API може виправити граматичні помилки у великих текстах, перевести текст у більш офіційний стиль або адаптувати його до специфічних вимог клієнта. Вказуючи інструкції для моделі, користувач може створити текст, який відповідає конкретним потребам. Наприклад, інструкція «Зроби текст більш професійним» змусить модель переглянути й відредагувати матеріал відповідно до встановлених формальних стандартів.

Функції аналізу тексту в OpenAI API забезпечують розуміння контенту через класифікацію, витяг ключових слів, оцінку тональності та семантичний аналіз, що дозволяє автоматизувати широкий спектр бізнес-завдань. Наприклад, класифікація тексту допомагає сортувати запити клієнтів за категоріями, такими як замовлення, технічна підтримка чи загальні питання, що полегшує їх подальшу обробку. Витяг ключових слів визначає основні терміни чи фрази, які можна використати для SEO-оптимізації чи аналізу відгуків. Оцінка тональності дозволяє виявляти емоційне забарвлення тексту, що важливо для моніторингу настрою клієнтів у відгуках або соціальних мережах. Семантичний аналіз розкриває контекст і смислові зв'язки, допомагаючи виявляти тенденції, автоматизувати системи рекомендацій чи знаходити приховані взаємозв'язки в текстових даних. Крім цього, API підтримує обробку текстів різними мовами, що робить його універсальним рішенням для глобальних бізнесів [26].

Робота з зображеннями в OpenAI API охоплює функції, які дозволяють створювати, редагувати та аналізувати графічний контент за допомогою моделей AI. Ці функції є надзвичайно корисними для широкого кола застосувань, таких як дизайн, автоматизація контенту, генерація нових ідей або навіть глибокий аналіз візуальної інформації для бізнес-задач. Завдяки цьому, бізнеси можуть швидко адаптуватися до змінюваних вимог ринку, створюючи індивідуалізовані маркетингові матеріали та оптимізуючи візуальний контент.

Однією з ключових можливостей API є здатність створювати зображення на основі текстових запитів, завдяки моделям на зразок DALL·E. Користувач задає текстовий опис, і модель генерує відповідне зображення (рис. 2.3), яке

відповідає вказаним параметрам. Це дозволяє автоматизувати процес створення графічного контенту, суттєво скорочуючи час та витрати на розробку. Це особливо корисно для створення маркетингових матеріалів, таких як банери, афіші або рекламні зображення, оскільки дає змогу отримати оригінальний контент без залучення дизайнерів. Крім того, можливість швидко створювати візуалізації є незамінною під час прототипування продуктів, дозволяючи перевірити ідеї та концепції ще на етапі планування.



Рисунок 2.3 – Приклад згенерованих зображень за запитом
«крісло у формі авокадо» [27]

OpenAI API підтримує редагування графіки, що дозволяє вносити зміни до існуючих зображень. Наприклад, користувач може змінити фон, додати нові елементи або покращити якість зображення. Ця функція широко використовується в електронній комерції для покращення фотографій товарів або у творчих індустріях для швидкого оновлення дизайну. Також OpenAI надає можливості аналізу візуального контенту, наприклад, для виявлення об'єктів, оцінки композиції, класифікації або OCR. Це дозволяє автоматизувати завдання, пов'язані з обробкою фотографій чи документів, і створює нові можливості для бізнесу, як-от аналіз поведінки клієнтів на основі зображень продуктів, що вони переглядають, а також персоналізацію рекомендацій і покращення маркетингових стратегій.

Залежно від текстових запитів, OpenAI API можна інтегрувати з іншими технологіями для генерації анімації або навіть 3D-моделей. Хоча такі функції вимагають більш складної інтеграції, вони ідеально підходять для геймдеву або розробки інтерактивного вмісту. OpenAI API підтримує широкі можливості

налаштування, що дозволяє адаптувати генерацію або аналіз зображень під конкретні задачі. Наприклад, у запитах можна уточнювати художній стиль, колірну гамму, роздільну здатність чи навіть тип графіки (цифровий арт, реалістичне зображення, 3D) [25].

Робота з кодом за допомогою OpenAI API відкриває широкі можливості для автоматизації програмування, налагодження, пояснення та розширення коду. Завдяки використанню спеціалізованих моделей, таких як Codex, API може генерувати фрагменти програмного забезпечення на основі текстових запитів, редагувати існуючий код, пояснювати складні логічні конструкції та навіть створювати документацію і тести [28].

Основний функціонал передбачає генерацію коду у відповідь на запити. Наприклад, користувач може описати проблему, і API надає готовий код для її розв'язання. Це особливо корисно для прискорення розробки та створення шаблонів. API також дозволяє автоматично оптимізувати код або виправляти помилки. Використовуючи його функції, можна покращити ефективність алгоритмів, запропонувати більш структуровані рішення або додати коментарі до складних ділянок. Крім того, API підтримує інтеграцію з існуючими інструментами, такими як системи управління версіями або IDE, полегшуючи автоматизацію робочих процесів. Наприклад, через API можна автоматично генерувати unit-тести, інтеграційні тести або сценарії перевірки коду. Це знижує ручну працю розробників і дозволяє їм зосередитися на ключових завданнях. Особливо корисною є функція пояснення коду, яка допомагає зрозуміти логіку та структуру програмного забезпечення. Це може бути цінним як для новачків, так і для досвідчених розробників, які працюють із чужим кодом. API здатен не лише роз'яснювати, як працюють певні алгоритми, але й пропонувати рекомендації для їхнього вдосконалення [29].

Робота з OpenAI API також передбачає написання документації. API може генерувати описи функцій, пояснення змінних та загальні огляди архітектури коду, що значно прискорює процес документування програм. Завдяки цьому, програмісти мають змогу зберігати час і водночас забезпечувати високу якість

технічних матеріалів. Використання API можливе через простий код, який інтегрується з різними мовами програмування. Наприклад, ініціалізація API вимагає ключ доступу, після чого можна відправляти запити на виконання завдань із генерації, редагування або аналізу коду. OpenAI API також надає доступ до інструментів обробки тексту, які можуть застосовуватися до програмного коду, забезпечуючи його узгодженість із описом задачі [30].

OpenAI API надає можливість працювати з аудіо та мовою, що дозволяє автоматизувати широкий спектр завдань, пов'язаних із розпізнаванням, генерацією та аналізом мовлення. OpenAI API підтримує розробку додатків, здатних працювати з голосовим введенням, синтезом мовлення та обробкою текстових даних. Однією з ключових функцій є перетворення голосу на текст. За допомогою технологій автоматичного розпізнавання мовлення, API може точно трансформувати голосове введення в текст, що ідеально підходить для стенографії, нотаток або інтерактивних систем, таких як голосові помічники. Інструменти OpenAI здатні розуміти багатомовний контент, що значно спрощує роботу з глобальними клієнтами. API також дозволяє виконувати синтез тексту в мовлення, що використовується для створення озвучених повідомлень, інтерактивних інтерфейсів або навчальних програм. Це забезпечує інтерактивність та покращує користувацький досвід для додатків, які спілкуються з клієнтами через аудіо [31].

2.4 Аналіз конкурентів OpenAI API та порівняння їх функціональності

У ході виконання дослідження основна увага була зосереджена на практичному аналізі двох провідних платформ – Google Gemini API та OpenAI API. Ці технології були вибрані завдяки їхній сучасності, доступності та відповідності до завдань проєкту. Обидва API забезпечують широкий спектр функцій, зокрема генерацію тексту, аналіз даних, створення інтелектуальних

систем рекомендацій та обробку великих обсягів інформації, що робить їх ідеальними для практичного тестування та впровадження.

Google Gemini API, як новітня платформа від Google, пропонує революційний підхід до обробки тексту, зображень та інших даних завдяки інтеграції багатовекторного аналізу та вдосконалених мовних моделей [32]. Ця технологія дозволяє значно покращити точність інтерпретації контексту та створення змістовних і релевантних відповідей, розширюючи можливості AI в різних сферах. Завдяки багатовекторному підходу, Gemini API здатний одночасно аналізувати різноманітні типи даних, такі як: текстові, зображення, аудіо та мультимедіа [33, 34].

OpenAI API, у свою чергу, вирізняється зрілим функціоналом, перевіреним надійністю та широкими можливостями для кастомізації моделей під специфічні потреби користувачів. Практичне тестування цих платформ дозволило оцінити їхні переваги, обмеження та сумісність з іншими інструментами, що стало основою для подальшого розроблення та обґрунтування вибору технологій у рамках проєкту [33, 35].

За різними параметрами було порівняно та обґрунтовано Google Gemini API та OpenAI API (додаток А, табл. А.1).

В результаті наведених обґрунтувань найбільше підходить OpenAI API. На основі проведеного аналізу, основним критерієм вибору API для кваліфікаційної роботи була доступність безкоштовного або економічно вигідного використання. Google Gemini API, який надає безкоштовний доступ, здавалося б, був ідеальним варіантом. Однак, на момент підготовки роботи, цей API не був доступним для використання в Україні через географічні обмеження. Використання Gemini API можливе лише через віртуальну приватну мережу (Virtual Private Network, VPN), що ускладнює процес інтеграції та розробки.

Зважаючи на ці обмеження, було прийнято рішення використовувати OpenAI API, який хоч і передбачає певну оплату за використання, проте є широко доступним в Україні, підтримується надійним сервісом та забезпечує стабільну роботу без додаткових технічних труднощів.

Висновки до розділу 2

У цьому розділі було проведено детальний аналіз технологій для реалізації нейроконсультанта, зосереджений на виборі найкращого API для виконання поставлених завдань. Зокрема, було детально розглянуто та проаналізовано два ключових інструменти – OpenAI API та Google Gemini API.

У ході аналізу було обґрунтовано вибір OpenAI API як основного інструмента для реалізації нейроконсультанта. Хоча Google Gemini API демонструє конкурентні можливості, але його недоступність в Україні через географічні обмеження створює значні технічні складнощі, що унеможлиблює його безпосереднє використання без додаткових засобів, таких як VPN. OpenAI API, своєю чергою, є доступним для українських користувачів, забезпечує високу якість виконання завдань, гнучкість інтеграції та пропонує широкий спектр функціональності для створення ефективного нейроконсультанта. Окрім цього, у процесі аналізу було відзначено переваги OpenAI API, які включають можливість обробки великих текстових задач, адаптивність до потреб користувача та підтримку мультимодальних функцій, таких як генерація тексту та створення зображень. Це дозволяє значно розширити функціональні можливості нейроконсультанта, забезпечуючи як високий рівень автоматизації, так і персоналізацію взаємодії з користувачами.

Таким чином, обраний API не лише відповідає технічним вимогам для реалізації проєкту, але й є оптимальним рішенням для довгострокового впровадження системи, враховуючи його масштабованість, доступність та відповідність сучасним технологічним стандартам.

РОЗДІЛ 3

ПРАКТИЧНА РЕАЛІЗАЦІЯ НЕЙРОКОНСУЛЬТАНТА ДЛЯ ІНТЕРНЕТ-МАГАЗИНУ

3.1 Розробка нейроконсультанта з використанням OpenAI API

Щоб створити нейроконсультанта за допомогою OpenAI API, спочатку виконується реєстрація на платформі. Для цього здійснюється перехід на офіційний сайт OpenAI [36]. На головній сторінці сайту розміщені опції для створення нового облікового запису або входу в систему, якщо обліковий запис уже існує (рис. 3.1).

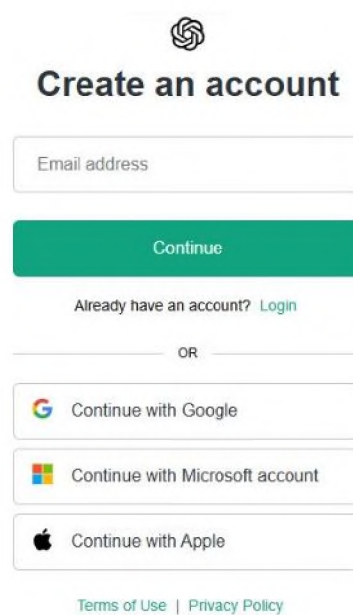
The image shows the OpenAI registration page. At the top is the OpenAI logo and the text 'Create an account'. Below this is a text input field labeled 'Email address'. Underneath the field is a green button labeled 'Continue'. Below the button is a link that says 'Already have an account? Login'. A horizontal line with the word 'OR' in the center separates this from the social login options. There are three buttons: 'Continue with Google' (with the Google logo), 'Continue with Microsoft account' (with the Microsoft logo), and 'Continue with Apple' (with the Apple logo). At the bottom, there are two links: 'Terms of Use' and 'Privacy Policy'.

Рисунок 3.1 – Вигляд форми реєстрації на сайті OpenAI [36]

Для створення нового облікового запису використовується кнопка Sign Up. Після натискання відкривається стандартна форма реєстрації, де зазначається електронна адреса, пароль та інша базова інформація. Також передбачена можливість реєстрації через облікові записи Google або Microsoft, що дозволяє уникнути введення даних вручну. Після заповнення форми підтверджується створення облікового запису через кнопку Create Account.

На вказану електронну пошту надходить лист із посиланням для підтвердження реєстрації. Після переходу за цим посиланням обліковий запис активується, і з'являється доступ до особистого кабінету користувача. Якщо підтвердження успішне, користувач автоматично перенаправляється до головного меню платформи для подальшої роботи (рис. 3.2).

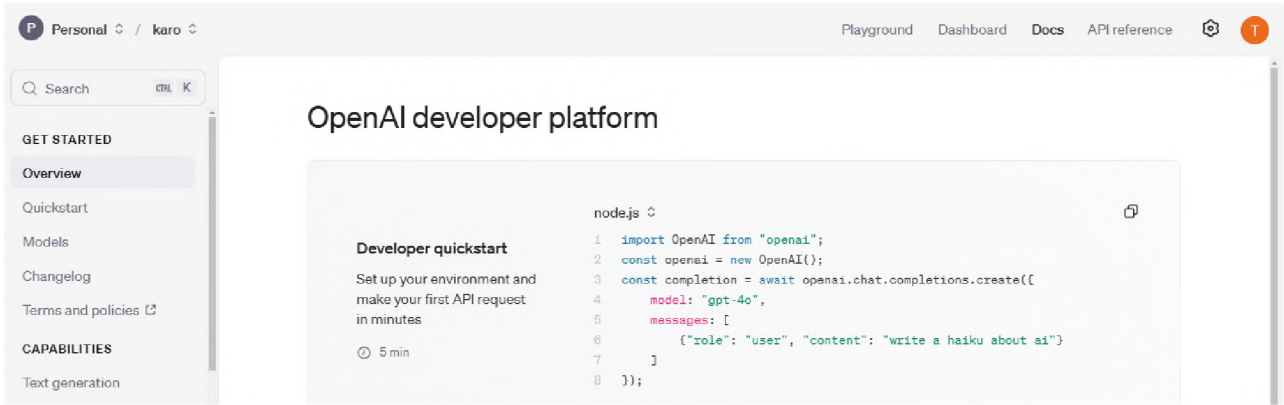


Рисунок 3.2 – Вигляд головного меню платформи OpenAI [36]

Для отримання доступу до API виконується генерація унікального API-ключа. У розділі API Keys передбачена функція створення нового ключа через кнопку Create new secret key (рис. 3.3). Система створює унікальний ключ, який необхідно скопіювати та зберегти у безпечному місці.

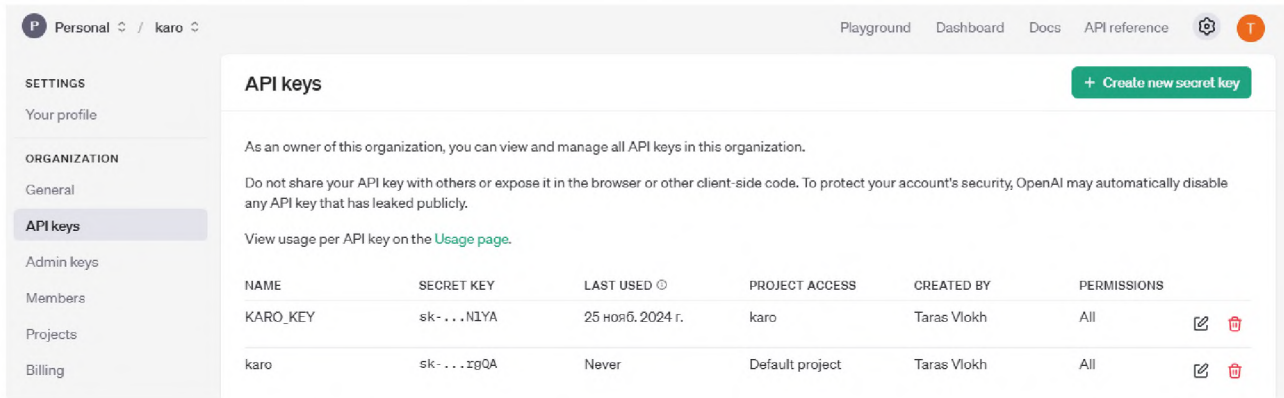


Рисунок 3.3 – Вигляд створеного API ключа OpenAI [36]

Цей ключ використовується для автентифікації під час запитів до API. Зберігати ключ рекомендовано у захищеному місці, наприклад, у

конфігураційному файлі проєкту або у спеціалізованому менеджері паролів. Після генерації ключ стає недоступним для повторного перегляду, тому його втрата вимагатиме створення нового ключа.

Наступний етап передбачає поповнення рахунку для активації доступу до API. Увійшовши в особистий кабінет, потрібно відкрити розділ Billing. У цьому розділі додаються платіжні дані, такі як номер банківської картки, термін дії та CVV-код. Дані зберігаються після натискання кнопки Save. Після цього у вкладці Deposit Funds зазначається сума поповнення, яка має становити не менше \$ 5. Після підтвердження платежу через кнопку Confirm Payment кошти зараховуються на баланс облікового запису (рис. 3.4).

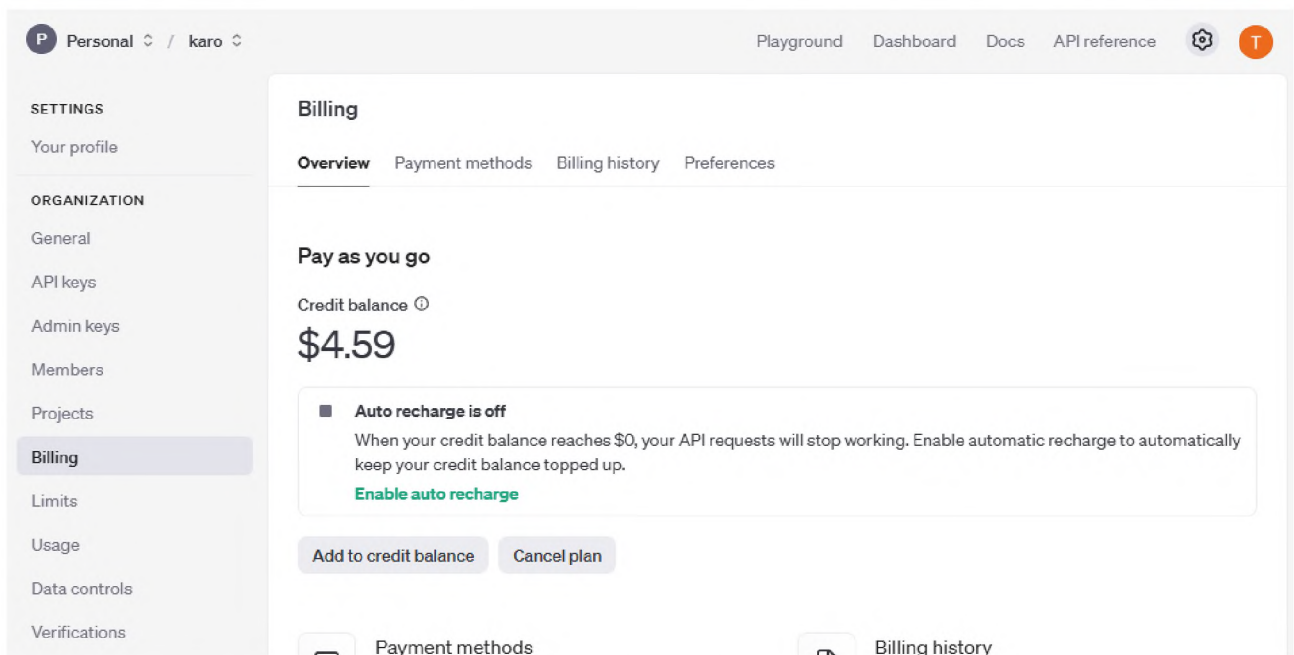


Рисунок 3.4 – Вигляд поповненого рахунку OpenAI [36]

Після виконання всіх зазначених кроків API-ключі OpenAI готові до використання для інтеграції у проєкт, оскільки вони вже згенеровані, а баланс поповнений для забезпечення стабільної роботи. На цьому етапі можна переходити до розробки функціонального коду нейроконсультанта, інтегруючи API у середовище розробки, та налаштовувати його відповідно до визначених вимог проєкту. Це дозволить ефективно використовувати можливості OpenAI для створення продукту.

3.2 Опис структури коду та основних функціональних компонентів

Для початку необхідно взяти ключ з сайту OpenAI (див. рис. 3.3). Після цього потрібно створити та заповнити структуру коду для отримання відповіді від API. На рис. 3.5 зображено початкову змінну для звернення до OpenAI API.

```
const openai = new OpenAI({
  apiKey: process.env.REACT_APP_OPENAI_API_KEY,
});
```

Рисунок 3.5 – Звернення до OpenAI API через створений ключ

Для передачі значення ключа було використано `.env`, який є допоміжним пакетом для роботи зі змінними середовища. Це дозволяє надійно захистити API-ключ від несанкціонованого доступу, оскільки він не зберігається безпосередньо у коді програми та залишається прихованим від сторонніх користувачів [37].

Після цього виконується асинхронна функція (рис. 3.6), яка виконується після натискання на кнопку «Знайти» на сайті (див. рис. 3.14).

```
const handleSubmitOpenAI = async () => {
  try {
    setFetching(true);

    const completion = await openai.beta.chat.completions.parse({
      model: "gpt-3.5-turbo",
      messages: [
        { role: "system", content: "Extract the filter information." },
        { role: "user", content: prompt + prompt2 },
      ],
      temperature: 0,
      top_p: 0,
    });
  }
};
```

Рисунок 3.6 – Початок виконання функції `handleSubmitOpenAI`

Параметри `temperature` та `top_p` мають значення нуль, через те що, потрібно повертати конкретні дані залежно від запиту користувача. Детальне пояснення основних параметрів API від OpenAI наведено в підрозділі 2.3. Змінна `prompt` зберігає текст, який ввів користувач для пошуку товару.

Змінна `prompt2` зберігає додаткові умови до запиту користувача, який дає зрозуміти нейроконсультанту, що йому потрібно повернути у відповідь, а саме, об'єкт конкретних даних. Також в цій змінній вказано, які є можливі значення до відповідних ключів об'єкта (рис. 3.7).

```
const prompt2 = `
You need to return an object as shown below. All keys need to be provided and should be empty
"category": [],
"brand": [],
"gender": [],
"childGender": [],
"childAge": [],
"type": [],
"color": [],
"sort": "",
"oldPrice": [],
"size": [],
each key can contain the data below (the default for displaying all value keys is empty):
category: [${categoryFilter}]
brand: [${brandArr}]
gender: [${genderArr}] (ALWAYS EMPTY WHEN NOT MENTIONED)
childGender: [${childGender}] (if the words boys or girls is mentioned, key "gender" must be s
childAge: [${childrenAge}] (if the child age is mentioned, key "gender" must be specified wit
type: [${clothesType}] (each category has different values: "clothes": ${clothesType} ; "shoe
color: [${colorArr}]
sort: ${sortFilterArr} (only one value)
oldPrice: ["oldPrice"] (when mention about discount)
size: [] (value for size not included)

(Only assign a value to the "gender" key if one of the following is explicitly mentioned IN t
If the request isn't about a topic for searching clothes, remove object above and return unc
`
```

Рисунок 3.7 – Вигляд змінної `prompt2`

В змінній `prompt2` описано, з якими саме ключами та ймовірними значеннями в них повинен повертати нейроконсультант. Ключ `size` не включений в список через дуже велику вибірку розмірів, що додає дуже багато символів для пошуку, що є економічно не вигідно, також складно пояснити нейроконсультанту для якого типу товару, які є ймовірні розміри, бо на різні типи товару розмір може мати такі самі значення та через це нейроконсультант постійно робить помилки, якщо запит стосується розміру.

Далі в змінній `resultJSON` виконується перетворення текстової відповіді від API у формат JSON (рис. 3.8). JSON – це формат обміну даними, який дозволяє організувати інформацію у вигляді структурованих об'єктів з ключами і

значеннями [38]. Це спрощує подальшу роботу з даними, оскільки їх можна легко зберігати, передавати та обробляти. У даному випадку це дозволяє зберегти отриману інформацію в об'єкт, що використовується для фільтрації товарів на сайті за визначеними параметрами (рис. 3.9).

```
const resultJSON = JSON.parse(completion.choices[0].message.content);
console.log(resultJSON);
setFilter(resultJSON);

if (FilterCondition(Object.entries({
  ...resultJSON
}))) {
  dispatch(setPromptFilter(resultJSON))
  navigate("/allitems")
}
else {
  setOutput("Немає в наявності такого товару або сталася помилка");
}
```

Рисунок 3.8 – Продовження виконання функції handleSubmitOpenAI

```
const [filter, setFilter] = useState({
  category: [],
  brand: [],
  gender: [],
  childGender: [],
  childAge: [],
  type: [],
  color: [],
  sort: "",
  oldPrice: [],
  size: [],
});
```

Рисунок 3.9 – Змінна filter, яка є об'єктом та зберігає дані для фільтрації товарів за певними критеріями

Функція FilterCondition (див. рис. 3.8), як параметр приймає значення з ключів об'єкта та всередині цієї функції йде перевірка чи є за такими фільтрами товари в базі даних чи ні (додаток Б, рис. Б.1-Б.6). Якщо функція знаходить товар, то зберігає отриманий об'єкт від нейроконсультанта в сховище dispatch(setPromptFilter(resultJSON)), яке дає змогу отримати доступ до цього об'єкту з фільтрами з будь-якого файлу та перенаправляє користувача на нову сторінку сайту із застосованими фільтрами, це виконує команда

`navigate(«/allitems»)`). Якщо немає знайдених товарів, то на сайті відображається текст: «Немає в наявності такого товару або сталася помилка». В кінці асинхронної функції здійснюється перевірка блоків `catch` та `finally` для обробки помилок і завершення запиту коректним чином (рис. 3.10).

```

    } catch (error) {
      setOutput("Немає в наявності такого товару або сталася помилка");
      console.error(error)
    } finally {
      setFetching(false)
    }
  }

```

Рисунок 3.10 – Кінець виконання функції `handleSubmitOpenAI`

Блок `catch` обробляє помилки, що виникають під час виконання функції. Наприклад, якщо користувач ввів некоректний запит, наприклад «2+2», або виникла проблема із запитом до API, `catch` перехопить цю помилку і відобразить повідомлення «Сталася помилка, повторіть запит знову».

Блок `finally` виконується завжди, незалежно від того, чи виникла помилка під час виконання коду, чи ні. У цьому випадку він використовується для активації кнопки пошуку після завершення операції, забезпечуючи її доступність для наступного використання.

Після успішного перенаправлення на нову сторінку, виконується окрема частина коду, де є нова змінна `filter`, куди зберігається об'єкт з даними для фільтрації, який знаходиться у сховищі, що був переданий до цього (рис. 3.11).

```

const [filter, setFilter] = useState({
  ...promptFilter
});

```

Рисунок 3.11 – Передані дані об'єкта в нову змінну `filter`

Далі йде функція `useEffect`, яка виконується один раз, відразу після того, як користувач потрапив на цю сторінку (рис. 3.12). В цій функції є інша функція `FilterCondition`, яка вже була згадана раніше (додаток Б, рис. Б.1-Б.6), але

відмінність в тому, що в цій функції ми отримуємо масив з товарами та відображаємо їх, а в попередньому випадку була лише перевірка розміру цього масиву, щоб запевнитися у наявності товарів.

```
useEffect(() => {
  FilterCondition( Object.entries({
    ...filter
  }));
}, []);
```

Рисунок 3.12 – Вигляд функції `useEffect`, яка повертає масив з товарами

Також, в функції `FilterCondition` заповнюються блоки з кожним окремим фільтром, щоб візуально було видно, які фільтри застосовані (рис. 3.13) та користувач міг змінити їх.

```
<SelectCountItems
  filterStatus={"filter"}
  arr={categoryFilterUA}
  arrFilter={categoryFilter}
  name={"Категорія"}
  keyName={"category"}
  filter={filter}
  setFilter={setFilter}
  filterCheck={filterCheck}
/>
```

Рисунок 3.13 – Приклад блоку для фільтрації товарів з типом «Категорія»

Блок `<SelectCountItems>` є окремим компонентом, що реалізований у вигляді окремого файлу зі своєю власною структурою мовою розмітки гіпертексту (Hypertext Markup Language, HTML) та логікою, яка також дає можливість реалізовувати каскадні таблиці стилів (Cascading Style Sheets, CSS). Використання цього компонента в коді через відповідний тег дозволяє створювати його копії в різних частинах додатка без необхідності повторного дублювання коду. Цей підхід забезпечує ефективність, зручність у розробці та підтримці проєкту, оскільки будь-які зміни, внесені до компонента `<SelectCountItems>`, автоматично застосовуватимуться до всіх його копій.

Ця властивість є однією з основних переваг React, що базується на компонентному підході. Завдяки цьому підходу код стає більш модульним, легко масштабованим і зрозумілим, а повторне використання компонентів суттєво знижує ймовірність помилок і скорочує час на розробку [39, 40].

Щоб була змога користуватися нейроконсультантом, потрібно графічно відобразити його із зрозумілим інтерфейсом (рис. 3.14). На рис. 3.15 наведено код, який відображає блок з нейроконсультантом.

ПОШУК ТОВАРІВ ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ

Введіть запит

ЗНАЙТИ

Рисунок 3.14 – Вигляд блоку з нейроконсультантом

```
return (
  <div className={style.chat}>
    <h2 className={style['title']}>ПОШУК ТОВАРІВ ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ</h2>
    <input
      className={style['input-field']}
      type='text'
      name='prompt'
      id='prompt'
      value={prompt}
      onChange={promptChange}
      placeholder='Введіть запит'
      maxLength={250}
    />
    <button
      className={({fetching || prompt.length == 0} ?
        style['main-btn-disabled'] : style['main-btn-inverted'])
      disabled={fetching || prompt.length == 0}
      onClick={handleSubmitOpenAI}
    >
      <p>{fetching ? "Йде пошук" : "Знайти"}</p>
    </button>
    <br />
    <div className={style['error']}>
      {output}
    </div>
  </div>
)
```

Рисунок 3.15 – Вигляд коду для відображення нейроконсультанта

У результаті користувач може взаємодіяти зі створеним нейроконсультантом для отримання персоналізованої допомоги.

3.3 Тестування та оцінка ефективності нейроконсультанта

При переході на сайт інтернет-магазину спортивного одягу [41], користувач отримує інтуїтивно зрозумілий інтерфейс, який забезпечує легкий доступ до всіх функцій системи. На головній сторінці представлені ключові елементи та створений нейроконсультант (рис. 3.16).

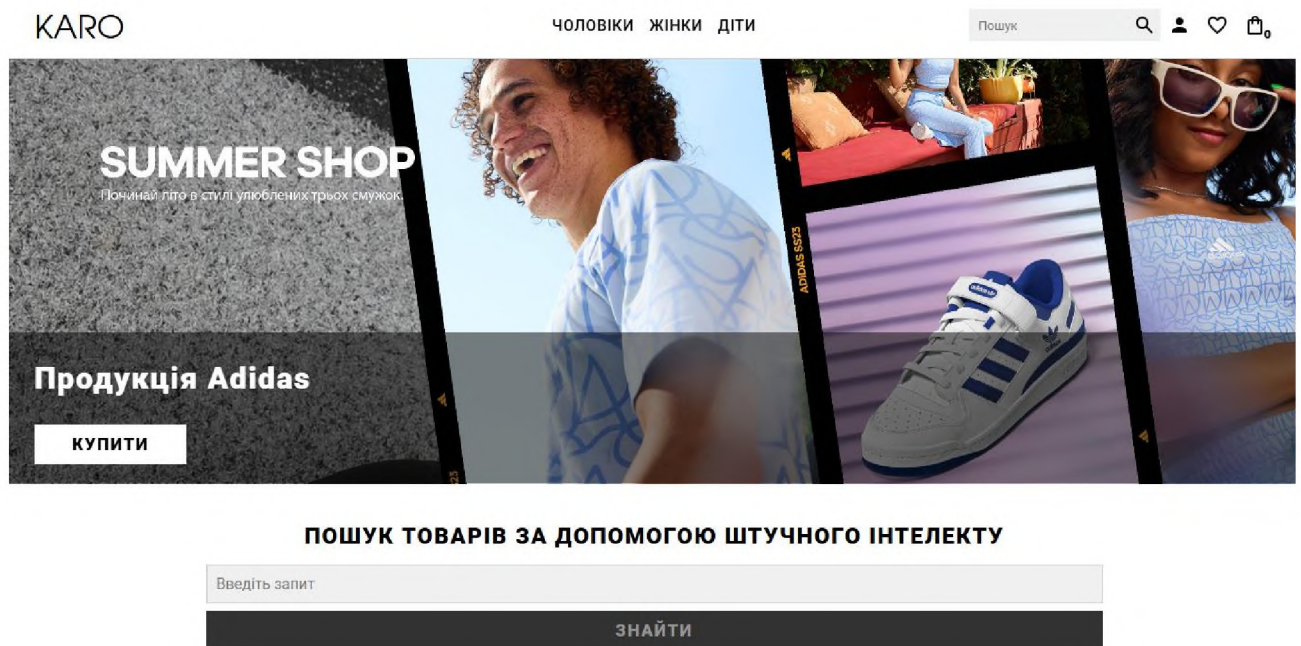


Рисунок 3.16 – Вигляд головної сторінки інтернет-магазину спортивного одягу

У межах пункту було проведено тестування двох моделей OpenAI: GPT-4o та GPT-3.5-turbo. Основна мета тестування полягала у визначенні функціональних можливостей моделей, а також у порівнянні їхньої вартості для використання в розробці нейроконсультанта.

Під час тестування було виявлено, що GPT-4o має суттєву перевагу у форматі відповідей: ця модель здатна безпосередньо генерувати відповіді у форматі JSON. Така функціональність значно спрощує інтеграцію з іншими частинами системи. Наприклад, це дозволяє легко передавати структуровані дані для фільтрації товарів на сайті, автоматизації аналізу інформації або персоналізації рекомендацій для клієнтів. У свою чергу, GPT-3.5-turbo такої

можливості не надає, що вимагає додаткової обробки текстових відповідей і їхнього програмного перетворення у формат JSON. Це може значно збільшувати навантаження на систему та вимагати додаткових технічних ресурсів для реалізації функціоналу.

Окрім функціональних можливостей, було проведено аналіз вартості використання моделей. Результати тестування показали значну різницю у вартості запитів. Використання GPT-4o коштує \$ 0.38 за 19 запитів (рис. 3.17), тоді як GPT-3.5-turbo коштує лише \$ 0.03 за 125 запитів (рис. 3.18). Така різниця є суттєвою, особливо якщо передбачається масове використання моделі. У довгостроковій перспективі GPT-3.5-turbo виявляється значно вигіднішою з економічної точки зору.

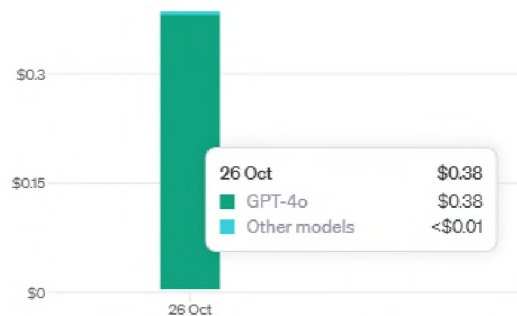


Рисунок 3.17 – Вартість моделі GPT-4o за 19 запитів

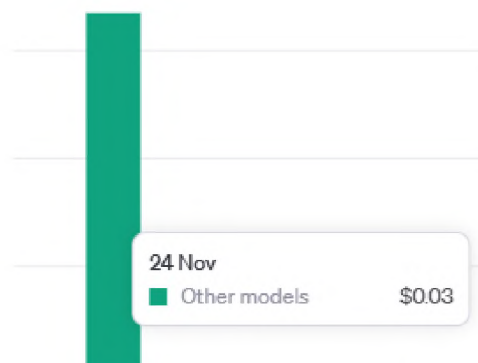


Рисунок 3.18 – Вартість моделі GPT-3.5-turbo за 125 запитів

На основі тестування можна зробити висновок, що GPT-4o, хоча і забезпечує розширені можливості, є дорожчим рішенням. Ця модель може бути

виправданою для проєктів із великим бюджетом або специфічними вимогами, які потребують використання JSON-відповідей безпосередньо. Натомість GPT-3.5-turbo пропонує оптимальний баланс між функціональністю та вартістю, що робить її ідеальним вибором для проєктів з обмеженими фінансовими ресурсами чи високими вимогами до обсягу запитів.

Якщо користувач введе запит у блоці «Пошук товарів за допомогою штучного інтелекту» (рис. 3.19) і натисне кнопку пошуку, нейроконсультант розпочне обробку введеної інформації.

ПОШУК ТОВАРІВ ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ

Купити чорну куртку

ЗНАЙТИ

Рисунок 3.19 – Введений текст для запиту

Після обробки запиту нейроконсультант застосовує визначені фільтри для пошуку відповідних товарів (рис. 3.20). У цьому процесі враховуються всі задані критерії, такі як категорія, ціновий діапазон або інші параметри. Це дозволяє точно підібрати асортимент, що відповідає запиту користувача.

```

▶ brand: []
▶ category: ['clothes']
▶ childAge: []
▶ childGender: []
▶ color: ['black']
▶ gender: []
▶ oldPrice: []
▶ size: []
  sort: ""
▶ type: ['jacket']
▶ [[Prototype]]: Object
Matching items count: 2
  
```

Рисунок 3.20 – Згенерований об’єкт з фільтрами від нейроконсультанта

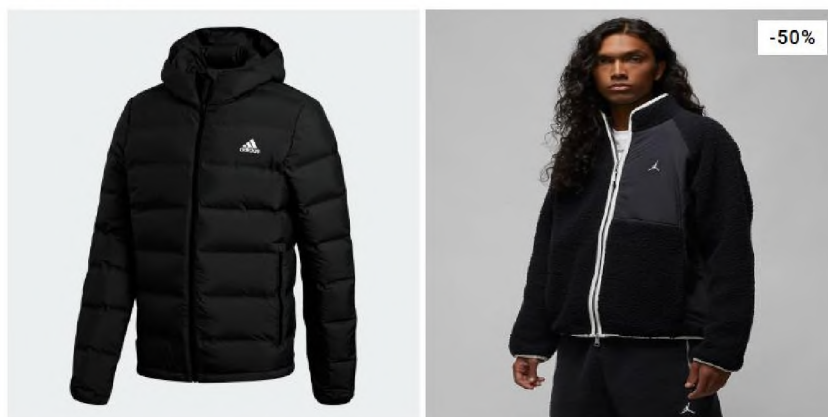
Оскільки перевірка цього об’єкта пройшла успішно, а саме система виявила два товари (див. рис. 3.20), що відповідають зазначеним параметрам

об'єкта, користувача було автоматично перенаправлено на сторінку з відфільтрованими товарами (рис. 3.21). Ця сторінка динамічно формується на основі результатів пошуку та дозволяє миттєво переглянути товари, які задовольняють усі критерії та змінити фільтри (рис. 3.22).

Товари

Категорія ▾	Тип товару ▾	Розмір ▾	Бренд ▾	Стать ▾	Колір ▾	Сортувати ▾	Скидки
-------------	--------------	----------	---------	---------	---------	-------------	--------

[Очистити все](#)



Пуховик з капюшоном Helionix Performance

Jordan Essentials

Рисунок 3.21 – Сторінка з відфільтрованими товарами

Товари

Категорія ▲	Тип товару ▾
Одяг Взуття Акcesуари	

Рисунок 3.22 – Вигляд блоку з фільтром

При цьому блоки з фільтрами автоматично налаштовуються відповідно до введеного запиту користувача, а ключові параметри, які вплинули на відбір товарів, виділяються жирним шрифтом для зручності та акцентування уваги. Користувач також має можливість змінити ці фільтри вручну, адаптуючи

результати пошуку до своїх потреб або вподобань. Такий підхід дозволяє забезпечити інтуїтивно зрозумілий інтерфейс, спрощуючи процес вибору товарів та зменшуючи час на пошук необхідного.

У процесі роботи нейроконсультанта можуть виникати ситуації, коли запит користувача не може бути виконаний через певні причини. Перший приклад – це випадок, коли запит не стосується тематики магазину або є некоректним. Наприклад, якщо користувач вводить запит, що не відповідає категоріям товарів магазину, нейроконсультант визначає невідповідність тематики та відображає повідомлення про помилку із поясненням, що запит не може бути оброблений (рис. 3.23). Це допомагає уникнути некоректних пошукових результатів і спрямовує користувача до правильного формулювання запиту.

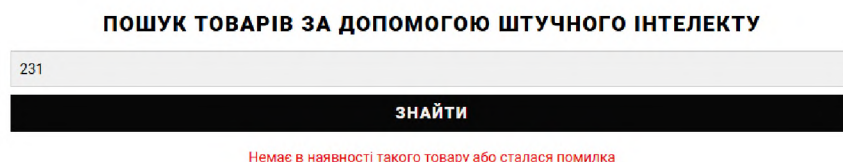


Рисунок 3.23 – Помилка запиту, яка не стосується теми

Другий приклад – ситуація, коли за заданими фільтрами чи ключовими словами не було знайдено жодного товару (рис. 3.24). Наприклад, якщо користувач шукає товар, якого немає в асортименті, система відобразить повідомлення про відсутність товару.

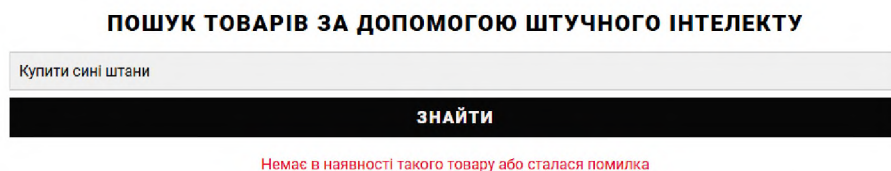


Рисунок 3.24 – Помилка запиту, яка не стосується теми

У подібному випадку користувачу пропонується внести зміни до запиту, наприклад, уточнити ключові слова або вказати більш конкретні параметри

пошуку. Це дозволяє системі надати більш точні та релевантні результати, які відповідатимуть потребам користувача.

Під час тестування було перевірено ще кілька запитів (рис. 3.25-3.26). У більшості випадків нейроконсультант обробляв ці запити коректно, формуючи відповідні об'єкти для фільтрації та перенаправляючи користувача на сторінку з відповідними товарами.

Товари

Категорія ▾

Тип товару ▾

Розмір ▾

Бренд ▾

Стать ▾

Колір ▾

Сортувати ▾

Скидки

[Очистити все](#)



Кросівки OZWEEGO

Кеди Suede Classic XXI Trainers

Nike Air Force 1 PLT.AF.ORM

Рисунок 3.25 – Результат за запитом «Кросівки зі скидкою»

Товари

Категорія ▾

Тип товару ▾

Розмір ▾

Бренд ▾

Стать ▾

Колір ▾

Сортувати ▾

Скидки

[Очистити все](#)



Кросівки OZWEEGO

Кросівки для бігу Runfalcon 2.0 Sportswear FY9496

Штани Dare To Woven Pants Women

Штани 3-Stripes Tapered Performance GE0667

Рисунок 3.26 – Результат за запитом «Купити кросівки та штани»

Однак під час тестування було виявлено, що іноді нейроконсультант може внести невірні дані до об'єкта, наприклад, був запит «Купити куртку», але

нейроконсультант повернув об'єкт із зайвим значенням ключа «gender: women», хоча в запиті не було нічого сказано про стать (рис. 3.27). Це показує, що, хоча OpenAI API є потужним інструментом для NLP, йому ще є куди розвиватися, особливо у сфері точності роботи з багаторівневими та специфічними запитами.

```

▶ brand: []
▶ category: ['clothes']
▶ childAge: []
▶ childGender: []
▶ color: []
▶ gender: ['women']
▶ oldPrice: []
▶ size: []
  sort: ""
▶ type: ['jacket']
▶ [[Prototype]]: Object

```

Рисунок 3.27 – Невірно згенерований об'єкт з фільтрами нейроконсультантом

Подальший розвиток OpenAI API може включати вдосконалення обробки складних запитів, підвищення надійності формування об'єктів для фільтрації та впровадження механізмів самокорекції, які допоможуть уникати помилок у роботі нейроконсультанта. Це дозволить досягти ще більш високої ефективності в електронній комерції та значно покращити взаємодію з користувачами.

3.4 Економічна оцінка ефективності запропонованих рішень

Розробка нейроконсультанта для онлайн-магазину потребує не лише технічної експертизи, але й раціонального підходу до використання фінансових ресурсів. Економічна оцінка запропонованих рішень дозволяє визначити, наскільки ефективно були використані наявні ресурси, а також оцінити доцільність інвестицій у розробку та подальше впровадження системи.

Розробка нейроконсультанта для онлайн-магазину була здійснена із мінімальними витратами. Основною та єдиною витратою стала активація доступу до OpenAI API, яка вимагала внесення \$ 5. Цей платіж забезпечив

можливість інтеграції моделі для реалізації функціональності нейроконсультанта, що дозволило ефективно покращити взаємодію з користувачами та оптимізувати процес обслуговування клієнтів.

Серед доступних моделей було обрано GPT-3.5-turbo. Це рішення обумовлене її доступною ціною та високою ефективністю для розв'язання завдань, пов'язаних із розробкою нейроконсультанта.

За інформацією, наведеною на офіційному сайті OpenAI, GPT-3.5-turbo пропонує найнижчу вартість використання серед сучасних моделей. Наприклад, ціна на генерацію 1 мільйона токенів для цієї моделі становить \$ 1.5 для запитів (input) та \$ 2 для відповідей (output). Це робить її ідеальним вибором для розробки, тестування та впровадження консультаційних рішень, особливо в межах обмеженого бюджету.

Однак у системі доступні й більш сучасні моделі, такі як GPT-4, які забезпечують вищу якість відповідей і кращу продуктивність. Але використання цих моделей значно дорожче. Наприклад, для моделі GPT-4-turbo за 1 мільйон токенів ціна становить \$ 10 для запитів і \$ 30 для відповідей, що може бути економічно невигідним для проєктів із великим обсягом даних або з обмеженим фінансуванням [42].

Таким чином, вибір моделі GPT-3.5-turbo забезпечив оптимальний баланс між продуктивністю та витратами, що є важливим фактором для впровадження нейроконсультанта у практичну діяльність.

Далі представлено гіпотетичний приклад економічного розрахунку витрат на створення нейроконсультанта в комерційних умовах, який може бути орієнтиром для майбутніх розробників.

Для оцінки загальної вартості розробки використовується формула:

$$B = P + \Phi + B + I + T + Z + O, \quad (3.1)$$

де B – сумарні витрати на розробку та обслуговування нейроконсультанта;

P – планування та дизайн;

Φ – frontend розробка;

B – backend розробка;

I – інтеграція API;

T – тестування;

Z – запуск системи;

O – поточне обслуговування.

Приклад ймовірних витрат на створення нейроконсультанта наведено в табл. 3.1, згідно даних з відкритих джерел [43-45].

Таблиця 3.1 – Орієнтовні витрати на розробку нейроконсультанта для онлайн-магазину

№ п/п	Найменування етапу робіт	Опис робіт	Вартість, грн
1	Планування та дизайн	Аналіз вимог, UX/UI дизайн	9000,00
2	Розробка інтерфейсу	Реалізація візуальних компонентів (HTML/CSS/ JavaScript)	12000,00
3	Backend розробка	Створення серверної логіки, налаштування баз даних	14000,00
4	Інтеграція OpenAI API	Генерація ключа, налаштування API	700,00
5	Тестування	Технічна перевірка, виправлення помилок	5000,00
6	Впровадження	Розгортання системи, налаштування хостингу	500,00
7	Поточне обслуговування	Оновлення, технічна підтримка	4000,00
8	Сумарні витрати		45200,00

Таким чином, орієнтовні сумарні витрати на розробку, впровадження та обслуговування нейроконсультанта для онлайн-магазину становлять 45200,00 грн, згідно формули (3.1).

Цей приклад демонструє можливий бюджет для реалізації подібного проєкту в умовах використання комерційних ресурсів. Наведені витрати можуть варіюватися залежно від складності завдань, обраних інструментів та технологій, а також кількості годин, витрачених на окремі етапи роботи.

Отже, завдяки мінімальним реальним витратам у \$ 5, розробка нейроконсультанта в рамках цього проєкту була виконана економічно ефективно. Гіпотетичний розрахунок демонструє, що в умовах комерційної реалізації витрати можуть бути значно більшими, але навіть тоді інтеграція доступної моделі OpenAI API залишається фінансово вигідним рішенням.

Висновки до розділу 3

В результаті було здійснено практичну реалізацію нейроконсультанта для інтернет-магазину із застосуванням OpenAI API. Детально розглянуто архітектуру рішення, ключові етапи інтеграції та тестування розробленого функціоналу. Було проведено порівняння моделей, під час якого встановлено, що одна з них здатна формувати відповіді у форматі JSON, що значно спрощує інтеграцію у процес фільтрації товарів. Інша модель, хоч і є економічно вигіднішою, має обмеження у структурованих відповідях, що ускладнює її використання в розробці, зокрема при обробці складних запитів.

Тестування показало, що нейроконсультант успішно обробляє більшість запитів користувачів, формуючи коректні об'єкти для фільтрації товарів і перенаправляючи користувача на сторінки з результатами пошуку. Проте виявлено, що іноді можуть виникати помилки, пов'язані з некоректним формуванням об'єкта фільтрації, навіть за умов правильного запиту. Це свідчить про наявність напрямів для подальшого вдосконалення та розвитку технології.

Загалом, реалізація нейроконсультанта продемонструвала значний потенціал у покращенні взаємодії користувача з інтернет-магазином, спрощенні процесу пошуку товарів і підвищенні зручності використання платформи. Проте інтеграція OpenAI API все ще має простір для подальшої оптимізації, що відкриває перспективи для майбутніх досліджень та вдосконалення системи.

ВИСНОВКИ

У процесі виконання кваліфікаційної роботи було проведено детальний аналіз сучасного стану ринку електронної комерції, зокрема інтеграції нейромережових технологій в інтернет-магазини, а також розроблено та протестовано функціонал нейроконсультанта для спортивного інтернет-магазину. На основі отриманих результатів можна зробити наступні висновки.

Проведено глибоке дослідження теоретичних основ створення нейроконсультантів та інтеграції AI в електронну комерцію, що дозволило структурувати знання про: роль AI в сучасних технологіях та його вплив на електронну комерцію; принципи роботи нейронних мереж і їх застосування для ефективної взаємодії з користувачами; наявні API та інструменти для розробки нейромережових рішень.

Здійснено вибір технологій та обґрунтовано їх застосування для реалізації нейроконсультанта, зокрема: використання API для обробки запитів користувачів; створення механізмів фільтрації товарів, що базуються на результатах роботи нейромережі; впровадження адаптивного інтерфейсу для покращення користувацького досвіду.

У процесі тестування нейроконсультанта було встановлено, що розроблена система здатна обробляти складні запити та повертати відповідні результати. Однак виявлено потенціал для вдосконалення, зокрема щодо зменшення ймовірності помилок у формуванні об'єктів для фільтрації.

Практична реалізація продемонструвала ефективність автоматизації процесів пошуку товарів, що значно спрощує взаємодію користувачів з платформою та підвищує їх задоволеність.

Порівняння використаних моделей AI засвідчило переваги та недоліки кожної, а також виявило напрями для розвитку технологій OpenAI API.

Економічно проєкт виявився мінімально затратним, оскільки всі основні інструменти розробки використовувалися безкоштовно. Однак для тестування та

роботи з OpenAI API було витрачено \$ 5, що забезпечило доступ до необхідного функціоналу для створення та оптимізації нейроконсультанта.

Практична цінність роботи полягає у створенні інноваційної, гнучкої та адаптивної системи, яка значно оптимізує процеси пошуку товарів у спортивному інтернет-магазині. Нейроконсультант, реалізований на основі сучасних технологій AI, забезпечує автоматизацію складних процесів, що дозволяє значно підвищити зручність користування сайтом для покупців. Завдяки використанню інтелектуальних алгоритмів, система не тільки формує релевантні результати пошуку, але й адаптується до специфічних запитів користувачів, враховуючи різноманітні фактори, такі як знижки, комбінації товарів тощо. Окрім цього, створення подібного інструменту демонструє реальні можливості впровадження нейромереж у сферу електронної комерції, що сприяє розширенню їх застосування для автоматизації процесів та покращення взаємодії клієнтів із платформою. Система дозволяє покупцям заощаджувати час на пошук потрібних товарів і знижує навантаження на адміністративну частину магазину. Однак, враховані недоліки під час тестування, такі як рідкісні помилки у формуванні об'єктів або неправильно відображені результати, вказують на напрямки для подальшого вдосконалення системи. Впровадження рекомендацій та розширення функціоналу в майбутньому дозволять створити ще ефективніші рішення, які відповідатимуть потребам сучасного споживача. Розробка нейроконсультанта підтверджує перспективність інтеграції AI та нейромереж у сферу електронної комерції. Це відкриває нові можливості для підвищення ефективності роботи інтернет-магазинів, зростання задоволеності клієнтів та покращення їхнього досвіду взаємодії з платформою.

Таким чином, результатами роботи є модель нейроконсультанта на основі генеративного штучного інтелекту; рекомендації щодо застосування штучних нейронних мереж для розробки нейроконсультантів. Вони можуть бути використані для подальших досліджень за даною тематикою та при проектуванні хмарних сервісів.