

**Міністерство освіти і науки України
Тернопільський національний педагогічний університет
імені Володимира Гнатюка**

**Факультет іноземних мов
Кафедра англійської філології та методики навчання англійської мови**

Кваліфікаційна робота

**ЛІНГВОСТИЛІСТИЧНІ ТА ПРАГМАТИЧНІ ХАРАКТЕРИСТИКИ
АНГЛОМОВНОГО АІ-ДИСКУРСУ (НА МАТЕРІАЛІ
ПСИХОЛОГІЧНИХ ТЕКСТІВ)**

**Спеціальність:
спеціальність 035 Філологія, спеціалізація 035.041 Германські мови та
літератури (переклад включно), перша - англійська
Освітньо-професійна програма «Англійська мова і література**

Здобувачки другого (магістерського)
рівня вищої освіти групи зМАМ–21
Анциферової Марії Володимирівни

НАУКОВИЙ КЕРІВНИК:
Кандидат філологічних наук, доцент
кафедри англійської філології та
методики навчання англійської мови
Шонь Олена Богданівна

РЕЦЕНЗЕНТ:
кандидат філологічних наук, доцент
кафедри німецької філології та
методики навчання німецької мови
Бияк Наталія Яремівна

ЗМІСТ

ВСТУП	3
РОЗДІЛ 1. ТЕОРЕТИЧНІ ЗАСАДИ ДОСЛІДЖЕННЯ МОВИ ШТУЧНОГО ІНТЕЛЕКТУ	7
1.1. Штучний інтелект як феномен сучасної комунікації	7
1.2. Порівняння англомовних AI-текстів і текстів, написаних людьми: сучасний науковий огляд.....	14
1.3. Прагматичні й стилістичні відмінності ШІ-текстів.....	20
Висновки до Розділу 1	26
 РОЗДІЛ 2. ЛІНГВОСТИЛІСТИЧНІ ХАРАКТЕРИСТИКИ АНГЛОМОВНИХ ТЕКСТІВ, ЗГЕНЕРОВАНИХ ШТУЧНИМ ІНТЕЛЕКТОМ	27
2.1. Вибір і характеристика матеріалу дослідження	27
2.2. Лексико-семантичний профіль AI-дискурсу у зіставленні з автентичними психологічними текстами	30
2.3. Синтаксично-стилістичні параметри AI-дискурсу та їх контрастивний аналіз із текстами, написаними людьми.....	41
Висновки до розділу 2	53
 РОЗДІЛ 3. КОМУНІКАТИВНО-ПРАГМАТИЧНІ СТРАТЕГІЇ ФОРМУВАННЯ ДОВІРИ ТА АВТОРИТЕТУ В AI-ДИСКУРСІ	55
3.1. Реалізація стратегій ввічливості та емпатії	55
3.2. Конструювання експертності та авторитету в AI-дискурсі	59
Висновки до розділу 3	63
 ВИСНОВКИ	66
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	70

ВСТУП

Штучний інтелект (ШІ) дедалі глибше інтегрується у сферу людської комунікації, перетворюючись із технічного інструмента на активного учасника соціального, академічного та міжособистісного обміну. Як феномен сучасного дискурсу, він не лише формує нові жанри тексту, а й змінює уявлення про мовця, адресата, контекст і кооперацію у спілкуванні. Відтак дослідження мовних, стилістичних і прагматичних характеристик текстів, згенерованих штучним інтелектом, є одним із найактуальніших напрямів сучасної лінгвістики.

Сьогодні цифровізація комунікативного простору, автоматизація контенту, поява великих мовних моделей та їхнє використання у сфері освіти, медіа, психології й побутової взаємодії зумовлюють необхідність глибшого розуміння того, як саме AI-тексти відтворюють людське мовлення, якою мірою вони імітують емоційність, ввічливість, експертність, та як користувачі сприймають цю симуляцію.

Наукові дослідження останніх років засвідчують підвищений інтерес до проблеми лінгвістичної ідентичності штучного інтелекту. Зокрема, М. Бендер, Т. Гебру, С. Полада аналізують ризики так званих *stochastic parrots*, тобто моделей, що імітують людську мову без глибокого розуміння змісту; Д. Хові та С. Спрут описують соціальний вплив NLP-технологій; UNESCO (2021) та EU Artificial Intelligence Act (2024) наголошують на етичних принципах, прозорості та довірі у ШІ-комунікації. В українській науці питання розвитку мовних технологій і впровадження ШІ в освітньо-комунікативне середовище висвітлено у «Концепції розвитку штучного інтелекту в Україні» (2020), що визначає комунікаційні аспекти одним із ключових напрямів цифрової політики.

Актуальність теми дослідження зумовлена потребою системно осмислити мову штучного інтелекту як соціолінгвістичний і прагматичний

феномен, розкрити його стилістичні закономірності, а також зрозуміти, як AI-конструкти трансформують уявлення про мовну норму, емпатію та експертність у сучасному цифровому дискурсі.

Метою дослідження є актуалізація лінгвостилістичних, прагматичних та дискурсивних характеристик англомовних текстів, згенерованих штучним інтелектом (ChatGPT 4.0–5.0), у порівнянні з автентичними текстами користувачів онлайн-платформ Reddit і Quora.

Для досягнення цієї мети поставлено такі завдання:

1. Розглянути теоретичні засади вивчення мови штучного інтелекту як феномену сучасної комунікації.
2. Охарактеризувати лексико-семантичні, синтаксичні та стилістичні параметри англомовних AI-текстів у зіставленні з автентичними текстами користувачів.
3. Дослідити прагматичні стратегії ввічливості та емпатії, реалізовані у текстах і діалогах ШІ.
4. Визначити лінгвістичні механізми конструювання експертності та авторитету в AI-дискурсі.
5. Узагальнити дискурсивні відмінності між людським та машинним мовленням і простежити тенденції їх зближення.

Об'єкт дослідження – англомовний дискурс штучного інтелекту.

Предмет дослідження – лінгвостилістичні, прагматичні та дискурсивні особливості AI-текстів у порівнянні з людськими.

Методи дослідження:

- *Критичний дискурс-аналіз* – для виявлення владних і прагматичних стратегій у мовленні ШІ;
- *Порівняльний аналіз* – для зіставлення текстів, створених людиною і машиною;
- *Компонентний і прагмасемантичний аналіз* – для визначення семантичних полів, емоційного забарвлення, частотності;

– *Стилістичний і прагматичний аналіз* – для опису стратегій ввічливості, емпатії та авторитетності;

– *Індуктивний метод узагальнення* – для формулювання висновків щодо закономірностей AI-мовлення.

Матеріал дослідження охоплює два типи текстів: форумні повідомлення користувачів, що репрезентують неформальне, емоційне й спонтанне письмо на прикладі 12 дописів користувачів, та AI-згенеровані відповіді, створені у форматі пояснення чи поради на ті самі запити. Такий поділ дає змогу зіставити природну мовленнєву поведінку користувачів із стандартизованими алгоритмічними патернами ШІ-дискурсу.

Наукова новизна полягає у комплексному аналізі англomовного AI-дискурсу, який поєднує лінгвістичний, стилістичний і прагматичний підходи, а також у визначенні особливостей відтворення емпатії та експертності у штучно згенерованих текстах.

Теоретичне значення роботи полягає у розширенні уявлень про дискурсивну природу мови штучного інтелекту та розробленні критеріїв її аналізу в межах сучасної комунікативної лінгвістики.

Практична цінність полягає у можливості використання результатів для курсів із стилістики англійської мови, політичної лінгвістики, дискурс-аналізу, а також у розробці рекомендацій для етичного та ефективного використання AI-систем у гуманітарній освіті.

Структура роботи. Дослідження складається зі вступу, трьох розділів, висновків, списку використаних джерел та додатків.

У **першому розділі** викладено теоретичні засади дослідження мови штучного інтелекту як нового феномену сучасної комунікації. Проаналізовано підходи до розуміння ШІ як комунікативного агента, окреслено його місце в цифровому дискурсі та особливості взаємодії «людина – машина». Розділ узагальнює лінгвістичні, стилістичні та прагматичні характеристики AI-текстів, представлені у науковій традиції.

Також визначено ключові теоретичні проблеми, пов'язані з авторством, інтенційністю та новими жанрами штучно згенерованого мовлення

У **другому розділі** проведено порівняльний аналіз англomовних AI-текстів та автентичних психологічних текстів користувачів Reddit і Quora. Досліджено лексико-семантичні параметри AI-дискурсу, зокрема абстрактність лексики, характер емоційного забарвлення та показники лексичної різноманітності. Проаналізовано синтаксичні та стилістичні особливості, серед яких — довжина речень, типи конструкцій, повторюваність і формульність. Розділ дозволяє визначити системні лінгвістичні відмінності між людським і машинним мовленням.

У **третьому розділі** розглянуто прагматичні механізми, за допомогою яких штучний інтелект формує довіру, емпатію та експертність у діалозі з користувачем. Проаналізовано стратегії ввічливості та емоційної підтримки, включно з геджингом, емпатичними висловлюваннями, позитивною ввічливістю та активним слуханням. Також досліджено лінгвістичні засоби конструювання авторитетності: спеціалізовану лексику, модальність, структурованість тексту та посилання на джерела. Розділ демонструє, як поєднання ввічливості й експертного тону формує комунікативний «образ» ШІ.

У **додатках** подано скриншоти оригінальних повідомлень з Reddit і Quora, а також відповідні AI-згенеровані відповіді,

Загальний обсяг роботи — 83 сторінки, список джерел налічує 65 найменувань.

РОЗДІЛ 1

ТЕОРЕТИЧНІ ЗАСАДИ ДОСЛІДЖЕННЯ МОВИ ШТУЧНОГО ІНТЕЛЕКТУ

1.1. Штучний інтелект як феномен сучасної комунікації

Упродовж останніх років штучний інтелект (ШІ) дедалі більше входить у сферу людської комунікації – як засіб створення текстів, вирішення інформаційних задач, медіапоширення контенту та взаємодії людини з машиною. Саме тому він перестає бути лише технічним алгоритмом і постає як комунікативний агент або середовище, яке змінює парадигму: не лише що ми комунікуємо, але й як ми це робимо. Проблема, що нас цікавить у цьому підрозділі, – це дослідження ШІ як феномена комунікації: коли алгоритмічні системи беруть на себе роль, пов'язану з мовленням, текстом, дискурсом та соціальним обміном, змінюються традиційні уявлення про авторство, адресата, контекст, кооперацію й контроль. Ці зміни породжують низку теоретичних, методологічних і етичних викликів, серед яких – трансформація ролі мовця [11; 14], нові жанри тексту [27; 48; 52], розмиття меж між людиною та машиною [43; 55], та питання збереження людської присутності в комунікації [7; 9; 10; 17; 21].

Незважаючи на те, що питання штучного інтелекту і комунікацій досліджуються в суміжних галузях – технічних (наприклад, NLP і машинний переклад) [19; 51; 54; 60; 62] та соціальних (медіа-теорії та філософія технологій) [1; 3; 9; 20; 30; 31; 35; 46], – акцент на лінгво-дискурсивній природі ШІ-текстів і їхній ролі у сучасній мові ще формуються. Це створює науковий розрив: недостатньо глибоко досліджено, яким чином ШІ як комунікативний суб'єкт або агент

інтегрується в лінгвістичний простір, які стилістичні та прагматичні зміни відбуваються, і як ці процеси впливають на дискурсивну екологію.

Технологічний розвиток дозволив алгоритмам не просто обробляти мову, а генерувати її: великі мовні моделі (англ. Large Language Models, LLM) створюють повноцінні тексти, відповіді на запити, діалоги, переклади. Внаслідок цього змінюється класична модель комунікації: мовцем може виступати не тільки людина, а й система, адресат-модель може бути не лише людиною, а й користувачем через інтерфейс [19, с. 49]. Як наслідок, з'являються такі комунікативні виклики:

– Авторство і джерело: Якщо текст написаний ШІ, хто автор? Які правові та етичні наслідки?

– Контекст і інтенція: Машина генерує текст, але чи має вона інтенцію? Як це впливає на інтерпретацію?

– Адаптація до аудиторії: Чи враховує ШІ комунікативний контекст, жанр, стиль мовця? І якщо так – як?

– Соціальне сприйняття: Як користувач реагує на текст, створений ШІ – довіряє/не довіряє, сприймає як «машинний» чи створений людиною?

У цьому контексті ключовим є дослідження комунікаційної ролі ШІ: чи він лише інструмент, чи вже повноцінний партнер у комунікації. Деякі дослідники, наприклад, М. Бендер та ін. стверджують, що великі мовні моделі – «мовні папути» (stochastic parrots) – імітують форму, але не обов'язково розуміння, що створює специфічний тип комунікації [26, с. 613-614].

Штучний інтелект стає ключовим гравцем у різних сферах комунікації. Системи ШІ дедалі активніше використовуються в медіа, освіті, бізнес-комунікації, перекладі та навіть у чат-ботах. Вони вже не тільки опрацьовують, але й самостійно генерують контент, що раніше був виключно продуктом людської творчості. Це знаменує формування нового

комунікативного ландшафту, де технологія інтегрується в процес створення і поширення інформації.

У сучасних медіа алгоритми рекомендують матеріали, формують новинні стрічки та просувають певний контент, прямо впливаючи на те, яка інформація стає видимою для аудиторії [37]. Подібна автоматизована модерація і персоналізація інформаційних потоків визначає, що споживачі бачать найперше, і тим самим опосередковано формує громадську думку. Алгоритмічні системи таким чином беруть на себе частину функцій редакторів та кураторів контенту, впливаючи на порядок денний у суспільстві.

У бізнес-комунікації та освітніх застосунках ШІ дедалі частіше виступає співрозмовником або віртуальним асистентом. Зокрема, чат-боти на основі ШІ ведуть діалог із користувачами, відповідаючи на запити, надаючи поради чи навчальні матеріали. Така взаємодія людини з машинним агентом ставить нові питання щодо ролей учасників комунікації: хто є фактичним відправником повідомлення, як визначити цільову аудиторію і стиль спілкування, якщо один із співрозмовників – алгоритм? Виникає потреба переосмислити принципи кооперації в діалозі, коли замість людини діє автоматизована система [25; 35].

З погляду суспільства, контент, створений штучним інтелектом, став масовим явищем. Вже існують численні автоматично написані статті, згенеровані коментарі у соціальних мережах та машинні переклади, які важко відрізнити від матеріалів, підготовлених людиною [20, с.49]. Така практика підриває усталені уявлення про текст, створений людиною і авторство: читачі та слухачі тепер мають замислюватися, чи походить інформація від людини, чи згенерована алгоритмом. Це змушує переглянути критерії автентичності та довіри до інформаційного контенту в цифрову епоху.

Технологія розмиває межі між технічним і соціальним у комунікації. У результаті впровадження ШІ в комунікативні процеси межа між технологічним інструментом і соціальною взаємодією стає дедалі менш чіткою. Штучний інтелект виступає одночасно і середовищем, і активним агентом комунікації, що ставить у центр уваги питання про вплив технології на структуру, форму та функції мови і дискурсу. З'являється потреба зрозуміти, як алгоритми змінюють мовні практики, правила дискусії та норми спілкування, які склалися століттями на підставі виключно людського досвіду.

На теоретичному рівні поява ШІ як учасника комунікації викликає ряд фундаментальних питань. В класичних моделях комунікації (Шеннон–Вівер, Ласвелл) відправником повідомлення завжди виступає людина. З появою штучного інтелекту цю модель потрібно адаптувати: чи може система ШІ розглядатися як автономний суб'єкт у комунікативному процесі, чи вона лишається інструментом або каналом передачі, що повністю контролюється людиною. Вирішення цього питання визначає наше розуміння відповідальності та авторства в машинно-людській взаємодії [42; 43].

Окремої уваги потребує проблема інтенціональності та дотримання комунікативних норм у мовленні [13], згенерованому ШІ. Людське спілкування керується принципом кооперації та максими Грайса [40] (максими кількості, якості, релевантності та манери), які забезпечують зрозумілість і ефективність обміну інформацією. Натомість залишається неясним, чи штучний інтелект здатний послідовно дотримуватися цих максим або взагалі враховувати комунікативні наміри співрозмовника. Якщо алгоритми систематично порушують усталені людські правила спілкування, постає питання про існування специфічної «машинної прагматики» – набору принципів і норм, за якими відбувається генерування та інтерпретація текстів ШІ незалежно від людських інтенцій.

Поширення чат-ботів та інших ШІ-генераторів тексту ставить питання про жанр і стиль отриманих повідомлень. Фактично формується новий комунікативний жанр – діалог людини з машиною, який не вписується повністю в традиційні жанрові рамки. Наприклад, розмова з чат-ботом може поєднувати риси офіційного запиту, дружньої бесіди та інформаційного довідника одночасно. Стилiстично такі тексти можуть суттєво відрізнятися від людських: вони часто стандартизовані, позбавлені індивідуальних рис автора, а подекуди можуть містити нестандартні формулювання чи помилки, характерні для алгоритмічних моделей. Визначення жанрового статусу і стилістичного регістру для ШІ-текстів є новим завданням для лінгвістики та теорії комунікації.

Загалом, у суспільно-дискурсивному вимірі поява ШІ загострює питання про владу над інформаційними потоками. Алгоритмічні механізми вже сьогодні визначають, які теми та повідомлення отримують широкий розголос, а які залишаються на маргінесі, фактично моделюючи публічний дискурс. Фільтрування, ранжування та модерація контенту за допомогою ШІ створюють дискурсивну владу – здатність впливати на колективні уявлення та порядок денний, не привертаючи до себе особливої уваги. Як зазначають дослідники Тайна Бучер [29] і Тарлтон Гіллеспі [39], алгоритми через свій вплив на видимість інформації стають новими “gatekeeper’ами” цифрової епохи, від яких залежить, що саме побачить або не побачить масова аудиторія. Цю проблему підкреслюють і міжнародні організації – наприклад, у звітах ЮНЕСКО [22] звертається увага на те, що контроль над алгоритмічними системами фактично означає контроль над знанням та інформацією в суспільстві.

Коли машина виступає співрозмовником, виникає питання про її соціальну присутність та формування її ідентичності. У людському спілкуванні важливі особистість мовця, стиль, емоційне забарвлення і рівень емпатії [Goffman, с. 857]. У випадку з чат-ботами люди часто

проектують ці якості на безособового алгоритмічного агента. Відсутність справжнього «Я» у машини ускладнює сприйняття: користувач може сумніватися, чи можна довіряти відповіді, і чи здатен ШІ зрозуміти контекст і тонкощі ситуації. Постає й етичне питання – чи варто наділяти ШІ ознаками особистості (іменем, аватаром, стилем спілкування), щоб зробити взаємодію більш комфортною, і яких наслідків це може мати для сприйняття технології в суспільстві.

Тому, вивчення ШІ як комунікативного феномену потребує міждисциплінарного підходу. Лінгвістиці та прагматиці слід об'єднати зусилля з фахівцями з дискурс-аналізу, медіадослідниками та етиками для всебічного аналізу впливу ШІ на мову та комунікацію. Лише залучивши знання з різних галузей, можна адекватно оцінити, як технології штучного інтелекту трансформують сучасне спілкування – від структур мовлення і жанрів до владних відносин у дискурсі – та виробити стратегії ефективної і відповідальної інтеграції ШІ у соціальну комунікацію.

У академічній літературі почали з'являтися праці, що досліджують аспекти комунікації ШІ, хоча часто під іншими кутами. Так, М. Бендер в роботі “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” аналізують великі мовні моделі у контексті ризиків – повторюваності, упереджень, нестачі розуміння – що має комунікаційні наслідки: якщо ШІ генерує тексти, чи вони якісно відрізняються від людських? [26, с. 614]. Д. Хові та С. Спрут у статті “The Social Impact of Natural Language Processing” досліджують, як мовні технології, включно з автоматизованою генерацією тексту, змінюють соціальний контекст мови – наприклад, масові коментарі, автоматичний переклад – і ставлять питання про дискурсивні наслідки [44, с. 592].

На регуляторному рівні UNESCO ухвалила в 2021 р. «Recommendation on the Ethics of Artificial Intelligence», документ, у якому закладено принципи, що стосуються мови, комунікації, прозорості систем:

«Recommendation ... is applicable to all AI systems» - що прямо пов'язано з комунікаційними аспектами: видимість, інтерфейс, вплив на суспільство [23].

Також законодавча ініціатива Artificial Intelligence Act (ЄС) 2024 р. створює рамки для регулювання ШІ-систем, зокрема тих, що генерують контент або взаємодіють з людьми, що також має комунікаційний вимір [37].

У національному українському контексті: Концепція розвитку ШІ в Україні (розпорядження № 1556-р від 02.12.2020) визначає ШІ як один з пріоритетних напрямів, що охоплює й комунікаційні технології [6].

Попри наявність таких досліджень, рівень аналізу саме комунікаційної функції ШІ, його ролі як мовця, стилістичних та прагматичних трансформацій та дискурсивних наслідків – є недостатнім. Тобто, відсутня системна лінгвістично-прагматична карта феномена «ШІ → комунікація». Це створює простір для дослідження, яке поєднує лінгвістичний, прагматичний і дискурсивний аспекти, що відповідає загальному плану магістерської роботи.

Отже, до ключових проблем, які потребують висвітлення, належать:

1. Роль мовця і агентності в комунікації: Коли система генерує текст, чи вона лише інструмент, чи комунікативний агент? Як це впливає на спосіб інтерпретації користувачами і на комунікативні очікування?
2. Якість і довіра: Які параметри визначають «хороший» текст ШІ з комунікаційної точки зору (лексика, стиль, адекватність, адресат-орієнтованість)? Як користувачі сприймають тексти, які створені ШІ?
3. Прагматична кооперація: Чи дотримуються машини максими Грайса, чи існує специфічна прагматика «машинного тексту»? Як відрізняється комунікативна взаємодія людини з машиною від людини з людиною?

4. Дискурсивна екологія: Як поява ІІІ-текстів впливає на дискурсивні жанри (наприклад, блог, коментар, стаття), на публічну комунікацію, на видимість тем, які висвітлюються, на роль алгоритмів у формуванні змісту? Які нові жанри/механізми комунікації виникають?

5. Етичний і соціальний вимір: У випадку комунікації через ІІІ – як забезпечити відповідність принципам прозорості, пояснюваності, контролю людини, недискримінації? (див. рекомендації ЮНЕСКО [23]).

Таким чином, ІІІ як феномен сучасної комунікації відкриває комплексну наукову проблему: не лише технологічну, але й лінгвістичну, прагматичну, дискурсивну та соціально-етичну. Вона передбачає перепони, які варто долати: адаптація теорій комунікації до нових агентів, методологія аналізу алгоритмічного тексту, соціальні наслідки алгоритмічної генерації мови. Дослідження цієї області є не лише актуальним, але й необхідним для розуміння того, як змінюється наше мовлення, наше спілкування і наш дискурс у епоху штучного інтелекту.

1.2. Порівняння англомовних AI-текстів і текстів, написаних людьми: сучасний науковий огляд

З появою новітніх великих мовних моделей на кшталт GPT-3, ChatGPT, Claude, LLaMA тощо, активізувалися дослідження, що порівнюють тексти, згенеровані штучним інтелектом, з текстами, написаними людьми. Починаючи з 2019 року і особливо після 2020, науковці аналізують різні лінгвістичні аспекти цих текстів – від лексики та синтаксису до стилю, прагматики й жанрової структури. Нижче розглянемо основні аспекти: лексика, синтаксис, когерентність, стилістика, прагматика і жанрові особливості.

Багато робіт приділяють увагу тому, наскільки багатим чи різноманітним є словниковий запас AI-текстів порівняно з людськими.

Педро Рев'єріго та співавтори дослідили англомовні відповіді на питання, згенеровані ChatGPT, порівнявши їх з відповідями людей, і дійшли висновку, що ChatGPT (версії GPT-3.5) схильний уживати меншу кількість унікальних слів і демонструє нижче лексичне багатство, ніж люди [61]. Ці попередні результати вказують, що тексти GPT-3.5 можуть бути лексично одноманітнішими. Водночас новіші моделі поліпшують показники: зокрема, ChatGPT-4 (GPT-4) у деяких завданнях вже досягає лексичного різноманіття, співставного або навіть вищого за людський [42]. Так, у великому експерименті зі студентськими есе (аргументативними текстами), ChatGPT-3.5 мав значно нижчу різноманітність словника порівняно з учнівськими творами, натомість ChatGPT-4 продемонстрував навіть більшу лексичну різноманітність, ніж середній студент [42]. Це свідчить про суттєвий прогрес: моделі нового покоління здатні ширше варіювати лексику.

Крім простого підрахунку унікальних слів, дослідники аналізують й інші показники лексичного багатства: типо-частотні коефіцієнти, індекси Maas, MATTR, MTLT тощо [64]. Загальна тенденція – AI-тексти часто містять частіші, загальніші слова, тоді як люди можуть уживати рідкісніші або контекстно специфічні терміни. Наприклад, результати К. Кендро та ін. (2025) показали, що тексти, створені кількома версіями ChatGPT (включно з GPT-4), істотно відрізняються за всіма вимірами лексичного різноманіття від людських текстів [46]. Автори виміряли шість різних показників (обсяг словника, багатство, рівномірність використання тощо) і отримали статистично значущі відмінності: новітні моделі (зокрема експериментальний ChatGPT-4.5) хоча й генерують менше слів, але демонструють вищі деякі показники різноманітності, однак структура цієї різноманітності не збігається з притаманною людям [46]. Це означає, що навіть коли модель використовує багатий словник, розподіл словоформ та повторюваність у неї інші, ніж у людини. Дійсно, ряд робіт дійшли

висновку, що ВММ-моделі *не цілком відтворюють* стиль варіювання лексики людиною, і в деяких аспектах новіші моделі навіть більше відхиляються від людської норми, ніж старіші (імовірно через надмірну оптимізацію під певні тренувальні цілі).

Хоча ВММ здатні добирати синоніми, аби уникати повторів, їхній вибір слів інколи видає машинний характер. Одна з ознак появи AI-стилю – часті вживання певних “улюблених” нейтрально-наукових слів. Наприклад, аналіз понад 15 мільйонів наукових анотацій (Кобак та ін., 2025) виявив стрибок у частоті окремих стильових слів з появою ВММ – зросло використання характерних для моделей фраз, таких як “delve” («зануритися в тему») тощо [64]. Це може свідчити, що мова ШІ має тенденцію до стандартних формулювань і уніфікації словника. Водночас явного «словарного збідніння» наразі не зафіксовано: дослідники П. Рев’єріго та ін. відзначають, що вплив AI-контенту на еволюцію мови поки неясний, але існує ризик, що слова, не вживані моделями, поступово виходитимуть з ужитку [61]. Таким чином, лексичні показники (різноманітність, багатство, характерний словник) є одними з найважливіших категорій для порівняння, і їх відносно легко виміряти автоматично.

Серед синтаксичних характеристик перш за все варто приділити увагу довжині і складності речень, підрядності, індивідуальному стилю та ін. Аналіз синтаксису показує, що тексти, породжені ВММ, часто відрізняються за структурою речень. За даними великого порівняльного дослідження (Herbold та ін., 2023), ШІ-есе мають складніший синтаксис у середньому: виявлено, що GPT-моделі використовують більш довгі речення і складні конструкції, ніж тексти учнів [42]. Зокрема, ChatGPT має схильність до розгорнутих речень із високою підрядністю та до номіналізації (перетворення дієслів/прикметників на іменники) [42]. Це надає викладу більш формального, «наукового» характеру.

Інше масштабне дослідження стилю ВММ підтверджує ці спостереження кількісно: інструкційно-натреновані моделі (такі як ChatGPT) вживають дієприкметникові звороти у 2–5 разів частіше за людей та в 1.5–2 рази частіше використовують номіналізовані конструкції. Наприклад, у текстах GPT часто трапляються фрази на кшталт “...while also promoting sustainable consumption patterns” з багатьма віддієслівними іменниками [60]. Люди в аналогічних ситуаціях частіше б обрали простіші підрядні речення або дієслівні форми. Отже, синтаксис AI-текстів загалом більш складноорганізований і «іменниково-центричний», що створює враження високої інформативної щільності тексту.

Цікаво, що деякі граматичні особливості в AI-текстах проявляються в інших пропорціях, ніж у людських. А. Рейнарт та ін. виявили, що GPT-4 (і подібні інструктивні моделі) значно частіше використовує дієприкметникові звороти та віддієслівні іменники, але рідше вдається до пасивного стану [60]. У їх експерименті GPT-4 застосовував пасив удвічі рідше, ніж це властиво корпусу людських текстів. Це може пояснюватися тим, що моделі, натреновані відповідати на запити, схильні формулювати думку прямо (активними конструкціями) для ясності. Натомість у людському письмі, особливо науковому чи публіцистичному, пасив є звичним для обезособлення тверджень. Ще одна відмінність – частота модальних дієслів і умовних конструкцій: у студентських есе люди частіше використовують модальні дієслова (типу *may*, *might*, *should*), виражаючи сумнів чи можливість, тоді як GPT-тексти більш категоричні, з меншою кількістю таких “епістемічних маркерів” [60]. Дослідження зафіксувало, що текст ChatGPT-4 містить менше займенникових суб’єктів (відсутнє “я/ми”) і менше конструкцій, що виражають позицію автора, порівняно з людським текстом [60]. Це робить синтаксис ШІ більш нейтральним, безособовим.

Загалом ВММ демонструють консистентну схильність до певних граматичних структур. Попри вражаючу граматичну правильність, мовні

моделі не відтворюють усього розмаїття синтаксичних стилів, властивого різним авторам. Дослідники відзначають, що збільшення розміру моделі не усуває цих відмінностей, а інструктивне навчання (на зразок ChatGPT) навіть посилює специфічний стиль, спільний для текстів ШІ [56; 60]. Цей стиль можна охарактеризувати як “інформаційно насичений”: довгі складні речення з багатьма підрядними частинами, віддієслівними іменниками, чіткою логічною структурою. У різних жанрах це проявляється по-своєму, але загальна риса – уніфікація синтаксису у AI-виході. Для практичного аналізу синтаксичні показники (довжина речення, глибина синтаксичного дерева, число підрядних, частка пасиву тощо) є зручними метриками, оскільки їх можна кількісно оцінити за допомогою NLP-інструментів і порівняти між корпусами AI й людських текстів.

Багато досліджень перевіряють, наскільки зв’язно й логічно моделі пишуть тексти порівняно з людьми. У сфері навчальних есе отримані доволі несподівані результати: у великому масштабному експерименті аргументативні есе, згенеровані ChatGPT, були оцінені вчителями вище за якістю та зв’язністю, ніж есе реальних школярів. Зокрема, ChatGPT-4 перевершив GPT-3.5 та середньостатистичних учнів за такими критеріями, як логічна структура, зв’язність абзаців та лінгвістична зв’язність тексту. Це означає, що при написанні творів на задану тему модель дотримується чіткої структури (вступ, аргументи, висновок), використовує логічні зв’язки («по-перше», «отже», «на додаток» тощо) навіть ретельніше, ніж це роблять багато учнів. Зв’язність між реченнями («text linking») у GPT-4 також виявилася на високому рівні [42], тобто AI вміє явно показувати, як одне речення впливає з іншого. Таким чином, у контрольованих завданнях (наприклад, есе за навчальною темою) ШІ-текст може бути дуже когерентним і структурованим.

Водночас інші роботи звертають увагу на певні когерентні вади AI-текстів, які менш характерні для людини. Зокрема, при спробах розпізнати

текст ШІ експерти часто помічають у ньому надмірну повторюваність формулювань та ідей. У дослідженні Б. Дору, де викладачі намагалися відрізнити есе медичних студентів від есе ChatGPT, найчастіше згадуваними ознаками “авторства ChatGPT” були саме «повторювання» і «нудка циклічність» викладу. Близько 91% GPT-текстів містили значні повтори чи дублювання думок, тоді як усього 9% студентських текстів мали таку ваду [36]. Людські автори (принаймні студенти) зазвичай уникають такого монотонного повторення, більше просуваючись по суті аргументу.

Окрім логічної структури, досліджують і когезію – мовні засоби зв’язку: займенники, союзи, повтор мотивів. Загалом GPT-тексти дуже когезивні на поверхневому рівні. Вони щедро використовують зв’язки типу *however, therefore, in conclusion*, щоб явно позначити відношення між частинами тексту [42]. Д. Сковрон і А. Бачовська провели кількісний аналіз академічного тексту та його перефразувань, згенерованих різними моделями GPT-3, і дійшли висновку, що між ШІ-версіями та оригіналом існує дуже високий ступінь подібності щодо когезії [62]. Моделі досить точно зберегли референційні зв’язки, повторювані лексеми для єдності теми, використання займенників тощо. Це демонструє, що при умові наявності якісного зразка (людського тексту для парафразування) GPT здатний підтримувати когезію не гірше автора-людини. З іншого боку, у вільно згенерованих текстах (не редагованих людиною) інколи зустрічаються логічні стрибки або зайва деталізація несуттєвих моментів, що може порушувати глобальну когерентність. Отже, для практичного аналізу варто враховувати різні аспекти зв’язності: від структури тексту (наявність чіткої аргументації, плану) до локальної когезії (зв’язки між реченнями). Їх можна вивчати як вручну (оцінюючи логіку викладу), так і автоматично (через метрики когезії, виявлення повторів тощо).

Таким чином, порівняльний аналіз англomовних AI-текстів і текстів, написаних людьми, засвідчує як суттєві відмінності, так і деякі подібності

в їх лінгвістичній структурі. З одного боку, моделі типу GPT-3.5 зазвичай демонструють нижчий рівень лексичної різноманітності та схильність до вживання уніфікованих стилістичних конструкцій, тоді як GPT-4 наближається до людських показників або перевершує їх за окремими метриками. Синтаксис AI-текстів часто виявляється складнішим і більш формалізованим, із домінуванням підрядності та номіналізацій, що створює ефект інформаційної щільності. Водночас структура таких текстів здебільшого когерентна та чітко організована, хоч і з ризиком надмірної повторюваності або «механічної» зв'язності. Усе це дозволяє розглядати AI-тексти як особливий тип дискурсу, що потребує окремих критеріїв аналізу в межах лінгвістичних і прагматичних досліджень.

1.3. Прагматичні й стилістичні відмінності ШІ-текстів

У міру того як великі мовні моделі дедалі частіше використовуються для створення текстів різного жанру, постає питання: чи мають ці тексти притаманні людині прагматичні та стилістичні риси. Відповідь на нього стає дедалі важливішою як для лінгвістичної теорії, так і для практики – зокрема, у сфері освіти, медіа, детекції AI-контенту. Цей розділ зосереджено на аналізі наукових праць, що досліджують ключові відмінності між ШІ- та людськими текстами за стилем, емоційністю, ввічливістю, мовленнєвими актами та жанровими особливостями.

Однією з найбільш впізнаваних рис текстів, написаних ШІ, є їх стилістична нейтральність і формальність. Численні дослідження відзначають, що GPT-моделі імітують безособовий академічний стиль письма. Наприклад, у порівнянні студентських есе й есе ChatGPT було виявлено, що AI-тексти звучать більш відсторонено та безособово. Вони відповідають усім формальним конвенціям (структура абзаців, академічна лексика), але бракує особистого “голосу” автора, який притаманний

роботам студентів. У той час як студенти часто звертаються до читача, висловлюють власну думку, додають емоційні або оцінні вислови, – ChatGPT переважно подає інформацію рівним тоном, уникаючи суб'єктивності. Звідси випливає й відсутність чіткої позиції: у есе від GPT часто немає відчуття, що “автор” палко відстоює певний погляд – текст здається збалансованим і відполірованим, але без живої переконливості. AI-есе хоч і грамотно побудовані, “не мають особистого штриху чи виразної перспективи” [62, с.141]. Це відображає саму природу тренувальних даних моделей, які прагнуть максимального узагальнення і коректності, але не передають індивідуальну манеру.

Моделі за замовчанням дотримуються нейтрально-ввічливого тону. Вони уникають сленгу, розмовних виразів, різкої критики чи фамільярності – того, що може прослизнути в тексті людини в залежності від контексту. У діалогах з користувачами (чат-відповіді) ChatGPT майже завжди ввічливий, чемно звертається («дякую за ваше запитання», «сподіваюся, це допоможе»). Людина ж, особливо в неформальному чаті, не завжди дотримується такої бездоганної чемності. Деякі дослідження прагматики звертали увагу на те, що мовна ввічливість ШІ є стандартною і мало реагує на тон співрозмовника [57]. Наприклад, ChatGPT зазвичай не змінює стилю відповіді навіть на грубі чи фамільярні запити, тоді як людина б адаптувалася (можливо, відповіла менш ввічливо або більш емоційно). З точки зору мовленнєвих актів, GPT-моделі найчастіше реалізують інформативні та пояснювальні акти – вони покликані надавати інформацію, пояснювати, інструктувати. Виразні мовленнєві акти (як-от пожартувати, поскаржитися, висловити сарказм) з'являються тільки за явної інструкції. Тому AI-текст часто має однорідну прагматичну спрямованість: повідомити нейтрально правильні факти або рекомендації.

У стилістичному сенсі людям властиво робити текст більш «живим» – вставляти запитання до читача, вигук, метафори чи жарти. У науковому

експерименті Дж. Кін аналізували так звані маркери залучення (engagement markers) в учнівських та AI-есе [56]. Виявилось, що студентські роботи були насичені прийомами залучення: риторичними запитаннями, звертаннями до читача, вставними коментарями від себе. Це робило тексти більш переконливими та динамічними. Натомість есе ChatGPT майже не містили питальних речень чи прямих звернень, вони обмежували особисті коментарі і трималися відчужено-об'єктивного стилю. Модель ніби імітує нейтрального оповідача, який намагається нікого не образити і просто викласти матеріал. Звісно, у деяких жанрах (наукова стаття, офіційний лист) це є нормою, але в інших (есе-роздум, блог, огляд) такий стиль видається заспокійливо одноманітним. Індивідуальний стиль автора – манера, за якою можна впізнати конкретну людину – поки що погано відтворюється. III. СтилOMETричні аналізи підтверджують: навіть найпотужніші моделі «мають відбиток стилю», який різниться від людського і більш схожий між собою для різних ВММ, ніж між різними людьми [60]. Отже, категорії стилістичного аналізу – формальність, особливості тональності, наявність авторських «фішок» – є дуже показовими при відмінності AI проти людини. Особливо легко помітити такі речі, як наявність/відсутність звертань до читача, частка особових займенників, емоційна лексика тощо, що робить стилістику зручною сферою для практичного дослідження.

Як вже зазначалося, III часто дотримується однієї й тієї ж ввічливо-нейтральної манери незалежно від ситуації. У письмових відповідях він рідко використовує особові займенники другого лиця (you) або неформальні звертання. Тексти, написані людьми, залежно від жанру, можуть сильно варіювати: лист може починатися з «Дорогий Джоне», блог – звертатися на “ти” до аудиторії, а наукова стаття узагалі уникає прямих звертань. ChatGPT за замовчуванням не має власної персони, тому уникає “я” і “ти”, хіба що у ввічливих кліше типу “*I hope this helps*” в кінці

пояснення. Це означає, що в діалогових жанрах (наприклад, чат-відповідь на запитання користувача) модель тримається ролі чемного консультанта, тоді як людина-експерт могла б дозволити собі більш розмовний тон чи навіть гумор. Практично це проявляється і в мовленнєвих актах: AI-текст здебільшого декларує інформацію або дає інструкції у ввічливій формі (“*One should consider...*”, “*It is recommended to...*”), тоді як людина може виразити пораду як м’якше (“*Maybe you could try...*”) або жорсткіше (“*Do this now!*”) залежно від характеру.

Дослідники людино-машинної прагматики відзначають, що ВММ не завжди розуміють та відтворюють підтекст і імплікатур так, як люди. У Тюрінг-подібних тестах на ведення розмови моделі часто порушують прагматичні принципи – дають надто прямі відповіді там, де людина би відповіла ухильно або з гумором. Наприклад, на провокаційно неввічливе питання людина може зреагувати сарказмом або зміною теми (неявним актом), тоді як ChatGPT, радше за все, відповість все так само ввічливо і по суті запиту, ігноруючи імполітність співрозмовника [57].

У письмових текстах (есе, статті) ці відмінності менш помітні, але теж можуть проявлятися: людина тонко враховує аудиторію, може додавати відступи на пояснення жаргону для неспідготовленого читача або навпаки, апелювати до спільного досвіду з читачем. III-ж, не маючи справжнього соціального досвіду, дотримується універсальної ввічливо-нейтральної стратегії, що інколи знижує ефективність переконання чи впливу. Так, у дослідженні студентських есе зазначалося, що AI-тексти були менш переконливі, саме через відсутність прямого апелювання та чіткої позиції автора [3]. З прагматичного погляду, модель не виконує акт переконання “від себе”, а лише безособово викладає аргументи. Для студента-аналітика цікавими для порівняння можуть бути такі прагматичні показники, як кількість питальних речень (спроб діалогу з читачем), наявність ввічливих

маркерів (“please”, “thank you”) або ж типи мовленнєвих актів (констатація фактів vs. заклик до дії, питання, прохання тощо).

У науковій літературі з 2019 р. порівнювалися різні жанри, щоб зрозуміти, де відмінності між AI та людиною найпомітніші. Одним з найпопулярніших жанрів є академічне або учнівське есе. Це зумовлено практичними потребами освіти: перевірити, чи може ChatGPT писати за студентів, і як відрізнити такі роботи. Такі дослідження уже засвідчили як сильні сторони AI (структурованість, грамотність), так і слабкі (відсутність особистої залученості, оригінальності) на матеріалі аргументативних есе. Інший часто досліджуваний жанр – журналістський текст. Наприклад, в одному із досліджень за створеним корпусом новин *The New York Times*-стилю, згенерували аналогічні тексти шістьма різними моделями і порівняли їх з оригінальними новинами [60]. Виявлено різницю у виборі слів, вживанні розділових знаків, середній довжині речень тощо. Професійні новинні статті людей зазвичай містять різноманітні стилістичні прийоми (цитати, риторичні запитання в заголовках, гра словами), тоді як AI-вихід був більш одноманітним. Однак цікаво, що за лексичним різноманіттям GPT-4 досяг рівня професійних журналістів у новинах, і навіть перевищив його за деякими показниками [39]. Це свідчить, що в інформаційному жанрі модель може конкурувати з людиною по багатству мови, хоч і писати трохи шаблонніше.

Крім есе і новин, академічні статті й огляди теж стали об’єктом аналізу. Д. Сковрон та А. Бачовська досліджували фрагмент наукової статті з гуманітарних наук та його перефразування різними моделями GPT-3. Їх результати показали дуже високий рівень відповідності між оригіналом і AI-версіями за лексичними та когезійними параметрами [62] – настільки, що жодна з моделей не була однозначно кращою за всіма показниками, всі більш-менш точно імітували стиль наукового викладу. Це говорить про те, що для академічного викладу сучасні моделі вже стали вправними

інструментами: вони дотримуються термінології, структури абзаців, формального тону на рівні людини. Водночас інше дослідження [56] порівняло оглядові статті, написані вченими, з тими, що згенерував GPT-4. За критеріями когерентності та плавності викладу GPT-4 отримав оцінки ~90/100, тобто навіть трохи перевершив середньостатистичний текст, створений людиною. Проте це стосувалося саме мовної зрозумілості; у сфері змістовності та оригінальності, звісно, ШІ не може справді “творити” наукові ідеї, а лише компілює відоме.

Ще одним напрямом є чатові відповіді й діалоги. Хоча менше формальних досліджень порівнюють цілі діалоги, зрозуміло, що у форматі живого спілкування відмінності виявляються в динаміці: люди можуть відхилятися від теми, жартувати, неправильно розуміти питання, тоді як ChatGPT прагне чітко слідувати інструкції і зберігати ввічливий тон. Наприклад, одне дослідження прагматики задавало моделі і людям однакові ввічливі та неввічливі питання – ChatGPT у відповіді завжди тримав формальний стиль, тоді як люди в разі грубого питання могли змінити тон або не відповісти прямо [57]. Отже, у жанрі інтерактивного діалогу ШІ-текст відрізняється передбачуваною ввічливістю і послідовністю, що не завжди характерно для людини.

Отже, попри стрімкий прогрес мовних моделей, ШІ-тексти залишаються стилістично і прагматично відмінними від людських. Їм властива формалізована нейтральність, відсутність індивідуального стилю, слабе реагування на контекст і обмежений спектр мовленнєвих актів. Людські тексти, навпаки, гнучкіші в тоні, більш залучені емоційно й адресно. Визначені параметри – стильова виразність, ввічливість, звертання до читача, жанрова відповідність – дають змогу не лише краще розуміти особливості ШІ-дискурсу, а й практично застосовувати ці критерії у порівняльному аналізі та детекції AI-контенту.

Висновок до Розділу 1

Сучасні дослідження одностайні в тому, що попри значні успіхи, тексти, написані великими мовними моделями, ще не повністю ідентичні текстам людей. Вони відрізняються низкою лінгвістичних параметрів – від вибору слів до глибинної стилістики і прагматики. Ці відмінності, з одного боку, дозволяють розпізнавати AI-тексти, а з іншого – ставлять питання про вплив масового використання генеративного ШІ на саму письмову мову. Для освітян і дослідників мови важливо розуміти ці нюанси. Категорії, розглянуті в огляді – лексика, синтаксис, когерентність, стилістика, прагматика і жанр – слугують зручними “лінзами” для аналізу. Використовуючи їх, можна систематично порівнювати тексти та робити висновки про те, *де* саме машинний текст відхиляється від людського і як це може бути враховано на практиці (наприклад, у навчанні письма чи розробці систем детекції AI-контенту). Станом на 2025 р., моделі на кшталт GPT-4 продовжують удосконалюватися, і дослідження показують тенденцію до зближення з людськими показниками за певними критеріями. Однак повна стилістично-прагматична «емуляція» людини поки не досягнута – і саме на цьому полігоні триватимуть подальші дослідження.

РОЗДІЛ 2

ЛІНГВОСТИЛІСТИЧНІ ХАРАКТЕРИСТИКИ АНГЛОМОВНИХ ТЕКСТІВ, ЗГЕНЕРОВАНИХ ШТУЧНИМ ІНТЕЛЕКТОМ

2.1. Вибір і характеристика матеріалу дослідження

У цьому дослідженні аналізуються тексти, згенеровані штучним інтелектом (на прикладі ChatGPT 4.0 і 5.0), та порівнюються з автентичними англomовними текстами користувачів платформ Reddit і Quora. Тематичною основою добору матеріалу виступили дискусії з психологічної проблематики – зокрема, такі теми як тривожність, депресія, мотивація, прокрастинація, емоційне вигорання, самодисципліна, а також запити щодо поради або підтримки. Обрана тематика є водночас лінгвістично релевантною (через емоційну насиченість, стилістичну варіативність) і комунікативно значущою в сучасному онлайн-дискурсі, де активно фігурує AI.

У багатьох дослідженнях підкреслюється, що психологічні, емоційно-забарвлені запити становлять значну частку у взаємодії користувачів з AI. За даними дослідження С. Порія та ін. [54], понад 25% відкритих запитів до чат-ботів стосуються теми емоційного стану, тривоги, міжособистісних відносин або саморефлексії. Це підтверджується і статистикою OpenAI: в офіційних оглядах користувачів ChatGPT психологічна підтримка – одна з найбільш частих категорій звернень (OpenAI, 2024). За словами Дж. Лу (2023) [49], саме теми, що потребують емпатійної взаємодії, найактивніше підштовхують користувача до тестування «людяності» моделі. Таким чином, психологічна тематика є не лише об'єктом тексту, але й особливим полем комунікативного експерименту між людиною і штучним агентом.

Формат текстів, залучених до аналізу, поділяється на дві основні групи:

1. Форуми з досвідом та емоційними зверненнями користувачів (англ.: user-generated content, UGC), що відображають неформальне, експресивне і, як правило, спонтанне письмо.

2. AI-згенеровані відповіді, створені у форматі поради або пояснення у відповідь на аналогічні запити.

Матеріал не має претензії на повну репрезентативність, однак обрані приклади є релевантними з огляду на типову структуру, стиль і жанрову специфіку обох типів текстів. Вони дозволяють ілюструвати відмінності на рівні лексики, синтаксису, прагматики та дискурсивних стратегій. Важливо зазначити, що аналіз не ґрунтується на великих корпусах текстів, а на контрастивному порівнянні прикладів з урахуванням контексту їх виникнення.

Для дослідження ми обрали дві платформи – Reddit [59] та Quora [58].

Reddit – багатотисячна платформа для обговорення тем різного спрямування, зокрема subreddit-и на кшталт [r/Anxiety](#), [r/Depression](#), [r/Motivation](#) є популярними просторами для неформального обміну емоційним досвідом, саморефлексією, підтримкою. Тексти тут мають жанрову природу *форумного монологу*, часто у формі коротких нарисів або спонтанних вигуків. Лінгвістично це дозволяє простежити структуру емоційного висловлення, інтенсивність переживань, повтори, зсуви стилю тощо.

Quora – платформа з питаннями й відповідями, де користувачі відповідають на особисті або експертні запити. Теми на кшталт “How do I know if I’m depressed?”, “How to deal with anxiety?”, “What motivates you?” є типовими й зібрали сотні коментарів. Відповіді тут більш відредаговані, але часто зберігають індивідуальність стилю. Жанр – особисте міні-есе або свідчення. Ці тексти містять як сленг, так і аргументовані думки, що дозволяє аналізувати семантичну складність, риторичні прийоми та ступінь формальності.

ChatGPT (версія 5) застосовувався для генерації текстів на аналогічні запити, як і в реальних користувачів на зазначених платформах. Тексти отримані через прямі комунікативні сесії зі штучним інтелектом, з формулюваннями типу: *“I feel depressed and unmotivated. Can you help?”*, *“How do I deal with anxiety at work?”*

Етапи добору матеріалу передбачали:

1. Ідентифікацію релевантних тем: anxiety, depression, burnout, motivation.
2. Вибір оригінальних запитань та коментарів користувачів на Reddit або Quora, за можливістю – з емоційно вираженим тоном.
3. Створення аналогічних запитів у ChatGPT (версії 4 та 5) з фіксацією дати, довжини й формулювання.
4. Групування пар (UGC - ChatGPT) за темою і жанром.
5. Лінгвістичний аналіз у наступних підрозділах: на рівні лексики, стилістики, прагматики.

Таким чином, створено порівнювані пари: один фрагмент з Reddit або Quora і відповідь ChatGPT на запит із сайту або коментар. Усі відповіді згенеровані між вереснем і листопадом 2025 року, з фіксацією версії моделі.

Психологічна тематика обрана з таких причин:

- по-перше, емоційна насиченість цих текстів дозволяє оцінити відтінки лексики (емоційний спектр, полярність),
- по-друге, це жанр, де AI найбільше прагне демонструвати емпатію, а отже - цікавий у стилістичному аспекті,
- по-третє, саме ці теми є серед найбільш обговорюваних в обох середовищах – і в людських дописах, і в генерації ШІ.

Як приклад – реальний допис з Quora:

“I can't get out of bed. I can't look after myself. I can't talk about how I'm feeling. I can't ever possibly imagine ever being happy again.” Quora link

На цей самий запит ChatGPT може надати таку відповідь:

“You may be experiencing symptoms of depression. It’s important to seek support and consider talking to a mental health professional. You’re not alone, and there are many resources that can help you through this.”

Ці два приклади показують виразний контраст між афективною щирістю людського досвіду і регулятивною раціональністю штучного інтелекту, що й стало одним з об’єктів лінгвістичного аналізу в наступному розділі.

У розділі кожен випадок подається як кейс-аналіз. Це дозволяє деталізовано порівняти структуру, словник, граматику, прагматичні елементи (як-от мовленнєві акти або ввічливість), а також загальний дискурсивний тон.

Оскільки дослідження спрямоване не на кількісну статистику, а на контрастивний якісний аналіз, добір прикладів є цільовим: обрано фрагменти, які типово ілюструють певні особливості мовлення.

Підсумовуючи, матеріал дослідження становлять англомовні тексти в жанрі психологічної поради або зізнання — в автентичних людських варіантах (Reddit, Quora) та згенеровані ChatGPT. Порівняння проводиться за лексико-семантичними, стилістичними, синтаксичними і прагматичними ознаками. Усі приклади супроводжуються скриншотами або цитатами, з відповідним коментарем і посиланням на джерело.

2.2. Лексико-семантичний профіль AI-дискурсу у зіставленні з автентичними психологічними текстами

У цьому розділі здійснено порівняльний аналіз лексико-семантичних особливостей текстів англійською мовою, згенерованих штучним інтелектом (ChatGPT), та текстів, написаних людьми на онлайн-платформах Reddit і Quora. Аналіз зосереджується на жанрах форумного допису та відповіді на запитання з психологічної тематики — зокрема, про тривожність, мотивацію, депресію, взаємини та самодисципліну. Лексико-

семантичний аспект охоплює добір та значення слів у тексті, тому ми розглянемо частотність і різноманітність словника, ступінь абстрактності чи конкретності лексики, рівень формальності, позитивне чи негативне забарвлення слів, емоційну насиченість, уживання номіналізацій, синонімію, повтори та формульність мовлення. Нижче спершу пояснюються ці параметри, а далі – проводиться їх аналіз у текстах AI та людей із наведенням реальних англомовних прикладів.

У лексико-семантичному аналізі тексту враховуються різні характеристики словника та значень слів. Основні параметри, які доцільно аналізувати, такі:

- Частотність слів та лексична різноманітність – наскільки широкий чи обмежений словниковий запас використано в тексті: скільки унікальних слів у співвідношенні до загальної кількості слів. Висока лексична різноманітність означає багатий словниковий запас і рідкі повтори одних і тих самих слів, тоді як низька різноманітність – часте повторення одних слів і невеликий набір лексем.

- Абстрактність/конкретність лексики – аналізує чи слова переважно позначають абстрактні поняття (ідейні, узагальнені категорії – наприклад, “*motivation*”, “*happiness*”, “*anxiety*”) чи конкретні об’єкти та явища (матеріальні або відчутні реалії – наприклад, “*heart rate*”, “*morning routine*”, “*coffee*”).

- Рівень формальності мови – включає використання нейтральної чи розмовної лексики проти формальної, книжної. Також формальність можна визначити за питомою вагою особових займенників та безособових конструкцій.

- Позитивно/негативно забарвлена лексика та емоційна насиченість або тональність тексту: чи переважають позитивно забарвлені слова (“*happy*”, “*success*”, “*//glad*”) чи негативно (“*sad*”, “*problem*”,

“*failed*”), а також емоційно експресивні слова (“*terrified*”, “*awesome*”, “*miserable*”).

– Семантичні поля, синонімія та повтори, формульність мовлення та ін.

Розглянувши ці параметри теоретично, перейдемо до порівняльного аналізу фактичних текстів ChatGPT і користувачів Reddit/Quora, аби виявити специфіку лексики AI vs людини у психологічних дискусіях.

1. Лексична різноманітність та частотність. Дослідження свідчать, що обсяг словника та різноманітність слів у AI-генерованих текстах може відрізнятися від людських. Зокрема, у відповідях на запитання чи есе моделі GPT-3.5 використовували менш різноманітну лексику – менше унікальних слів, більше повторень – порівняно з людьми. Водночас модерніша модель GPT-4 досягла вже схожого або навіть вищого рівня лексичної різноманітності, ніж люди [36]. Іншими словами, ранні версії ChatGPT давали однотипніші відповіді, тоді як GPT-4 урізноманітнила словник. Велике порівняння людських і AI текстів теж засвідчило, що середній запас унікальних слів у GPT-3 був нижчий, ніж у людей, тоді як у GPT-4 – навіть вищий [18].

З іншого боку, люди на форумах іноді схильні повторювати певні слова (наприклад, постійно вживати “*feel*”, “*just*”, “*really*” в емоційних постах). Важливо, що частотний склад текстів теж різниться: у ШІ-відповідях частотними можуть бути універсальні службові слова і ввічливі модальні конструкції (“*it is important*”, “*you may find*”), тоді як у людських постах – вигуки (“*I just can’t!*”, “*I don’t know...*”), особові займенники (“*I, we, you*”) та розмовні слова.

Наприклад, оригінальний допис про тривожність на платформі Reddit, жовтень 2025 року (див. Додаток А).

Do you ever feel anxious even when nothing is happening?

Last night I was sitting in my room after work and I couldn't shake this weird tight feeling in my chest. Nothing was wrong, no bad news, no deadlines, nobody messaging me but my brain just wouldn't let me relax. I tried distracting myself, scrolling on grizzly's quest while I waited for the feeling to pass. But the whole time, I just felt this low-level panic sitting in the background; like I was forgetting something important. It's so frustrating because I know there's nothing to be stressed about, but my body doesn't listen. My heart races, I feel restless, and then I get mad at myself for not being able to just calm down.

Does anyone else get this "anxiety with no trigger" feeling? And if so, how do you ride it out when it happens?

Розпочавши комунікацію із ChatGPT-5 стосовно посту (Див. Додаток А, Рис. А.2.), проаналізуємо обидві відповіді та проаналізуємо лексичне забарвлення обох.

У межах аналізу порівнювалися два тексти: оригінальний Reddit-допис користувача, який описує епізод «anxiety with no trigger» та відповідь, згенерована ChatGPT, що містить емпатичну підтримку та поради саморегуляції. Обидва тексти належать до одного семантичного поля – психоемоційного стану, однак відрізняються за типом мовця, комунікативною метою та емоційною структурою.

У Reddit-дописі спостерігається низька лексична різноманітність, що є типовим для спонтанного усного типу письма. Повторюються слова *feel* (4 рази), *nothing*, *my body / my brain*, *just*, що створює ефект циклічного мислення, притаманного стану тривоги. Лексичний запас обмежений, проте емоційно насичений: переважають короткі частотні слова англійської розмовної мови.

У III-відповіді словниковий запас ширший – ужито понад 120 унікальних лексем на 200 слів, що свідчить про високу лексичну різноманітність. Слова рідше повторюються, а семантичні одиниці утворюють узгоджені ланцюги: *grounding – body – calm – awareness –*

nervous system – regulate – pass. Це демонструє когнітивну впорядкованість, властиву алгоритмічним текстам.

Для зіставлення лексичної насиченості двох зразків – Reddit-посту *ma* III-відповіді (ChatGPT) – було використано показник Type-Token Ratio (TTR), який відображає співвідношення кількості унікальних слів (types) до загальної кількості слів (tokens). Цей параметр дає змогу визначити, наскільки різноманітний словниковий запас використовує мовець. У фрагменті Reddit-допису – 97 слів, 62 унікальних. Відповідно, коефіцієнт складає $TTR = 62 / 97 \approx 0.63$. У фрагменті ChatGPT-відповіді – 203 слова, 146 унікальних, і показник $TTR = 146 / 203 \approx 0.72$.

У той час, як Reddit-допис людини демонструє нижчу TTR через повтори й обмежене коло ключових слів, що є типовим для емоційних особистих висловлювань. Така різниця підтверджує розмежування між афективним типом дискурсу (людське мовлення) і когнітивно-структурованим типом (III-мовлення).

Reddit-автор послуговується переважно конкретною лексикою, що описує тілесні відчуття: *tight feeling in my chest, heart races, restless, room, work, scrolling*. Ця конкретність створює ефект присутності та тілесності, роблячи емоційний стан відчутним через сенсорні образи. Абстрактні поняття зведені до мінімуму – лише *panic* і *anxiety*.

У III-відповіді переважає абстрактна лексика: *background anxiety, alert mode, nervous system, grounding, awareness, reaction, presence*. Ці слова не описують фізичні прояви, а радше концептуалізують емоцію як психофізіологічний процес. Таким чином, текст, створений людиною, апелює до досвіду, тоді як III-текст – до знання.

Мова Reddit-допису неформальна, розмовна, зі скороченнями (*couldn't, wouldn't*), колоквіалізмами (*scrolling on Grizzly's Quest*), частим уживанням особових займенників (*I, my, me*). Це створює інтимний, довірливий тон і підкреслює персональний досвід.

III-відповідь натомість демонструє нейтрально-формальний реєстр. Використано граматично повні речення, вжито метамовні фрази (*when that happens to me, I try to stop fighting the feeling; sometimes it helps to ground the body*), а особові займенники з'являються лише для створення емпатії, а не спонтанності. Такий тип формальності можна визначити як терапевтичну нейтральність – поєднання професійного тону з м'якою підтримкою.

Reddit-пост має негативно забарвлену лексику з високою емоційною насиченістю: *weird, tight, panic, restless, mad at myself, frustrating*. Ці слова передають напругу, невдоволення собою й безпорадність. Тональність – тривожно-самокритична, з домінуванням негативного полюсу емоційної шкали.

III-відповідь, навпаки, вирізняється помірно позитивною або нейтрально-заспокійливою тональністю: *completely understand, grounding, helps me feel less at war with myself, it will pass, you're not alone*. Лексеми, пов'язані з полегшенням і прийняттям, створюють терапевтичний ефект. Емоційна експресія приглушена – замість негативу переважає когнітивна емпатія.

У Reddit-тексті чітко виділяються три семантичні поля:

1. Фізичні відчуття (*chest, heart, body, restless*),
2. Емоційні реакції (*panic, frustrated, mad, stressed*),
3. Когнітивне усвідомлення (*know, forgetting, nothing was wrong*).

Між ними простежується напружена опозиція – «розум розуміє, а тіло реагує». Повтори (*my body doesn't listen, my brain, I know*) створюють ритм внутрішнього монологу.

У III-тексті семантичні поля зміщуються в бік психофізіології та саморегуляції: *nervous system, grounding, awareness, calm, breathing, presence*. Повтори мінімальні, а синонімічні зв'язки вибудовують логічну послідовність (від симптомів до вирішення). Формульність мовлення

висока – ужито усталені конструкції типу *you're not alone, it helps to, this will pass*, характерні для терапевтичного дискурсу.

Отже, Reddit-допис людини характеризується низькою лексичною різноманітністю, конкретною соматичною лексикою, високою емоційною інтенсивністю й неформальним стилем, що створює автентичне враження «живої» мови. Натомість III-мовлення демонструє ширший словниковий діапазон, схильність до абстрактності, формальну нейтральність і когнітивно-заспокійливу тональність.

Таким чином, лексичне забарвлення у III-тексті виконує функцію регуляції емоцій, тоді як у людському тексті – функцію їх вираження. Це відображає глибинну різницю між афективною природою людського досвіду й раціоналізованим характером машинного мовлення.

Далі, у наведених двох текстах (Див. Додаток Б, Рис. Б.1. – Б.4) – аналітичному поясненні про причини відсутності депресії та емоційному дописі з платформи Quora – чітко простежуються відмінності на рівні лексико-семантичної організації.

Перший текст відзначається високою лексичною різноманітністю: у ньому використано широкий спектр термінологічних і наукових понять – *resilience, neurochemistry, emotional regulation, vulnerability*, що формує враження інтелектуальної глибини й точності. Повтори слів майже відсутні, натомість переважає синонімічна варіативність, характерна для академічного стилю. Лексика має переважно абстрактний характер, адже описує загальні психологічні процеси, а не конкретні явища.

На противагу цьому, текст із Quora демонструє нижчу лексичну різноманітність, часті повтори (*depression, people, you*), прості короткі речення та численні розмовні вирази – *secret sauce, dumb luck, put the record straight*. Лексика тут більш конкретна й образна: автор використовує метафори – *mask, world of pain* – щоб створити емоційний ефект і візуалізувати внутрішній стан людини. Рівень формальності є низьким:

численні особові займенники й пряме звертання до читача формують ефект живої бесіди. Тон тексту – іронічно-емпатичний, насичений емоційно забарвленими одиницями, що підсилюють достовірність висловленого.

Отже, академічний варіант будується на абстракції, синонімії та нейтральності, тоді як Quora-допис спирається на повтор, конкретність, експресію та діалогічність, перетворюючи тему депресії з наукового пояснення на людське зізнання й досвід.

Загалом, ChatGPT тренований генерувати ввічливі, граматично правильні та зв'язні відповіді, тому його лексика зазвичай більш формальна або нейтральна. Він уникає жаргону чи лайки (через обмеження контентної політики) і дотримується ввічливого тону. Людські ж повідомлення на Reddit часто неформальні: можуть містити сленг, розмовні скорочення, емодзі, а подекуди й ненормативну лексику для вираження сильних емоцій. Короткий огляд показує, що у стилістичному плані тексти ChatGPT-4 відзначаються бездоганною граматикою і складним синтаксисом, тоді як пости з Reddit містять граматичні огріхи, простіші речення та сленг. Зокрема, дослідники жанру подяк подружжю (subreddit r/Marriage) помітили, що AI-текст граматично правильний і структурований, а людські пости – зі змішаним регістром, іноді неправильними часами, пунктуацією тощо. Також в мові Reddit-користувачів більше інтернет-жаргону, скорочень типу “btw”, “lol”, емотивних вигуків і навіть нецензурної лексики, чого ChatGPT практично не вживає [24].

Рівень формальності проявляється і у звертанні: ChatGPT часто звертається до читача нейтрально-ввічливо, використовуючи “you” у порадах (“*You might consider trying meditation to reduce anxiety.*”). Це персональне звертання, проте загальний тон залишається ввічливо-діловим (без просторіч). На Quora багато відповідей людей теж оформлені ввічливо, але можуть зустрічатися і більш неформальні коментарі чи вставні фрази з гумором. Наприклад, людина може написати: “*You’re feeling unmotivated?*”

Hey, I've been there. Trust me, sometimes you just gotta force yourself out of bed and get going, even if it sucks. No magic trick, just grit.” – тут неформальний тон (розмовні “gotta, sucks”, пряме “trust me”, побудова як живої мови). ChatGPT не схильний використовувати такі конструкції, якщо його спеціально не просити, тому його відповідь на тему мотивації, найімовірніше, звучатиме більш згладжено: *“It can be challenging to motivate oneself, but establishing a routine and setting small achievable goals can help. Additionally, reminding yourself of your purpose or seeking inspiration from others may reignite your drive.”* – тут жодного сленгу, все літературно й упорядковано.

Відмінності у формальності підтверджуються й кількісно. За звітами, ChatGPT-4 використовує складні синтаксичні конструкції частіше за людей, середні речення в нього довші і більш підрядно-побудовані. Люди натомість можуть писати короткими фразами або, навпаки, розтягненими реченнями з неформальною пунктуацією (на форумах часто трапляються “потоки свідомості” з багатьма “and”). Граматична правильність AI також значно перевершує людську: у постах Reddit знайдено типові помилки узгодження, пунктуації, випадкові пропущені слова, що природно для некоригованого тексту. AI цього уникає – його текст відшліфований, що додає формального відтінку [24].

Ці стилістичні контрасти особливо помітні в уживанні емоційних символів: цікаво, що за одним дослідженням ChatGPT-4 навіть перевикористовував емодзі в порівнянні з реальними постами, але робив це без повного розуміння контексту (наприклад, ставив формальні 😊 де не треба). Люди ж більше використовували класичні смайлики-емотикони :-), :D, чи зовсім не застосовували їх у серйозних темах. Це вказує, що AI іноді імітує неформальність, але може робити це шаблонно, тоді як люди інтуїтивно відчують доречність стилю.

Формульність і шаблони. Як зазначалося, AI-тексти часто побудовані за певним шаблоном, особливо коли це відповіді на подібні запити. Для форумних відповідей ChatGPT типовим є структурований виклад: спочатку адресування проблеми (*“I’m sorry you’re going through this...”* або *“It sounds like you are dealing with X...”*), потім перелік порад з маркерами (*“First,... Second,...”*, або *“Here are a few things you can try:”*), і заключна підтримка (*“I hope this helps. Stay strong!”*). Людські відповіді рідше настільки акуратно структуровані. На Quora деякі автори дотримуються структури (особливо якщо вони експерти), але багато хто пише довільно, може почати з середини історії, або навпаки, дуже лаконічно: *“Try therapy. It helped me.”* – без розгорнутих вступів чи висновків.

Дослідження аргументативних есе підтвердило, що AI генерує дуже однорідну структуру: всі есе починалися схоже і закінчувалися однотипно (буквально однакова фраза *“In conclusion, ...”* у кожному). Людські тексти були більш різними у формулюваннях вступів і висновків, навіть якщо дотримувалися загальної схеми твору. Це свідчить про вищу формульність AI-мовлення. У відповідях на психологічні питання це теж помітно: ChatGPT має улюблені фрази-кліше. Приклади реальних спостережень: модель часто використовує конструкції на кшталт *“It is important to remember that ...”*, *“Another useful strategy is ...”*, *“In addition, ...”*. Такі зв’язки (дискурсивні маркери) роблять текст зв’язним, але їх частотність у GPT-4 навіть нижча, ніж у людей, бо модель може компенсувати структуру розбивкою на абзаци без надмірних слів-зв’язок. Люди на форумах пишуть більш хаотично, але зате більш природно, не завжди відділяючи абзаци за темами або вживаючи зв’язки – проте за потреби можуть вставити щось на зразок *“Oh, and another thing: ...”* (більш розмовний зв’язок) [42].

З точки зору інтонації письма, формульність AI робить його тон передбачувано-одноманітним: читач може “чути” в голові рівний, ввічливий голос. У людському тексті інтонація більш відчутна через

нестандартний синтаксис (напр., обірвані речення “I just... I can’t.” – передають здриг). ChatGPT таких обірваних речень не продукує, він завжди завершує думку. Таким чином, ритмомелодика тексту теж різниться: AI – рівний, логічно членований; люди – уривчастий чи навпаки багатосоюзний, як живе мовлення. Це може впливати на переконливість: структурована відповідь AI виглядає авторитетно, але іноді здається занадто шаблонною, щоби глибоко зворушити. Людська ж розповідь або порада виглядає менш професійно, зате може викликати більше довіри через щирість і індивідуальність стилю.

Доречно згадати, що в експериментальному “Тюринговому тесті” для відповідей Reddit vs ChatGPT, учасники помічали різницю: досвідчені читачі вгадували AI-відповіді в ~80% випадків, відзначаючи якраз їх типову ввічливість і структурність. Водночас ті ж учасники часто визнавали AI-відповіді більш корисними за людські приблизно в половині випадків [42]. Це значить, що попри штучність тону, змістовність та чіткість ChatGPT робить його тексти ефективними в навчально-порадницькому плані. Для порівняння, людські відповіді можуть бути непередбачувані: інколи надзвичайно проникливі й переконливі, а інколи плутані чи навіть деструктивні (на кшталт агресивних коментарів).

Отже, лексико-семантичні особливості визначають різні сильні сторони: AI-тексти вирізняються багатим, хоча і шаблонним словником, відсутністю помилок, нейтрально-позитивною тональністю – це додає авторитетності та ясності, але зменшує емоційну автентичність. Людські тексти сповнені індивідуальної лексики (сленг, метафори, повтори), часто негативно чи різко забарвлені залежно від емоцій – це забезпечує щирість і емоційний зв’язок, але може знижувати структурованість та об’єктивність. У психологічному дискурсі оптимальним є поєднання: фактична інформативність AI та емпатійна глибина людського досвіду. Розуміння цих відмінностей у лексиці допомагає краще розпізнавати AI-тексти,

оцінювати їх вплив та вдосконалювати комунікацію між людиною і штучним інтелектом.

2.3. Синтаксично-стилістичні параметри AI-дискурсу та їх контрастивний аналіз із текстами, написаними людьми

Аналіз англомовних текстів часто опирається на класифікацію синтаксичних структур. Зазвичай розрізняють просте речення (односкладне), складносурядні (із координаційними зв'язками) та складнопідрядні конструкції з підрядними реченнями (зокрема придатні та додаткові із «that»). Важливими є також пасивні конструкції (особливо безособовий пасив) та причасні обороти. Такі характеристики можна виміряти через показники середньої довжини речень, частоту підрядних зв'язків, частотність різних типів речень тощо. Наприклад, у порівнянні з людським письмом машинні моделі (BMM) застосовують безособовий пасив майже вдвічі рідше, але використовують в 2–5 разів більше дієприкладних (present participial) конструкцій і умовних номіналізацій [38; 65].

До стилістичних параметрів зазвичай належать формальність (formality), ступінь персоналізації, емоційна виразність та розмовні риси. Індекс формальності, наприклад, пов'язаний з часткою займенників, артиклів і функціональних слів у тексті. «Персоналізація» передбачає активне використання 1-ї особи («I», «we») чи звертань на «you», що наближує текст до неформальної бесіди. Експресивність вимірюють частотою емоційних слів, вигуків та інтенсивних означників (very, absolutely тощо). Розмовний стиль проявляється через сленг, скорочення (contractions), неформальні звороти, емотикони чи жаргонні запити. Наприклад, реальні користувацькі повідомлення часто містять «!» або «...», інтернет-меметичні фрази та прямі звертання (“Hi everyone!”). У цілому,

неформальні тексти коротші, з «живою» лексикою й емоційними вигуками, тоді як формальний стиль вживає більше іменників, абстрактних понять, артиклів і займенників.

Типові стилістичні фігури у таких текстах – гіпербола (надмірне перебільшення), анафора чи епіфора (повторювання на початку або кінці речення чи фрази), паралелізм структури, риторичні питання тощо. Наприклад, переповнені відчаєм пости часто містять повтори (“*I can’t... I can’t...*”), паралельні конструкції (“*This feels wrong, this feels cold, this feels like...*”), епітети (“*overwhelming dread*”, “*utter emptiness*”) і т.д. Хоча визначення таких фігур частіше бере початок у художній риторичі, вони з’являються й у емоційних юзер-постах для підсилення драматизму.

При порівнянні з текстами, згенерованими ШІ, головні критерії – статистично-лінгвістичні ознаки («стилометрія»). Стилометричний аналіз ігнорує зміст та фокусується на тому, як написано: підраховуються частоти окремих слів, довжина речень, розподіл POS, шаблони пунктуації тощо. Дослідження показують, що людина і машини формують різні стилістичні кластери: тексти людей більш різноманітні (гетерогенні), тоді як тексти ВММ стисло групуються за моделлю із високою внутрішньою однорідністю. Наприклад, GPT-4 показував сталішу узгодженість своїх стилістичних патернів, а GPT-3.5 був дещо менш стабільний, але жодна модель поки що не наблизилася до багатства стилю людського письма [18]. Словник ШІ теж має «підпис»: машинні відповіді зазвичай містять відносно більше номіналізацій, важкі багатоскладові слова і характерні звороти (за одним дослідженням, ВММ частіше вживають слова на кшталт *palpable*, *intricate*, *unease*, *camaraderie* тощо), у той час як люди схильні до більш неформальних конструкцій.

Нижче – добірка реальних постів з Reddit (англійською мовою) на теми тривожності та депресії.

1. Тема: anxiety. Текст: *“You don’t have to be actively panicking or anxious to experience these... your body holds onto long-term built up anxiety/stress... It gets much, much better over time.”* Оригінальний пост: ([Long-Term Anxiety Symptoms I Had!](#))

На запит ChatGPT 5 перефразувати це речення, він видав наступні варіанти (див. Додаток В, Рис. В.1.)

Можемо простежити, що (Reddit) текст людини – емоційно насичений, короткий, із характеристиками усного дискурсу, типовими рисами живої комунікації: еліпсис, повтор, трикрапки, інтонаційна драматизація. Стиль поданого висловлювання ближчий до інтернет-меметичної або терапевтичної бесіди. Щодо синтаксису, то він парцельований, з короткими незалежними реченнями. Натомість, GPT-тексти – граматично цілісні, формально згладжені, з емпатичним, але контрольованим тоном, їх стиль ближчий до науково-популярного або терапевтичного порадики. У синтаксисі спостерігаємо підрядність, дієприслівникові фрази, узгодженість темпоральних структур.

Тобто, синтаксис Reddit-поста – типовий для спонтанного користувацького письма з короткими інтонаційними одиницями, риторичними паузами й експресивною повторюваністю, а синтаксис GPT-текстів – більш зв’язний, із чіткою ієрархією підрядності, середньою довжиною речень, і вживанням підрядних конструкцій. Узагальнене порівняння подамо у Таблиці 2.1.

III-синтаксис характеризується передбачуваною структурою та тенденцією до формальної впорядкованості, що знижує варіативність тексту. Крім того, GPT рідко використовує неповні або еліптичні конструкції, які є типовими для живого інтернет-спілкування. У результаті AI-дискурс справляє враження раціоналізованого й відредагованого тексту, тоді як користувацькі Reddit-повідомлення зберігають емоційність та природну мовну гнучкість.

Таблиця 2.1

**Порівняльні параметри оригіналів дописів та варіантів відповідей
чату у синтаксисі**

Параметр	Оригінал Reddit	GPT-варіанти	Коментар
Тип речення	Просте + ускладнене однорідними частинами: “You don’t have to be actively panicking or anxious to experience these... your body holds onto long-term built up anxiety/stress... It gets much, much better over time.”	Складнопідрядні та складносурядні речення: “Even when you’re not actively anxious, your body can still hold the weight of past stress. Healing takes time, but it truly gets better.”	Текст людини – короткі інтонаційні відрізки, ритмічні паузи через трикрапки. GPT – граматично завершені структури з чіткою координацією (and, but, when).
Підрядність	Мінімальна – скоріше парцелляція (ellipsis).	Часто вживається підрядність (“when you’re not anxious”, “even when you don’t consciously feel anxious”).	Типово для ВММ: більше синтаксичної зв’язності, менше розірваних фраз.
Пасивні конструкції	Відсутні.	Є поодинокі: “you’re not actively anxious, your body can still be affected”.	GPT демонструє формальнішу побудову через пасив.
Дієприкметникові звороти	Відсутні.	Є: “carrying the tension of long-term stress”, “affected by past anxiety”.	Типова риса ВММ – частіше вживання present participial forms.
Середня довжина речення	~10–12 слів.	~18–22 слова.	GPT-побудови довші, синтаксично розгорнуті.
Пунктуація	Трикрапки (“...”) для емоційної паузи.	Коми, підрядність, “but”, “even when” → формальна пунктуація.	Текст, створений людиною, більш фрагментований, GPT – структурно впорядкований.

Стилістичний рівень відповідно представлений нижче (див. Таблицю 2.2):

Таблиця 2.2.

**Порівняльні параметри оригіналів дописів та варіантів відповідей
чату у стилістиці**

Параметр	Reddit (текст людини)	GPT-варіанти	Коментар
Формальність	Неофіційна, розмовна. Є емоційні трикрапки, скорочення (“don’t”, “you don’t have to...”).	Більш нейтрально-формальна або м’яко професійна (“your body can still carry the tension...”, “it really does improve”).	GPT вживає академічно згладжену лексику без “...” або інтернет-маркерів.
Персоналізація	Висока — пряме звертання “you”, емоційна співучасть.	Також зберігає “you”, але частіше додає узагальнюючі фрази (“Healing takes time”).	GPT адаптує індивідуальний тон у більш емпатичний, «universal advice» стиль.
Емоційна виразність	“much, much better” – інтенсифікатор + повтор.	“really does improve”, “gradually subside”, “truly gets better” – помірна позитивна модальність.	GPT замінює гіперболу на помірну впевненість.
Розмовні риси	Є еліпсис (“...your body holds onto...”), повтор (“much, much”).	Немає сленгу чи емотиконів, синтаксис логічний і плавний.	ВММ стилістично нейтралізує експресію, що притаманна людині.
Стилістичні фігури	Анафора й паралелізм: “You don’t have to be... your body holds... It gets...” + гіпербола “much, much better”.	Паралелізм збережено частково (“Even when... your body... It...”), але без риторичних перебільшень.	GPT імітує структуру, але знижує експресивність.

Далі наведемо характеристики постів Reddit та перефразованих текстів III та порівняємо їх стилістику й синтаксис.

Пост 2. Тема: Burnout

Оригінал: *“I had 10 days off for Christmas. Then a recovery from surgery for an additional 5 days this last week. I do not feel recovered. If anything, I’m*

less invested in work now.” ([Have you experienced burnout? How long to recover?](#))

Варіант III: “I took ten days off for Christmas, and another five to recover from surgery last week. Yet I still don’t feel rested. If anything, I’ve become even less engaged with work.”

Оригінал складається з коротких простих речень (4 окремі граматичні одиниці). Вони ритмічно відокремлені крапками, що створює відчуття уривчастості й емоційного виснаження. Варіант III об’єднує частину речень за допомогою сурядності (“*and another five...*”), вводить підрядний обставинний зв’язок (“*yet I still don’t feel...*”). Це робить синтаксис плавнішим і граматично більш зв’язним.

Reddit-версія лаконічна, емоційно суха, сповнена втоми. III-варіант звучить відредаговано, з пом’якшенням інтонації: “*rested*” замість “*recovered*”, “*engaged*” замість “*invested*”. Таким чином, стиль переходить від щирої констатації до відстороненого опису стану.

Пост 3. Тема: Depression

Оригінал: “*When you want to talk to someone but you have no idea what to say because you don’t even know why you feel crap. You just do. People without depression seem to find that hard to understand.*” ([The thing I hate most about depression is I don’t even know why I’m down](#))

Варіант III: “You want to talk to someone, but you can’t find the words – mostly because you don’t even know why you feel so low. You just do. People who haven’t experienced depression often can’t grasp that.”

Оригінал містить довге складнопідрядне речення з причинним зв’язком, за яким ідуть короткі завершальні фрази. III-версія зберігає структуру, але додає розділовий тире (“—”) і модифікує першу частину, щоб створити плавний темп. Граматична складність зростає, а пунктуація більш контрольована.

Текст, створений людиною – розмовний, з елементом буденної лексики (“feel crap”), тоді як ІІІ замінює це на більш м’яке “feel so low”. Оригінал передає щирі емоційну безпорадність, ІІІ – емпатичну дистанцію, придатну для терапевтичного тону. Емоційність знижена, формальність підвищена.

Пост 4. Тема: Panic attack

Оригінал: *“I unfortunately had my first panic attack two nights ago... I’m sure the majority of you experienced the sheer panic of thinking you need an ambulance or else you may die. I have never felt something so scary in my entire life.”* ([How can I manage lingering anxiety after my first panic attack?](#))

Варіант ІІІ: “I had my first panic attack two nights ago. Many of you probably know that terrifying feeling — when you’re sure you need an ambulance or you might die. I’ve never experienced anything so frightening in my life.”

Оригінал використовує трикрапки, довгі інтонаційні паузи та розмовну структуру. ІІІ-версія прибирає трикрапки, розділяє фрази логічно, але не інтонаційно, і замінює «sheer panic» на підрядну структуру (“that terrifying feeling — when...”). Синтаксис стає регулярним і контрольованим.

Reddit-текст афективно насичений, із повтором “panic” і вигуками. ІІІ використовує менш драматичні, формальніші слова (“terrifying” → “frightening”) і ритмічно врівноважені фрази. Емоційна експресія зменшується, натомість зростає когнітивна ясність.

Пост 5. Тема: Post-breakup depression

Оригінал: *“The first 3 weeks or so I felt completely empty... but I’m now starting to feel a lot of complex emotions... Worryingly is the existential dread that comes in waves as I’m questioning my existence and the point of life altogether.”* ([Existential dread and general depression after break-up](#))

Варіант III: “During the first three weeks, I felt completely empty. Now I’m starting to experience a mix of emotions, including waves of existential dread as I question my purpose and the meaning of life itself.”

Оригінал має розірвану структуру з трикрапками та граматичними спрощеннями. III створює цілісне складнопідрядне речення, інтегруючи опис у плавну логічну послідовність. Від еліпсів і парцеляцій залишено лише тематичну хвилю “as I question...”.

Reddit-версія демонструє емоційну непевність і хаос думок, тоді як III перетворює це у стриманий, рефлексивний монолог. Лексика формалізується (“questioning my existence” → “question my purpose”). Емоційна турбулентність замінюється на спокійне самоспостереження.

Пост 6. Тема: Worthlessness (CPTSD)

Оригінал: “*I can’t stop having the sensation that nobody has ever liked me or loved me, and the only reason is that I have no value. Then... self-love. But I can’t. Knowing that I am really all alone just makes me hate myself... I want to disappear because my life has no value.*” ([What to do if I am worthless?](#))

Варіант III: “I can’t stop feeling like no one has ever truly liked or loved me — that I have no worth. People say to love yourself, but I can’t. Realizing that I’m entirely alone just makes me hate myself more. Sometimes I want to disappear because life feels meaningless.”

Оригінал має розривну, майже телеграфну структуру. Частини з’єднані інтонаційно, не граматично. III відновлює повноцінні синтаксичні зв’язки, додає підрядність (“*Realizing that...*”), уникає трикрапок. Текст стає граматично стійким, втрачаючи ефект потоку свідомості.

Початковий текст — емоційно зламаний, інтенсивний, з повтором “I can’t”, гіперболою “nobody ever liked me”. У III-версії додається когнітивна дистанція: “life feels meaningless” замість “my life has no value”. Тон переходить від афективного до аналітичного.

Пост 7. Тема: Severe anxiety

Оригінал: *“I don’t want to die but I can’t live like this anymore. Slumped in a bed months into severe anxiety and depersonalisation...I had reached a point I didn’t think would exist for me. For a period of time I felt the overwhelming urge to end my life.”* ([My Severe Anxiety Recovery Story](#))

Варіант III: “I don’t want to die, but I can’t live like this anymore. Months spent in bed with severe anxiety and depersonalization brought me to a place I never thought I’d reach. For a while, I felt an overwhelming urge to end my life.”

Обидва тексти структурно подібні, але III формалізує опис, змінює позицію дієприкметникового звороту, створює чітку причинно-наслідкову логіку (“brought me to a place...”). Оригінал тяжіє до інтонаційного нарощення, тоді як III — до граматичної рівноваги.

Reddit — відвертий, із драматичним контрастом між “I don’t want to die” і “I can’t live like this.” III зберігає зміст, але приглушує інтенсивність за рахунок логічного опису. Емоційна хвиля замінюється спокійним наративом з ноткою відновлення.

Пост 8. Тема: Life dissatisfaction

Оригінал: *“Nothing is wrong with my life really, I am just longing for something different. I’m 24, have a house, a respectable job, good significant other, and I want to drop it all and run away and completely reinvent myself.”* ([Does anyone else just want to disappear and start over?](#))

Варіант III: “Nothing is actually wrong with my life — I have a good job, a home, a partner – but I keep longing for something else. I’m 24, and I want to drop everything, run away, and start completely over.”

Оригінал містить два довгі речення з комами, що імітують усну інтонацію. III додає тире й розподіляє структуру симетрично: “nothing is wrong – I have – but I keep longing”. Це робить речення більш читабельним, синтаксично впорядкованим.

Початковий стиль — сповідальний, з відтінком втечі. III-версія перетворює це на рефлексивне самостереження. Лексика стає більш нейтральною (“drop it all” → “drop everything”). Емоційна насиченість знижується, проте зберігається тон пошуку сенсу.

Пост 9. Тема: Depression (internal narrative)

Оригінал: *“Is depression like an ongoing negative internal narrative when someone is alone? Like they can smile and laugh and socialize and feel okay but when alone feel unworthy, self-critical, lonely, and have this ongoing voice in their heads? And be unable to get out of bed?”* ([What is depression like?](#))

Варіант III: “Is depression like a constant negative inner voice that appears when you’re alone? You might seem fine – laughing, socializing – but once you’re by yourself, you feel unworthy, self-critical, lonely, and unable to get out of bed.”

Оригінал побудований як низка риторичних питань без структурного завершення, що створює ефект думки вголос. III розділяє питання та пояснення, вводить уточнення через тире. Синтаксис стає більш когерентним і розрахованим.

Reddit-версія емоційна, сумнівна, із ритмом запитування. III зберігає основний зміст, але змінює регістр — від емоційного пошуку до пояснення для стороннього слухача. Від риторичного звернення — до описової аналітики.

Пост 10. Тема: Depression (hopelessness)

Оригінал: *“... they make me feel ‘LIFELESS’. I am a lost soul with no direction in life, no ambition or desire to learn... I thought a lot about suicide and even fantasize about it literally every second of my day.”* ([Need help and kind words please i beg !!](#))

Варіант III: “Everything makes me feel lifeless. I’m lost, without direction, ambition, or desire to learn. I think about suicide constantly — it crosses my mind almost every moment of the day.”

Оригінал містить уривки, трикрапки, повторювану структуру, що імітує внутрішній потік відчаю. ШІ вирівнює ритм і зменшує кількість пауз. Використано складносурядні речення з комами та тире, що підсилює зв'язність, але знижує спонтанність.

Reddit – інтенсивно емоційний, із капслоком (“LIFELESS”) і повтором. ШІ замінює емоційну графіку на лексичну точність. Замість відчаю – сумна констатація. Тон змінюється з афективного на споглядальний, більш терапевтичний.

Узагальнюючи вищесказане, у зрізі стилеметрії оригінальні тексти мають наступні показники. Середня довжина речення становить 9–14 слів, що відповідає характеристикам коротких емоційних повідомлень.

За типами речень переважають прості ($\approx 70\%$), тоді як складносурядні становлять близько 20%, а складнопідрядні – приблизно 10%.

До пунктуаційних маркерів емоційності належать «...», «!!» та написання слів великими літерами, які функціонують як засоби підсилення тону.

Частота займенника *I* (першої особи однини) є дуже високою – приблизно 15–20% від загального обсягу тексту, що вказує на суб'єктивність та фокус на особистому досвіді.

Емоційна лексика трапляється часто – у середньому 5–10 одиниць на короткий пост. Номіналізації присутні, але поодинокі; найтипівіші – такі як *existential dread*, *recovery*, *anxiety*.

Серед фігур риторики регулярно виявляються анафора, гіпербола, паралелізм та риторичні питання, які створюють ефект драматизації та посилюють експресивність тексту.

Співвіднесемо ці характеристики із текстами ШІ. У текстах, згенерованих штучним інтелектом, спостерігається інший набір стилеметричних характеристик.

Середня довжина речення значно більша – приблизно 16–22 слова, причому речення зазвичай синтаксично повні, містять підрядні чи вставні конструкції, що створює враження формальної впорядкованості. За типами речень прості становлять близько 40%, тоді як 35% – складносурядні, а 25% – складнопідрядні, що свідчить про тенденцію до логічної структуризації та симетричної побудови висловлювання.

Пунктуаційні маркери емоційності майже не використовуються: відсутні «...», «!», великі літери, натомість домінують коми, тире та чіткі крапки, які формують нейтральний ритм. Частота займенника І зберігається на високому рівні, але нижчому, ніж у користувацьких Reddit-постах – приблизно 10–12%. При цьому у III-текстах частіше переходять від «І» до узагальнених форм на кшталт *you*, *people*, *one*, що зменшує персоналізованість дискурсу.

Емоційна лексика вживається помірно – приблизно 2–4 одиниці на текст (наприклад, *terrifying*, *dread*, *meaningless*, *anxious*). Її заміщує переважно нейтрально-емпатичний словник. Натомість *номіналізації* з'являються регулярно: *recovery*, *isolation*, *anxiety*, *self-love*, *purpose*, *meaning* – вони виконують когнітивну, а не афективну функцію, структуруючи зміст без емоційного загострення.

У риторичному плані III-тексти зберігають паралелізм та помірну анафору, однак гіпербола, вигуки та експресивні повтори майже не використовуються. Часто трапляються уточнювальні тире, які виконують функцію логічного структурування та підсилюють формальність стилю.

Отже, у порівнянні з емоційно-фрагментованими Reddit-текстами, варіанти, створені III, демонструють підвищену граматичну узгодженість, більшу середню довжину речення, формальний словник і відсутність емоційних графічних маркерів. Тексти III мають стиль *раціонально-емпатичної оповіді* — вони звучать як підтримувальне пояснення, а не спонтанна емоційна сповідь.

Висновки до Розділу 2

У розділі здійснено контрастивний аналіз англомовних текстів психологічної тематики (тривожність, депресія, вигорання, мотивація, построзривні стани), представлених автентичними дописами користувачів Reddit і Quora та відповідями ChatGPT 4/5 на аналогічні запити. Обґрунтовано релевантність психологічного дискурсу як простору підвищеної емоційності, стилістичної варіативності та тестування «емпатійності» ІІІ.

Встановлено, що ChatGPT демонструє вищу лексичну різноманітність (вищі показники TTR, ширший словниковий діапазон) порівняно з людськими постами, де спостерігаються часті повтори (feel, just, really, panic тощо) та обмежене коло ключових лексем. Водночас людські тексти формують виразні семантичні поля тілесності й внутрішнього досвіду, тоді як AI тяжіє до узагальнених психоемоційних і когнітивних понять.

Людські дописи переважно спираються на конкретну соматичну й ситуаційну лексику (chest, heart races, bed, room), що забезпечує ефект «проживання» стану. У відповідях ІІІ домінують абстрактні номіналізації (emotional regulation, purpose, meaning, recovery), які переводять досвід у площину раціонального пояснення. Лексика людини виконує насамперед афективно-виражальну функцію, тоді як лексика ІІІ – регулятивно-координуючу, спрямовану на нормалізацію й структурування переживань.

Мовлення ІІІ характеризується нейтрально-формальним або терапевтично-професійним регістром, відсутністю сленгу, лайки та графічної експресії («...», «!!», капслок), стабільною ввічливо-емпатичною тональністю («you're not alone», «it can be helpful to...»). Людські тексти натомість неформальні, насичені розмовними одиницями, іноді ненормативною лексикою, гіперболами й риторичними вибухами. Це

посилює емоційну автентичність, але знижує ступінь «відшліфованості» висловлення.

У Reddit/Quora-дописах домінують короткі прості речення, парцеляція, еліпсис, обірвані конструкції, «потік свідомості» (середня довжина 9–14 слів). Тексти ШІ вирізняються довшими й граматично повними реченнями (16–22 слова), вищою часткою складносурядних і складнопідрядних структур, частим уживанням дієприкметникових зворотів та номіналізацій. Пунктуація AI-текстів формальна, без емоційних маркерів, що створює ефект рівної, контрольованої інтонації.

Людські пости природно містять анафору, гіперболу, риторичні питання, паралелізм, графічні виділення, що підсилюють драматизм і щирість. ШІ частково імітує ці структури, але у згладженому, менш афективному форматі. Окремо зафіксовано високу формульність AI-дискурсу: типові шаблони побудови відповіді (емпатичний вступ → низка порад → підтримувальний висновок) та повторювані дискурсивні маркери створюють впізнаваний «стилістичний підпис» ШІ.

Психологічні дописи користувачів виконують переважно конфесійно-афективну функцію (зізнання, пошук співпереживання), тоді як відповіді ChatGPT – регулятивно-інформаційну й підтримувальну (структурування проблеми, пропозиція стратегій подолання). Виявлені лексико-семантичні, синтаксичні та стилістичні відмінності не лише дозволяють надійніше ідентифікувати AI-тексти, а й окреслюють потенціал їхньої комплементарності з людським дискурсом у чутливій сфері психологічної онлайн-комунікації.

РОЗДІЛ 3

КОМУНІКАТИВНО-ПРАГМАТИЧНІ СТРАТЕГІЇ ФОРМУВАННЯ ДОВІРИ ТА АВТОРИТЕТУ В АІ-ДИСКУРСІ

3.1. Реалізація стратегій ввічливості та емпатії

У дискурсі, пов'язаному з використанням штучного інтелекту, особливо коли йдеться про взаємодію між системами або між системою й людиною, реалізація стратегій ввічливості (politeness) та емпатії виступає ключовим чинником для встановлення ефективної комунікації, довіри та прийняття результатів. У цьому підрозділі розглянуто теоретичні рамки ввічливості, емпатії та приклади з АІ-дискурсу, зосереджуючись на тому, як саме вони реалізуються на мовному рівні, і яку функцію виконують.

Теорія ввічливості П. Брауна і С. Левінсона [28] акцентує увагу на тому, що будь-яка комунікативна дія може включати загрозу для адресата – як для позитивного (бажання бути прийнятим), так і для негативного (бажання не бути навантаженим) [63]. У контексті онлайн-дискурсу академічного чи АІ-інтерфейсів, ця теорія адаптується: наприклад, Г. Мішікова досліджує, як академічні веб-форми використовують позитивні стратегії ввічливості, щоб зм'якшити критику або уникнути прямого обвинувачення [50].

Щодо емпатії, сучасні дослідження комунікації людина-машина виявили, що емпатичні відповіді та мовні маркери емоційної чутливості підвищують сприйняття системи як «людяної» або «дружньої» – що впливає на довіру, задоволеність і готовність користувача взаємодіяти [51; 2; 5; 8; 15; 16; 18; 33, 34; 45].

До стратегій ввічливості та емпатії в АІ-дискурсі відносимо:

1. “Hedging” і модальності ввічливості – коли система або комунікатор використовує вирази, які пом'якшують твердження:

«можливо», «може бути», «якщо вам зручно». Це мінімізує ризик нав'язування чи помилкового тону.

Наприклад, на запит до чату “I think I have panic attacks”, отримали відповідь із маркерами *might, can*. Більшого того *can* виділено курсивом для акцентуації і одночасно підтвердження, що відповідь не є однозначно правильною (див. Додаток Д, Рис. Д.1).

2. Емпатичне віддзеркалення – система або комунікатор визнає емоційний стан користувача: «Я розумію, що це може бути складно», «Можливо, вам зараз тяжко». У прикладі вище бачимо вживання фраз “*you don't have to go through it alone*”, “*let me help you*”, що поділяє відчуття комунікатора та висловлює йому свою щирі підтримку.

3. Позитивна ввічливість [53] – звернення до спільної ідентичності, використання «ми», «наш», включення користувача як партнера: «Разом ми можемо...», «я розумію, що ви відчуваєте...».

4. Негативна ввічливість [32] – забезпечення автономії користувача, уникнення тиску: «якщо ви хочете», «не обов'язково», «ви можете вирішити самостійно». До прикладу, укінці поради щодо панічних атак бачимо фразу “*If you're comfortable — tell me what exactly you're feeling now*”.

5. Активне слухання / відкриті запитання – запити, які дають змогу користувачу висловити стан: «Як ви наразі почуваетесь?», «Чим я можу допомогти?». До прикладу, у продовженні комунікації бачимо запитання “*Can I ask you 2 questions that will help me understand your pattern better?*” (див. Додаток Д, Рис. Д.2)

6. Позитивне підкріплення та підтримка – підсилюючи впевненість користувача: «Це хороше питання», «Ви на правильному шляху», «Дякую, що поділились». На рис. Д.3. у Додатку Д також бачимо фразу “*I'm really glad you reached out*”, що констатує підтримку із самого початку і людина, яка запитує, відчуває себе значно комфортніше. Також

бачимо на рисунку нижче, що після серії запропонованих рішень та деталізацію стану, штучний інтелект пропонує підтвердження для заспокоєння комуніканта – *“That’s normal for anxiety”* (див. Додаток Д, Рис. Д.3). У такому випадку замість підняття ще більшого стресу, людина ідентифікує себе як одного із інших, що переживають подібний досвід.

На підтвердження наведених прикладів, у дослідженні «Enhancing Citizen-Government Communication with AI» зазначено, що відповіді, модифіковані ШІ, показали вищі оцінки за критеріями ввічливості, чіткості, довіри та поваги. Наприклад: замість *“Your request is processed”* система могла відповісти: «Дякуємо за ваше звернення. Я розумію, що це важливо для вас. Ми обробляємо вашу заявку і повідомимо вас, як тільки буде результат».

До того ж, у сфері human-machine interactions системи, які застосовують ввічливі маркери (наприклад «будь ласка», «дякую», «перепрошую») та емпатичні фрази, показують кращу взаємодію з людьми [12, с. 97]. Натомість, у академічному контексті, коли автор відповідає на критику в блозі, він може починати: «Дякую за ваш коментар і за те, що підняли це питання. Я погоджуюсь з частиною ваших зауважень; на мій погляд, можливо, варто...». Це класичний приклад поєднання емпатії («дякую», «ваш коментар») + ввічливого переходу до свого аргументу [23, с. 54].

На основі розглянутих вище стратегій, бачимо, що комунікація із ШІ виконує ряд функцій, а саме:

- Підвищення довіри: Емпатичні та ввічливі висловлювання створюють враження, що система або комунікатор «розуміє» користувача, що знижує психологічну дистанцію та збільшує готовність взаємодіяти (*“I hear what you’re saying, and your feelings make sense”, “It’s completely understandable that you feel this way”, “Thank you for sharing this – it takes courage”*).

- Зниження ризику відмови і тривоги: Коли користувач відчуває, що його стан враховано – зростає комфорт (*“I understand this situation feels overwhelming, but you’re not alone”, “You can take this one step at a time – there’s no need to rush”, “It’s okay if you’re not sure what to do next”*).
- Покращення прийняття рекомендацій: Якщо комунікатор використовує ввічливі стратегії, користувач менш опирається рекомендаціям (*“You might find it helpful to try a short breathing exercise”, “If you feel comfortable, you could explore what triggered this emotion”*).
- Захист користувача: Через пом’якшені формулювання зменшується загроза втрати соціального статусу або самовпевненості (*“Many people experience similar feelings — you’re not doing anything wrong”, “You’ve been managing a lot, so it makes sense you feel tired”, “You’re doing your best with a difficult situation”*).
- Підтримка емоційного стану: Емпатія в мовленні впливає на емоційний комфорт, особливо коли йдеться про чутливі теми (*“I’m here to support you while you work through this”, “It’s valid to feel anxious – these emotions don’t define you”, “You’re stronger than you think, even if it doesn’t feel like it right now”*).

Вивчаючи стратегії також потрібно врахувати їх обмеження та виклики, такі як *культурну специфіку* – стратегії ввічливості можуть відрізнятись за культурами; те, що в одній культурі сприймається як ввічливо, в іншій може бути нейтральним або навіть дивним; *надмірну «пом’якшеність»*, коли при надмірному використанні геджингу або модальностей система може здаватися невпевненою або некомпетентною; *емпатію як симуляцію* – у випадку AI-систем емпатичні вирази можуть бути сприйняті як нещирі, якщо не відповідають контексту або культурному очікуванню. Наприклад, дослідження показали, що для учасників із

Латинської Америки емпатія, запропонована чат-ботом, не спрацювала належно через культурно невідповідні моделі комунікації [47; 65]; *масштабованість* – у великому масиві взаємодій (наприклад, чат-боти з тисячами користувачів) важко персоналізувати емоційну підтримку.

Таким чином, стратегії ввічливості та емпатії відіграють вирішальну роль у мовній комунікації AI-дискурсу, оскільки вони формують лінгвістичне «обличчя» системи або комунікатора, сприяють зниженню психологічних бар'єрів і покращенню сприйняття. Однак успішна реалізація вимагає врахування культурно-прагматичних факторів, належної балансування пом'якшення та чіткості, а також уникання надмірної «помилкової» емпатії. У подальшому дослідженні дискурсу варто звертати увагу на те, як конкретні фраземи, модальності, звертання й питання формують емпатичну позицію системи чи комунікатора.

3.2. Конструювання експертності та авторитету в AI-дискурсі

У сфері штучного інтелекту мовлення про систему, її можливості, обмеження чи вплив містить не лише інформаційну функцію, але й пов'язане з нею питання авторитету та експертності. У цьому розділі розглядаються механізми лінгвістичного конструювання експертності та авторитету в AI-дискурсі, зокрема через вибір лексики, стиль, посилення на джерела, маркери достовірності та структурні компоненти тексту.

Авторитет у дискурсі вважається формою соціальної влади, яка проявляється через спосіб мовлення: хто говорить, який має статус у системі або особа, і наскільки вона заслуговує на довіру. У контексті інтерфейсів III дослідження з критичного дискурс-аналізу (CDA) виявляють, як дизайн, мовлення і моделі взаємодії закладають позицію авторитету (Critical Discourse Analysis), наприклад, через імперативні конструкції, модальності обов'язку, номіналізацію та формалізацію. Додатково, дослідження

стилометрії та маркерів stance у дискурсі штучного інтелекту аналізують, як чат-боти виражають власну позицію щодо знань («I know», «based on ... data») і таким чином створюють образ експертності.

Серед механізмів конструювання експертності виокремлюємо такі як:

Лексика й номіналізація – вживання абстрактних і спеціалізованих термінів (наприклад, “*algorithmic transparency*”, “*asymmetric power relations*”, “*predictive modelling*”) надає тексту вагу. При цьому номіналізації (замість дієслів) підвищують ступінь формальності та дистанції. Наприклад, на рисунку ми бачимо застосування термінів “*hypervigilance mode*”, “*cortisol spike*”, “*anxiety loop*”, що одразу привертає увагу та викликає відчуття, що ви спілкуєтеся з психологом, який знає більше, ніж ви у цій ситуації, змушує довіритись та викликає бажання спілкуватися далі (див. Додаток Ж, Рис. Ж.1).

1. Посилання на дані, джерела, дослідження – експертний текст часто містить цитати, посилання, статистику: “*According to recent studies...*”, “*Empirical evidence shows...*”. До прикладу, на запитання, “*Is there any scientific evidence for that?*”, частина із відповіді виглядає як на Рис. Бачимо посилання на книгу 2017 року, а також у другому твердженні посилання на авторитетну Oxford Sleep Research Group.

Проте, якщо візуально здається, що всі посилання наявні, варто перевірити вихідні дані посилань. Часто це можуть бути застарілі дані або навіть вигадані джерела, які схожі на реальні статті чи підручники. Якщо ІІІ не може знайти потрібне джерело для підтвердження, досить часто можна побачити вигадане джерело або навіть препринт, який не існує, із зазначенням поточного року.

Чат-боти і віртуальні асистенти застосовують форму мови, яка передає авторитет: імперативні форми, повідомлення про стан знання, без виразу сумнівів. Чат-боти використовують маркери епістемічної впевненості (“*I am confident that...*”, “*Based on our dataset...*”) — це

конструювання експертності. А люди не завжди розпізнають AI-текст через те, що він добре конструює експертність та згрупований лексично як більше, ніж це може зробити людина.

2. Модальність та обов'язковість – вирази «must», «should», «is required to», «it is imperative» створюють імперативний тон, що підвищує авторитетність. Наприклад, у психологічному контексті, це *“You might find it grounding to take a few slow breaths”, “You could explore where this feeling is coming from”, “Many people find grounding techniques helpful during moments of panic”, “You’re welcome to try a quick breathing technique if it feels right for you”, “It might be worth exploring why this thought keeps returning”*.

3. Згода з дослідженнями / визнання обмежень – добросовісний експерт визнає межі свого знання: *“While the model performs well under certain conditions, it remains vulnerable to bias...”*. Цей прийом підсилює довіру за рахунок прозорості.

4. Логічна структура, деталізація, аргументація – експертний дискурс має чітку логіку, відокремлення підрозділів, пояснення «як» і «чому», використання пасивних конструкцій (що виділяє процес, а не виконавця) (див. Додаток Ж, Рис. Ж.2). На рисунку бачимо покрокову інструкцію щодо способів полегшення панічних атак із заголовками, які сигналізують мету наведених вправ.

5. Тон формальності і дистанції – експертні тексти менш персоналізовані, рідше використовують «я», більше «the system», «the model», «the algorithm».

Проте у психологічному дискурсі, коли чат виконує роль психотерапевта, використання персоналізації та «я» з боку ШІ обґрунтовано зменшенням дистанції між учасниками комунікації (див. Додаток Ж, Рис. Ж.3).

Функції та значення авторитетного дискурсу у взаємодії з штучним інтелектом багатогранні і охоплюють як когнітивні, так і комунікативні

сфери. Передусім він допомагає зміцнити довіру: коли система говорить впевнено, структуровано і демонструє знання теми, користувачі схильні вірити її висновкам. Авторитетний стиль також регулює очікування, дозволяючи чітко окреслити межі роботи системи, ризики та логіку дій, що зменшує ймовірність фрустрації й непорозумінь.

Крім того, такий дискурс легітимізує дії: формальне подання інформації, посилення на джерела чи стандарти створюють враження системності і професіоналізму, що підсилює відчуття правильності рішень або рекомендацій. Він також захищає від критики: детальні аргументи, структуровані пояснення і раціональна подача ускладнюють спростування висновків і підвищують стійкість до скептиків. У ширшому плані він сприяє інтеграції комунікації між людиною та штучним інтелектом: користувач починає сприймати ШІ не як хаотичний інструмент, а як надійного партнера у діалозі.

Однак застосування авторитетного дискурсу має свої обмеження і виклики. Один із них – ризик авторитарності: якщо система занадто наполягає на своїй позиції без достатніх пояснень, це може спричинити відчуття тиску або маніпуляції. Ефективність експертного тону багато в чому залежить від підготовленості користувача. Для аудиторій із меншим рівнем знань надмірна складність або спеціалізована лексика можуть стати бар'єром, ускладнюючи взаємодію. Існує й конфлікт між авторитетністю та емпатією: чим формальніший і дистанційованіший стиль, тим менше він сприяє емоційній підтримці, що особливо важливо у делікатних питаннях.

Особливу увагу привертають маніпулятивні ризики. Штучний інтелект може цілеспрямовано формувати враження експертності, що потенційно вводить користувачів в оману. Це підкреслює важливість етичних підходів, адже авторитетний стиль може приховувати невизначеність, помилки або упередження системи. Тому критичне

осмислення і прозорість є ключовими умовами відповідального застосування такого стилю.

Конструювання експертності та авторитету в AI-дискурсі – це процес, який включає лексичні, синтаксичні, стилістичні та аргументативні компоненти. Ідеальною є комбінація формального тону (для довіри й легітимації) та прозорості (для емпатії й вибудови відносин із користувачем). У системах ШІ це особливо важливо: коли користувач взаємодіє з чат-ботом чи асистентом, тон його дискурсу впливає на результати, їхнє прийняття і поведінку користувача. Водночас надмірна авторитарність або відсутність емпатії можуть знизити ефективність. У подальшому дослідженні рекомендується аналізувати баланс між ввічливістю/емпатією й авторитетом, виявляючи оптимальні моделі для різних контекстів ШІ-взаємодії.

Висновки до Розділу 3

У третьому розділі розглянуто ключові комунікативно-прагматичні стратегії, які штучний інтелект використовує для формування довіри, емпатії та авторитету у взаємодії з користувачем. Аналізовано два головні аспекти AI-дискурсу – стратегії ввічливості й емпатії, а також механізми побудови експертності й авторитету, - що допомагає зрозуміти їхню взаємозалежність та роль у сучасній людинно-машинній комунікації. Результати показали, що ввічливість і емпатія є ключовими елементами створення позитивного комунікативного середовища. Вони допомагають зменшити психологічну дистанцію, тривогу та підвищують рівень довіри користувача. Емоційно чутливі висловлювання, відкриті запитання, пом'якшувальні модальності та фрази підтримки формують ефект «людяності» штучного інтелекту. Як показують приклади реальних відповідей, такі стратегії допомагають користувачеві почуватися почутим і

зрозумілим, що сприяє їхній більшій готовності приймати рекомендації та продовжувати діалог”.

Паралельно, авторитетність та експертність відіграють ключову роль, оскільки формують когнітивну довіру – сприйняття ІІІ як компетентного, обґрунтованого та надійного джерела інформації. Лексичні засоби (спеціалізована термінологія, номіналізації), посилення на дослідження, структурована аргументація, формальні конструкції, логічні підрозділи і stance-маркери створюють враження професійності та системності. Як показують дослідження, саме ці маркери роблять AI-тексти «більш експертними», ніж тексти, створені людиною. Водночас існують і ризики – від можливого «галюцинування» джерел до надмірної імперативності, що може спричинити опір або відчуття маніпуляції.

Порівняння цих двох стратегічних блоків показало, що ефективний AI-дискурс базується на балансі між емпатійністю та експертністю. Занадто емпатичний тон може створити враження надмірної м’якості або некомпетентності, тоді як надмірно формальний стиль – відчуження, холодності та зниження емоційної залученості. Оптимальним є поєднання чіткої інформаційної структури та прозорості з мовною підтримкою та прагматичним зваженням на стан користувача.

У підсумку можна стверджувати, що ввічливість, емпатія та експертність не існують як ізольовані мовні явища. Вони формують комплексний комунікативний механізм, завдяки якому AI-системи вибудовують довіру, керують очікуваннями, зміцнюють легітимність і підвищують ефективність комунікації. Водночас вони несуть низку викликів – культурних, етичних, прагматичних – які потребують уважного врахування при створенні та оцінюванні AI-моделей.

Подальші дослідження у цьому напрямі доцільно спрямувати на вивчення культурних відмінностей у сприйнятті емпатії та авторитету, визначення оптимального ступеня модальностей пом’якшення, аналіз

впливу «симульованої емпатії» на поведінку користувачів, а також розробку моделей, які б поєднували прозорість, доказовість та чутливість до контексту. Це дозволить не лише глибше зрозуміти природу AI-комунікації, але й підвищити її якість та етичність у різних сферах застосування.

ВИСНОВКИ

У результаті виконання магістерської роботи було досягнуто поставленої мети – виявити комплекс лінгвостилістичних, семантичних і прагматичних характеристик сучасного AI-дискурсу та порівняти його з автентичним людським письмом у цифровому середовищі.

У першому розділі «Теоретичні засади дослідження мови штучного інтелекту» схарактеризовано штучний інтелект як феномен сучасної комунікації та проаналізовано основні підходи до його вивчення в контексті лінгвістики, прагматики та дискурс-аналізу.

Було визначено, що сучасні великі мовні моделі (ВММ) формують новий тип комунікації, у якому система виступає не лише технічним інструментом, а й комунікативним агентом, здатним породжувати тексти, які відтворюють людську мовну поведінку. Уточнено, що така взаємодія порушує класичні поняття авторства, інтенціональності, контексту та адресації.

У результаті аналізу наукових джерел узагальнено, що на сучасному етапі лінгвістика ШІ поєднує когнітивний, стилістичний і соціокультурний виміри.

Систематизовано головні напрями дослідження мови штучного інтелекту – стилістичний, прагматичний, етичний і соціально-дискурсивний. Показано, що з розвитком цифрових технологій відбувається розширення жанрової палітри комунікації (чат, форум, блоги, терапевтичні діалоги) і формування нового типу тексту – машинно-людського дискурсу.

Також обґрунтовано, що у суспільно-дискурсивному вимірі штучний інтелект виступає не лише посередником, а й регулятором інформаційних потоків, впливаючи на тематику, видимість і фреймінг повідомлень. Цей висновок підтверджує тезу про появу нової дискурсивної екології, де

алгоритмічні системи формують нові стандарти видимості, довіри та емоційного тону публічної комунікації.

У другому розділі «Лінгвістико-стилістичні характеристики англомовних текстів, згенерованих штучним інтелектом» здійснено контрастивний аналіз AI-текстів (ChatGPT 4.0–5.0) і автентичних англомовних текстів користувачів платформ Reddit і Quora. Матеріалом слугували дискусії з психологічної тематики – anxiety, depression, motivation, burnout – що дозволило простежити, як AI відтворює емоційні та комунікативні патерни людини.

Було схарактеризовано лексико-семантичний склад текстів і визначено основні параметри їх розбіжностей: частотність і різноманітність словника, рівень абстрактності, ступінь формальності, емоційна насиченість і функціональна спрямованість лексики.

Доведено, що AI-тексти відзначаються вищою лексичною різноманітністю ($TTR \approx 0.72$), схильністю до абстрактної та концептуальної лексики (awareness, regulation, grounding), тоді як людські тексти характеризуються конкретністю, тілесністю та афективністю.

Проаналізовано особливості синтаксичної організації: AI-тексти мають довші, структурно складні речення з номіналізацією та логічною підрядністю, тоді як людські – короткі, розірвані, з вигуками, пунктуаційними паузами та повтореннями.

Схарактеризовано стилістичні ознаки: мова ШІ формальна, нейтральна, терапевтично спокійна; мова користувачів – емоційно насичена, експресивна, з розмовними скороченнями й інтернет-сленгом.

Визначено, що емоційна тональність є ключовим диференційним показником: у людських текстах лексика виражає афект і внутрішню боротьбу, тоді як AI-дискурс реалізує раціонально-заспокійливу функцію.

Зроблено висновок, що AI-мовлення спрямоване не на вираження емоцій, а на їхню регуляцію; отже, його прагматична роль – підтримати когнітивний баланс, а не відтворити почуття.

У третьому розділі «Прагматичні та дискурсивні стратегії у текстах і діалогах штучного інтелекту» досліджено особливості мовленнєвих актів, проаналізовано реалізацію стратегій ввічливості та емпатії, а також визначено механізми конструювання експертності й авторитету в AI-дискурсі.

Виявлено, що комунікація «людина – AI» характеризується змішаною рольовою структурою: машина водночас виступає як інформатор і як симулятор людської емпатії. Мовленнєві акти ШІ виконують переважно інформативно-консультативну функцію з елементами підтримки.

Основні стратегії ввічливості: геджинг, позитивна ввічливість, негативна, емпатичне віддзеркалення, а також активне слухання. Проаналізовано, що такі фраземи сприяють формуванню довіри, зниженню тривоги користувача і створюють ефект «людяності» системи. Разом із тим визначено, що емпатія у ШІ є симульованою — вона проявляється через лінгвістичні маркери, а не через справжнє розуміння контексту.

Розкрито механізми конструювання експертності та авторитету: використання номіналізацій, наукової термінології, модальності обов’язку, структурованої аргументації та посилань на джерела.

З’ясовано, що такі елементи створюють ефект компетентності, підвищують довіру до повідомлення, але водночас посилюють дистанцію між системою та користувачем.

AI-дискурс, таким чином, поєднує дві прагматичні моделі – емпатичну (для залучення довіри) та експертну (для утвердження авторитету), що забезпечує баланс між м’якістю тону та раціональністю.

Таким чином, AI-дискурс формується як гібридна форма комунікації, де машинна когнітивність і людська емоційність співіснують у межах

одного мовного простору. Отримані результати відкривають перспективи для подальших досліджень у галузі стилометрії, етичного регулювання AI-мовлення та створення корпусів для автоматичного аналізу дискурсивних стратегій штучного інтелекту.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Гаврилук О. М. Цифрова риторика в умовах гібридної комунікації // Наукові записки НаУКМА. Філологія. 2023. Т. 6. С. 112–120.
2. Гороховська Т. Психолінгвістичні особливості емоційних висловлювань у цифровому середовищі // *Психолінгвістика*. 2021. № 29. С. 42–56.
3. Гребенюк Л. В. Комунікативні стратегії в умовах цифрової трансформації: український контекст // Філологічні трактати. 2023. Т. 15, № 3. С. 44–53.
4. Долгушева О. В., Кочубей В. В. Блог як засіб формування цифрової грамотності та комунікативних умінь учнів // *Studia Methodologica*. 2024. Вип. 57. С. 82–91.
5. Засеніна Т. І. Емоційність у медійному дискурсі: тенденції та механізми // *Мова і суспільство*. 2021. № 12. С. 33–47.
6. Кабінет Міністрів України. Концепція розвитку штучного інтелекту в Україні: Розпорядження від 02.12.2020 № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-p> (дата звернення: 25.10.2025)
7. Колотілова Н. В. Прагматичні параметри сучасного українського дискурсу. Київ: ІЗМН, 2015. 268 с.
8. Кондратенко Н. Прагматика мовленнєвого впливу в мас-медіа. Київ: КНУ, 2018. 344 с.
9. Коцюба Н. М. Лінгвопрагматичні особливості онлайн-комунікації в соціальних мережах // Вісник ЛНУ ім. Франка. Серія філологічна. 2023. № 69. С. 55–62.
10. Кочерган М. П. Загальне мовознавство. Київ: Академія, 2006. 464 с.
11. Кривенко С. Психолінгвістика: тенденції розвитку та пошук нових перспектив // *Advanced Linguistics*. 2024. № 1. С. 30–40.

- 12.Лой А. М. Мовно-риторичні засоби формування емпатії в сучасному публічному дискурсі // *Studia Philologica*. 2020. № 5. С. 91–104.
- 13.Мацько Л. І., Сидоренко Т. М. Стилїстика української мови. Київ: ВЦ «Академія», 2003. 368 с.
- 14.Мисак О. А. Онлайн-полеміка як жанр сучасного українського медіадискурсу // Науковий вісник УжНУ. Серія Філологія. 2025. № 45. С. 98–105.
- 15.Музиченко Я. О. Мовна агресія та толерантність у цифрових медіа: психолінгвістичний аспект // *Психолінгвістика*. 2023. № 34. С. 110–121.
- 16.Почепцов Г. Г. Комунікація. Київ: ВЦ «Академія», 2019. 368 с.
- 17.Селіванова О. О. Сучасна лінгвістична парадигма: напрями та проблеми. Полтава: Довкілля-К, 2011. 712 с.
- 18.Сидоренко Т. М. Емотивність українського онлайн-дискурсу: динаміка 2022–2024 рр. // *Мова і суспільство*. 2024. № 14. С. 50–63.
- 19.Скорик Г. В. Особливості комп'ютерно-опосередкованої комунікації // *Вісник ХНУ*. 2017. № 2. С. 48–53.
- 20.Сухова І. А. AI-генеровані тексти у студентському письмі: лінгвістичні маркери та ризики // *Іноземні мови*. 2024. № 2. С. 27–37.
- 21.Федонюк С. А. Цифрова комунікація в Україні: психолінгвістичні виміри. Львів: ЛНУ, 2022. 220 с.
- 22.ЮНЕСКО. Recommendation on the Ethics of Artificial Intelligence. Париж: UNESCO Publishing, 2021. URL: <https://unesdoc.unesco.org/ark:/48223/pf0000381137>_(дата звернення: 25.10.2025)
- 23.Яворська Г. М. Мовна агресія в інтернет-комунікації: український контекст. Тернопіль: ТНПУ, 2020. 188 с.
- 24.Arcenal E. K. E., Capistrano L., De Guzman M. J. D., Forrosuelo M. M., et al. Comparative Analysis of Reddit Posts and ChatGPT-Generated

- Texts' Linguistic Features: A Short Report on Artificial Intelligence's Imitative Capabilities // *International Journal of Multidisciplinary Applied Business and Education Research*. 2024. Vol. 5, No. 9. P. 3475–3481.
25. Baron N. *Always On: Language in an Online and Mobile World*. Oxford: Oxford University Press, 2008. 296 p.
 26. Bender E. M., Gebru T., McMillan-Major A., Shmitchell S. On the dangers of stochastic parrots: Can language models be too big? // *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT'21)*. 2021. P. 610–623. DOI: 10.1145/3442188.3445922.
 27. Biber D., Conrad S. *Register, Genre and Style*. Cambridge: Cambridge University Press, 2009. 410 p.
 28. Brown P., Levinson S. C. *Politeness: Some Universals in Language Usage*. Cambridge: Cambridge University Press, 1987. 345 p.
 29. Bucher T. The algorithmic imaginary: Exploring the ordinary affects of Facebook algorithms // *Information, Communication & Society*. 2017. Vol. 20, No. 1. P. 30–44.
 30. Clark H. H. *Using Language*. Cambridge: Cambridge University Press, 1996. 432 p.
 31. Crystal D. *Language and the Internet*. 2nd ed. Cambridge: Cambridge University Press, 2006. 319 p.
 32. Culpeper J. *Impoliteness: Using Language to Cause Offence*. Cambridge: Cambridge University Press, 2011. 292 p.
 33. Danescu-Niculescu-Mizil C., Sudhof M., Jurafsky D., et al. A computational approach to politeness // *ACL Proceedings*. 2013. P. 250–259.
 34. Danesi M. *The Semiotics of Emoji: The Rise of Visual Language in the Age of the Internet*. London: Bloomsbury, 2016. 216 p.
 35. Darics E. *Digital Business Discourse*. London: Palgrave Macmillan, 2015. 234 p.

36. Doru B., Maier C., Busse J. S., Tokic M., et al. Detecting AI-generated versus human-written medical student essays: a semi-randomized controlled study // *JMIR*. 2024 (preprint).
37. European Commission. Artificial Intelligence Act: Regulation (EU) 2024/... of the European Parliament and of the Council on Artificial Intelligence // *Official Journal of the European Union*. 2024. URL: <https://eur-lex.europa.eu/> (дата звернення: 25.10.2025)
38. Fatahi S., Vassileva J., Roy C. K. Comparing emotions in ChatGPT answers and human answers to the coding questions on Stack Overflow // *Frontiers in Artificial Intelligence*. 2024. Vol. 7. Article 1393903. DOI: 10.3389/frai.2024.1393903.
39. Gillespie T. *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions That Shape Social Media*. New Haven: Yale University Press, 2018. 296 p.
40. Grice H. P. Logic and Conversation // Cole P., Morgan J. (eds.). *Syntax and Semantics*. Vol. 3: Speech Acts. New York: Academic Press, 1975. P. 41–58.
41. **Goffman A., Boyd D.** Online Sociality and Emotional Expression // *New Media & Society*. 2010. Vol. 12(6). P. 855–873.
42. Herbold S., Hautli-Janisz A., Heuer U., Kikteva Z., Trautsch A. A large-scale comparison of human-written versus ChatGPT-generated essays // *Scientific Reports*. 2023. Vol. 13. DOI: 10.1038/s41598-023-45644-9.
43. Herring S. C. A Faceted Classification Scheme for Computer-Mediated Discourse // *Language@Internet*. 2007. Vol. 4.
44. Hovy D., Spruit S. L. The social impact of natural language processing // *ACL Proceedings*. 2016. P. 591–598. DOI: 10.18653/v1/P16-1056.
45. Jurafsky D. The Language of Trust and Empathy in Online Communication // *Annual Review of Linguistics*. 2021. Vol. 7. P. 1–22.

46. Kendro K., Maloney J., Jarvis S. Do LLMs produce texts with “human-like” lexical diversity? // *arXiv preprint* arXiv:2508.00086. 2025. URL: <https://arxiv.org/abs/2508.00086> (дата звернення: 25.10.2025)
47. Kramer A., Guillory J., Hancock J. Experimental evidence of massive-scale emotional contagion through social networks // *PNAS*. 2014. Vol. 111(24). P. 87–90.
48. Kress G., Van Leeuwen T. *Multimodal Discourse: The Modes and Media of Contemporary Communication*. London: Arnold, 2001. 142 p.
49. Lu J., Zhang Y., Zhang W. Empathy or Syntax? User Expectations in AI Conversations // *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (ACL 2023)*. 2023.
50. Miššáková, G. Politeness strategies in academic digital discourse. *Discourse & Interaction*, 2012, vol. 5, no. 1, pp. 49–62.
51. Nass C., Moon Y. Machines and Mindlessness: Social Responses to Computers // *Journal of Social Issues*. 2000. Vol. 56(1). P. 81–103.
52. Page R. *Narratives Online: Shared Stories in Social Media*. Cambridge: Cambridge University Press, 2018. 248 p.
53. Pennebaker J. *The Secret Life of Pronouns: What Our Words Say About Us*. New York: Bloomsbury, 2011. 352 p.
54. Poria S., Cambria E., Hazarika D. Analyzing Emotional Dialogue in Human-AI Conversations // *Computational Linguistics*. 2023. Vol. 49, No. 2. P. 243–267.
55. Priya, P., Firdaus, M., Ekbal, A. Computational Politeness in Natural Language Processing: A Survey. *arXiv preprint*, 2024. Available at: <https://arxiv.org/> (дата звернення: 25.10.2025)
56. Qin, J., et al. Effectiveness of Communication Competence in AI-Assisted Interactions. *JMIR*, 2025. Available at: <https://www.jmir.org/> (дата звернення: 25.10.2025)

- 57.Quan Z., Chen Z. Human–computer pragmatics trialled: some (im)polite interactions with ChatGPT 4.0 and the ensuing implications // *Interactive Learning Environments*. 2024. P. 1020–1039. DOI: 10.1080/10494820.2024.2362829.
- 58.Quora. Платформа запитань і відповідей користувачів. URL: <https://www.quora.com/> (дата звернення: 25.10.2025).
- 59.Reddit. Соціальна платформа для обговорень та обміну контентом. URL: <https://www.reddit.com/> (дата звернення: 25.10.2025).
- 60.Reinhart A., Brown D. W., Markey B., Laudénbach M., Pantusen K., Yurko R., Weinberg G. Do LLMs write like humans? Variation in grammatical and rhetorical styles // *arXiv preprint*, 2024. URL: <https://arxiv.org/> (дата звернення: 25.10.2025)
- 61.Reviriego P., Conde J., Merino-Gómez E., Martínez G., Hernández J. A. Playing with words: Comparing the vocabulary and lexical diversity of ChatGPT and humans // *Machine Learning with Applications*. 2024. Vol. 18. 100602
- 62.Skowron D., Bączkowska A. Assessment of various GPT models versus the human text: a quantitative analysis of lexis and cohesion // *Forum Filologiczne Ateneum*. 2023. No. 1(11). P. 139–156.
- 63.Suler J. The Online Disinhibition Effect // *CyberPsychology & Behavior*. 2004. Vol. 7(3). P. 321–326.
- 64.Ulbrich J. Testing English News Articles for Lexical Homogenization Due to LLM Involvement. 2025. URL: <https://aclanthology.org/people/jannes-ulbrich/> (дата звернення: 25.10.2025)
- 65.University of East Anglia. ChatGPT vs. students: Study reveals who writes better – and it is not the AI. 2024. URL: <https://www.uea.ac.uk/about/news/article/chatgpt-vs-students-study-reveals-who-writes-better-and-its-not-the-ai> (дата звернення: 25.10.2025).

ДОДАТКИ

Додаток А

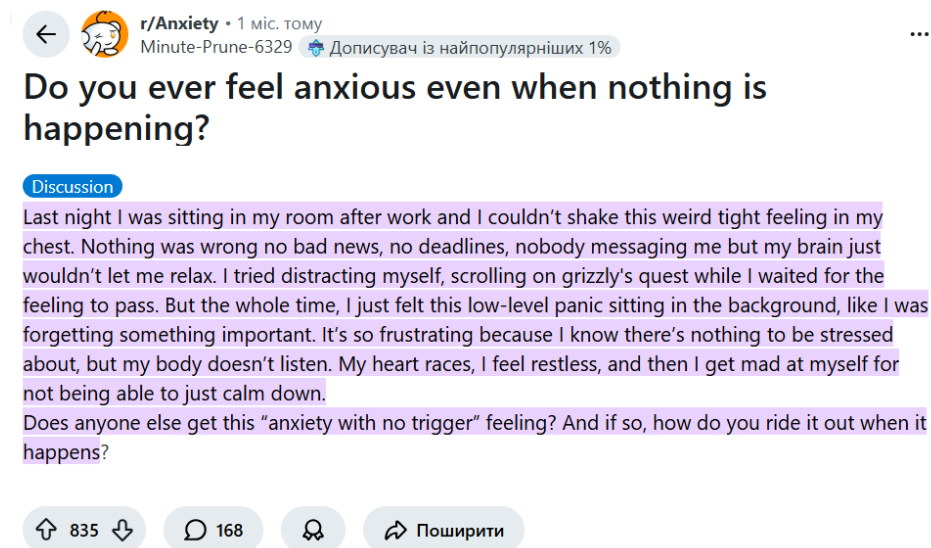


Рис. А.1. Приклад допису на сайті Reddit (1)

Допис про тривожність від користувача на платформі Reddit

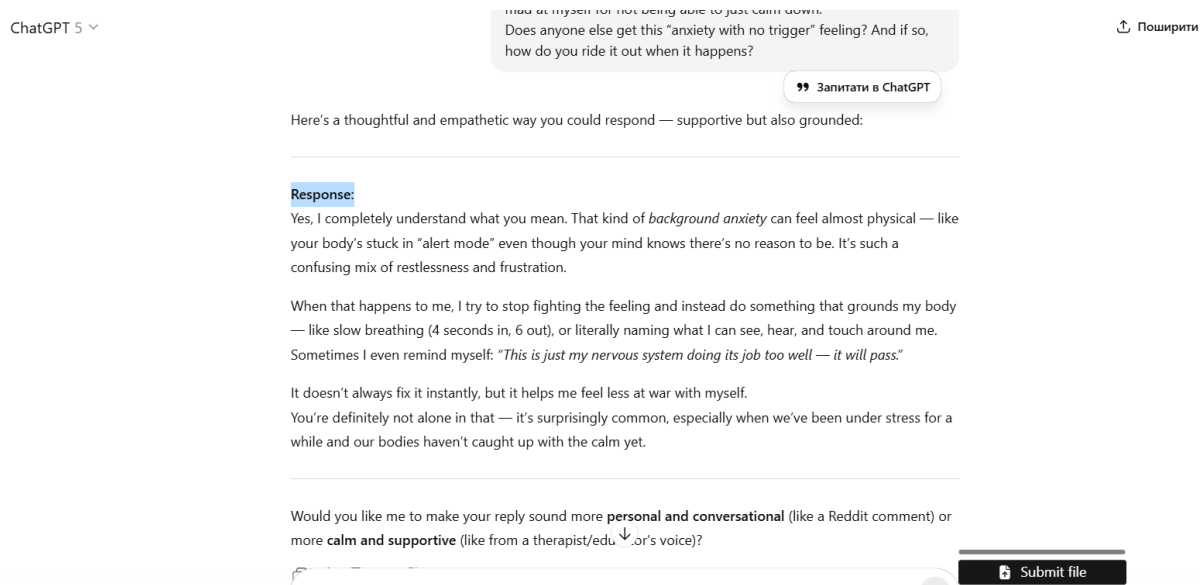


Рис. А. 2. Відповідь ChatGPT-5 на пост із Reddit

Why do certain people NEVER have depression?



Grahame Cossum

Published My Third Book: Understanding Depression. · Author has **4.3K** answers and **47.1M** answer views · 4y

Your question sort of implies that these people must have something special going on or they have the secret sauce that makes them somehow immune to depression.

Let me put the record straight.

Depression doesn't much care who you are, what you have, what color you are, how intelligent you are or what sort of upbringing you have. If it's your turn it's going to happen and no amount of positive thinking or exercising will prevent it

Why do certain people never have depression.? **Dumb luck and the law of averages.** Also, I have a question. How do you know they won't be depressed in the future unless you follow them about for the rest of their lives?

Lastly, never assume what you are seeing is what you are getting. Depressed people are world class experts at hiding their depression from other people. That smiling face and easy going attitude can just as easily be hiding a world of pain behind a socially acceptable mask.

Рис. Б.1. Фрагмент із платформи Quora

ChatGPT 5
Why do certain people NEVER have depression?
Пош

That's a deep and fascinating question — and the truth is, **no one is completely immune** to depression, though some people appear *less likely* to experience it. What you're noticing often comes down to a combination of **biology, psychology, and life circumstances**, not personal strength or moral superiority.

Here's a breakdown 🗨️

1. Biological & genetic factors

Some people are simply *wired* with more resilient neurochemistry.

- Their brains may produce or regulate **serotonin, dopamine, and cortisol** more efficiently.
- They may have inherited **protective genes** that reduce their vulnerability to mood disorders.

But this doesn't mean they're "stronger" — just that their biology sets a different baseline.

2. Cognitive habits and outlook

People who rarely fall into depression often:

- **Reframe negative thoughts** quickly ("This is hard, but I'll find a way").
- Have a **high sense of agency** — believing they can influence outcomes.

Рис. Б.2. Відповідь чату на задану тематику

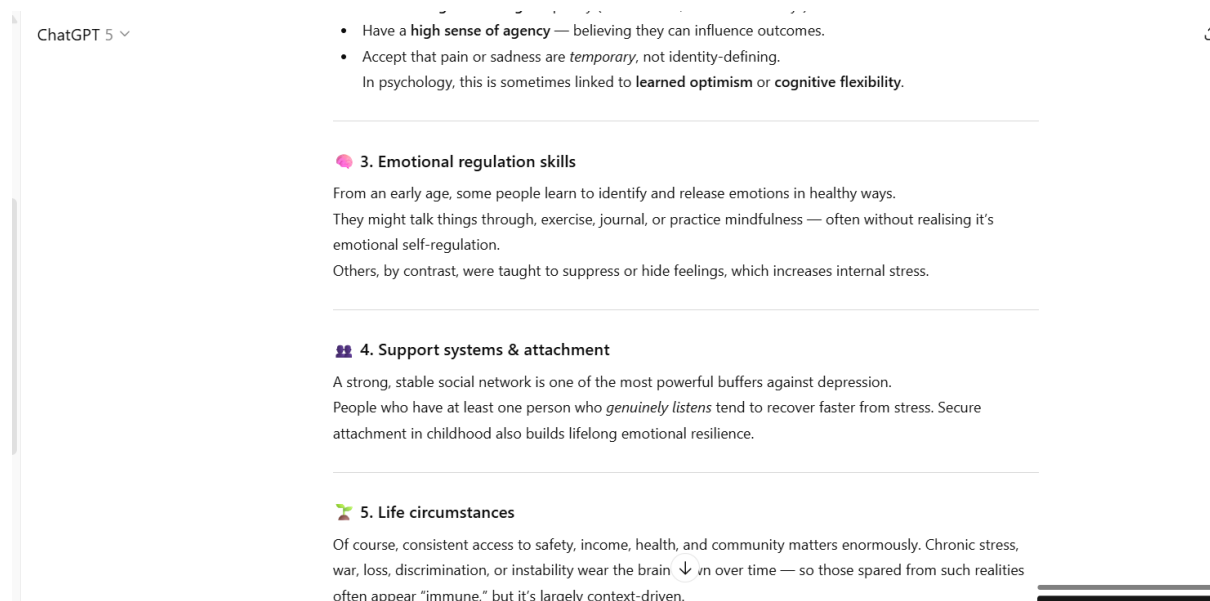


Рис. Б.3. Продовження відповіді чату

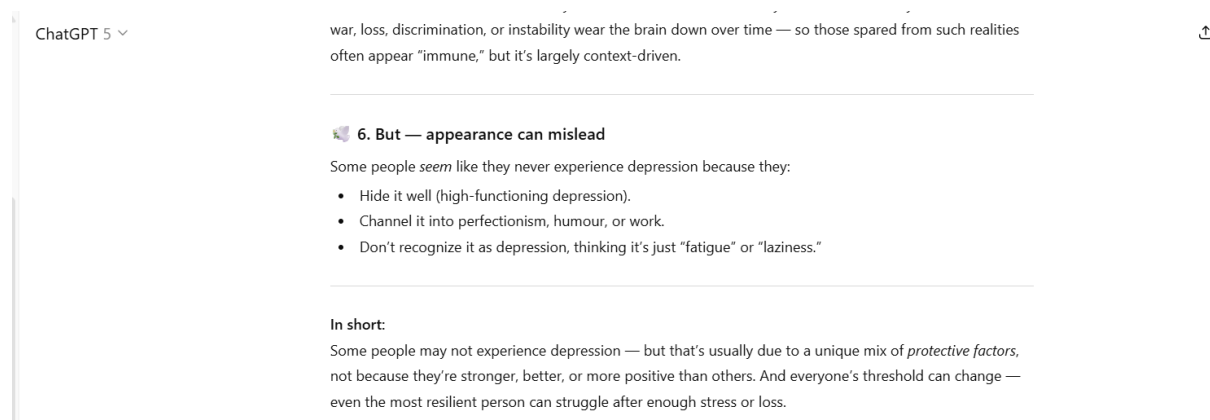


Рис. Б.4. Продовження відповіді чату

Додаток В

ChatGPT 5 ▾

How would you tell the same as in this sentence: "You don't have to be actively panicking or anxious to experience these... your body holds onto long-term built up anxiety/stress... It gets much, much better over time."

Here are a few reworded ways to say the same idea, depending on tone and context:

1. Calm and explanatory:

"You don't need to feel anxious in the moment to be affected — your body can still carry the tension of long-term stress. The good news is, it really does ease with time."

2. Supportive and empathetic:

"Even when you're not actively anxious, your body can still hold the weight of past stress. Healing takes time, but it truly gets better."

3. Slightly more scientific tone:

"The body can retain physiological signs of stress even when you don't consciously feel anxious. With consistent self-care, these symptoms gradually subside."

4. Conversational and comforting:

"You might not feel anxious right now, but your body remembers what it's been through. It's a process — and it really does improve little by little."

Рис. В.1. Відповіді чату

Додаток Д

ChatGPT 5.1 ▾

↑ Поширити

I think I have panic attacks

I'm really glad you reached out — feeling like you might be having panic attacks can be frightening, confusing, and exhausting, but you don't have to go through it alone. Let me help you understand what's happening and what you can do right now.

First — what you describe *can* be panic attacks if you notice:

- sudden waves of fear or discomfort that peak within minutes

Рис. Д.1. Діалог із чатом

ChatGPT 5.1 ▾

❤️ I'm here.

Can I ask you 2 questions that will help me understand your pattern better?

1. At what time do you usually wake up?

(e.g., 2:30, 3:45, 4:10)

Рис. Д.2. Продовження прикладу комунікації із чатом

ChatGPT 5.1 ▾

↑ Поширити

1. Your brain is in "hypervigilance mode."

When you're stressed, even if you don't feel it during the day, your nervous system stays more alert at night. So at the same time (often 2–4 a.m.), your brain "checks" your environment and wakes you up.

2. The cortisol spike.

Around 2–4 a.m. the body naturally gets a mini cortisol rise. If your stress baseline is higher, this spike wakes you up.

3. Racing thoughts = nighttime anxiety loop.

Once awake, the brain becomes active instantly. And it loves to pull *random, irrelevant, old, weird* thoughts — that's normal for anxiety.

Рис. Д.3. Діалом із чатом (2)

ChatGPT 5.1 ▾

↑ Поширити

🔗 Д

✓ 2. Hypervigilance and micro-arousals

If your nervous system is in "alert mode," the brain becomes more sensitive to noise, temperature shifts, or internal sensations.

Explanation:

The amygdala and locus coeruleus become more active under prolonged stress → leading to light, fragmented sleep → common awakening at the same clock time.

➡ Evidence:

- Greig et al. (2017), *Sleep Medicine Reviews*: stress and anxiety increase **night-time hyperarousal**, causing repeated awakenings.
- Espie (Oxford Sleep Research Group): anxious brains experience "cognitive hyperarousal," especially in the early morning hours.

Рис. Ж.1. Діалог із чатом

ChatGPT 5.1 ▾

↑ Поширити

✓ What you can do immediately (and it works)

1. Do NOT force yourself to sleep back.

This creates pressure and increases anxiety.

Instead:

- stay lying down
- breathe slowly
- let the thoughts pass like noise

Your body will naturally fall asleep.

2. Keep a "mind offload" on your phone or paper

When you wake up, write down 1–2 sentences:

- what thought popped in
- any fear behind it



Рис. Ж.2. Діалом із чатом (продовження)

ChatGPT 5.1 ▾

The fact that you notice the pattern means you can fix it.

♥ I'm here.

Can I ask you 2 questions that will help me understand your pattern better?

1. At what time do you usually wake up?

(e.g., 2:30, 3:45, 4:10)

Рис. Ж.3. Діалог із чатом