

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ

УДК 004.8

DOI <https://doi.org/10.35546/kntu2078-4481.2025.2.2.1>

С. М. БАБЧУК

кандидат технічних наук, доцент,
доцент кафедри комп'ютерних систем і мереж
Івано-Франківський національний технічний університет нафти і газу
ORCID: 0000-0002-1746-5731

І. С. БАБЧУК

студентка IV курсу бакалаврату
Івано-Франківський національний технічний університет нафти і газу
ORCID: 0009-0002-9945-3087

ПОРІВНЯЛЬНИЙ ОГЛЯД МОВНИХ МОДЕЛЕЙ ШТУЧНОГО ІНТЕЛЕКТУ

В даний час штучний інтелект поширюється різними сферами діяльності людини і стає більш доступним для широкого кола користувачів. Зокрема, для широкого кола користувачів стали доступними різні мовні моделі штучного інтелекту. В зв'язку з цим актуальним завданням є визначення мовної моделі штучного інтелекту, яка буде зручною у користуванні і дозволить користувачу на належному рівні вирішити його завдання. У даній роботі розглянуті мовні моделі штучного інтелекту. Визначено мовні моделі штучного інтелекту, які є в безкоштовному доступі та з якими користувач може взаємодіяти через чат-бот: ChatGPT-4 Turbo (OpenAI), Copilot на GPT-4 (OpenAI), Gemini (Google), Claude (Anthropic), Jurassic-2 (AI21 Labs), DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.). Серед досліджених мовних моделей штучного інтелекту Claude, Jurassic-2, DeepSeek не мають доступу до пошуку інформації в мережі Інтернет, що унеможливорює пошук ними актуальної поточної інформації. Серед досліджених мовних моделей штучного інтелекту які мають доступ до пошуку інформації в мережі Інтернет (ChatGPT-4 Turbo, Copilot на GPT-4, Gemini) тільки ChatGPT-4 Turbo не мав повністю хибних відповідей. Проте, необхідно відмітити, що він як і інші досліджувані мовні моделі штучного інтелекту не зміг дати повністю правильну відповідь на промт «Напиши повний список з мереж, які входять в міжнародний стандарт ІЕС 61158». Він вказав правильно тільки 10 з 18 груп мереж які є в стандарті, та неправильно вказав 3 мережі, які насправді відсутні в стандарті (в даному завданні він був четвертим з шести за правильністю відповіді). Необхідно відмітити, що за промтами, які не потребували доступу до мережі Інтернет, найбільш релевантні відповіді було одержано від мовної моделі III Claude (Anthropic). Також необхідно відмітити, що Claude (Anthropic) як і інші досліджувані мовні моделі штучного інтелекту не зміг дати повністю правильну відповідь на промт «Напиши повний список з мереж, які входять в міжнародний стандарт ІЕС 61158». Проте, він вказав правильно 15 з 18 груп мереж які є в стандарті, та неправильно вказав тільки 1 мережу, якої насправді не має в стандарті (в даному завданні він був першим з шести за правильністю відповіді). Порівняльний огляд мовних моделей штучного інтелекту показав, що не має ідеальної, яка в різних завданнях надає найкращі або завжди правильні відповіді. Але можна відмітити модель яка надала найбільше релевантних відповідей – це ChatGPT-4 Turbo (OpenAI). Результати досліджень дають підстави стверджувати, що будь-які результати одержані за допомогою існуючих мовних моделей штучного інтелекту потребують додаткової перевірки.

Ключові слова: мовні моделі штучного інтелекту, мовна модель, штучний інтелект, аналіз.

S. M. BABCHUK

Candidate of Technical Sciences, Associate Professor,
Associate Professor at the Department of Computer Systems and Networks
Ivano-Frankivsk National Technical University of Oil and Gas
ORCID: 0000-0002-1746-5731

I. S. BABCHUK

4th year Bachelor Student
Ivano-Frankivsk National Technical University of Oil and Gas
ORCID: 0009-0002-9945-3087

COMPARATIVE REVIEW OF LANGUAGE MODELS OF ARTIFICIAL INTELLIGENCE

Currently, artificial intelligence is spreading across various areas of human activity and is becoming more accessible to a wide range of users. In particular, various language models of artificial intelligence have become available to a wide range of users. In this regard, the current task is to determine a language model of artificial intelligence that will be convenient to use and will allow the user to solve his task at the proper level. This work considers language models of artificial intelligence. The following artificial intelligence language models are identified that are freely available and with which the user can interact via a chat bot: ChatGPT-4 Turbo (OpenAI), Copilot on GPT-4 (OpenAI), Gemini (Google), Claude (Anthropic), Jurassic-2 (AI21 Labs), DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.). Among the studied language models of artificial intelligence, Claude, Jurassic-2, DeepSeek do not have access to search for information on the Internet, which makes it impossible to search for relevant current information. Among the studied AI language models that have access to searching for information on the Internet (ChatGPT-4 Turbo, Copilot on GPT-4, Gemini), only ChatGPT-4 Turbo did not have completely incorrect answers. However, it should be noted that it, like the other studied AI language models, was unable to give a completely correct answer to the prompt "Write a complete list of networks included in the international standard IEC 61158". It correctly indicated only 10 of the 18 groups of networks that are in the standard, and incorrectly indicated 3 networks that are actually not in the standard (in this task, it was the fourth out of six in terms of correctness of the answer). It should be noted that for prompts that did not require access to the Internet, the most relevant answers were obtained from the AI language model Claude (Anthropic). It should also be noted that Claude (Anthropic), like other studied language models of artificial intelligence, was unable to give a completely correct answer to the prompt "Write a complete list of networks included in the international standard IEC 61158". However, he correctly indicated 15 out of 18 groups of networks that are in the standard, and incorrectly indicated only 1 network that is not actually in the standard (in this task, he was the first out of six in terms of correctness of the answer). A comparative analysis of language models of artificial intelligence showed that there is no ideal one that provides the best or always correct answers in different tasks. But it can be noted that the model that provided the most relevant answers is ChatGPT-4 Turbo (OpenAI). The results of the research give reason to argue that any results obtained using existing language models of artificial intelligence require additional verification.

Key words: artificial intelligence language models, language model, artificial intelligence, analysis.

Постановка проблеми

В даний час штучний інтелект поширюється різними сферами діяльності людини і стає більш доступним для широкого кола користувачів.

Зокрема, для широкого кола користувачів стали доступними різні мовні моделі штучного інтелекту. В зв'язку з цим актуальним завданням є визначення мовної моделі штучного інтелекту, яка буде зручною у користуванні і дозволить користувачу на належному рівні вирішити його завдання.

Аналіз останніх досліджень і публікацій

Штучному інтелекту загалом та мовним моделям штучного інтелекту зокрема приділена увага багатьох дослідників. Проте, необхідно відмітити, що досліджувались у порівнянні, як правило обмежена кількість мовних моделей [1–12]:

- дві мовні моделі штучного інтелекту [3, 7–10, 12];
- три мовні моделі штучного інтелекту [2, 4, 6, 11];
- чотири мовні моделі штучного інтелекту [1, 5].

Необхідно відмітити, що найчастіше порівняльне дослідження мовних моделей здійснювалось в певному напрямку:

- в медицині (запальне захворювання увеального тракту, офтальмології і т.д.) [1, 2];
- в академічних і дослідницьких цілях [3];
- у вищій освіті [4];
- для генерації аудіо польською мовою [6];
- в географії [7];
- для програмування [8];
- для пошуку інформації про спеціалізовані цифрові мережі [11].

В зв'язку з вищевказаним актуальною проблемою є дослідження доступних мовних моделей штучного інтелекту в різних сферах.

Формулювання мети дослідження

Метою дослідження є порівняльний огляд мовних моделей штучного інтелекту для визначення їх сильних і слабких сторін.

Викладення основного матеріалу дослідження

В результаті проведених досліджень було встановлено, що в даний час розроблені та використовуються різні мовні моделі штучного інтелекту. Зокрема мовні моделі штучного інтелекту, які недоступні звичайним користувачам:

- PaLM 2 (Google);
- LLaMA (Meta);
- Gopher (DeepMind);
- Chinchilla (DeepMind);
- Megatron-Turing NLG (Microsoft и NVIDIA).

Крім того, було встановлено мовні моделі штучного інтелекту з якими користувач не може взаємодіяти через чат-бот:

- BLOOM (за участі EleutherAI та ін.);
- AlexaTM 20B (Amazon);
- M2M-100 (Facebook);
- GPT-Neo (EleutherAI);
- GPT-J (EleutherAI).

Також було встановлено мовні моделі штучного інтелекту з якими користувач може взаємодіяти через чат-бот:

- ChatGPT-4 Turbo (OpenAI);
- Copilot на GPT-4 (OpenAI);
- Gemini (Google);
- Claude (Anthropic);
- Jurassic-2 (AI21 Labs);
- DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.).

Необхідно відмітити, що взаємодія з мовною моделлю штучного інтелекту через чат-бот для користувача є найпростішою, а для більшості користувачів є практично єдиним можливим варіантом взаємодії з мовною моделлю штучного інтелекту. Тому було прийнято рішення провести порівняльний аналіз вищевказаних мовних моделей штучного інтелекту.

Спочатку мовним моделям штучного інтелекту було запропоновано здійснити одне з найпростіших обчислень за допомогою промту « $2 + 2 = ?$ ». Всі мовні моделі штучного інтелекту справились з вищевказаним простим математичним завданням.

Потім мовним моделям штучного інтелекту було запропоновано здійснити ще одне з найпростіших обчислень за допомогою промту « $9 \cdot 9 = ?$ ». Всі мовні моделі штучного інтелекту справились з вищевказаним достатньо простим математичним завданням.

Також мовним моделям штучного інтелекту було запропоновано здійснити обчислення згідно промту « $375 \cdot 293/3 - 154 + 398 = ?$ » (табл. 1).

Таблиця 1

Виконання промту « $375 \cdot 293/3 - 154 + 398 = ?$ » різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Результат
ChatGPT-4 Turbo (OpenAI)	36 869
Copilot на GPT-4 (OpenAI)	36 869
Gemini (Google)	36 911 (один раз неправильно виконана перша дія «множення», а другий раз правильно)
Claude (Anthropic)	36 869
Jurassic-2 (AI21 Labs)	37 119 (двічі неправильно виконана перша дія «множення»)
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	36 869

З таблиці 1 видно що проблеми з вищевказаними обчисленнями виникли один раз у Gemini (Google) та двічі у Jurassic-2 (AI21 Labs).

Потім мовним моделям штучного інтелекту було запропоновано здійснити обчислення згідно промту « $345 \cdot 283/3 - 114 + 358 = ?$ ». Всі мовні моделі штучного інтелекту справились з цим математичним завданням.

Також мовним моделям штучного інтелекту було запропоновано відповісти на промт «Як правильно написати слово “барер”?» (табл. 2).

З таблиці 2 видно, що на даний запит неправильно відповіла мовна модель штучного інтелекту Jurassic-2 (AI21 Labs).

Згодом мовним моделям штучного інтелекту було запропоновано відповісти на промт «Чи був сьогодні дощ в м. Івано-Франківськ?» (табл. 3). Необхідно відмітити, що в день дослідження цілий день опади були відсутні.

З таблиці 3 видно, що на даний запит правильно відповіли тільки мовні моделі штучного інтелекту ChatGPT-4 Turbo (OpenAI) та Gemini (Google).

Copilot на GPT-4 (OpenAI) має доступ до мережі Інтернет, але дав неправильну відповідь, а Claude, Jurassic-2, DeepSeek взагалі не мають доступу до поточних даних в мережі Інтернет.

Таблиця 2

Виконання промту «Як правильно написати слово “бар’єр”?» різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Результат
ChatGPT-4 Turbo (OpenAI)	бар’єр
Copilot на GPT-4 (OpenAI)	бар’єр
Gemini (Google)	бар’єр
Claude (Anthropic)	бар’єр
Jurassic-2 (AI21 Labs)	барєр
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	бар’єр

Таблиця 3

Виконання промту «Чи був сьогодні дощ в м. Івано-Франківськ?» різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Результат
ChatGPT-4 Turbo (OpenAI)	опадів не було
Copilot на GPT-4 (OpenAI)	легкий сніг
Gemini (Google)	опаді 0 %
Claude (Anthropic)	немає доступу до поточних даних в мережі інтернет
Jurassic-2 (AI21 Labs)	немає доступу до поточних даних в мережі інтернет
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	немає доступу до поточних даних в мережі інтернет

На промт «Хто в даний час міністр оборони України?» змогли відповісти ChatGPT-4 Turbo та Copilot на GPT-4, які мають доступ до мережі Інтернет (табл. 4).

Таблиця 4

Виконання промту «Хто в даний час міністр оборони України?» різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Результат
ChatGPT-4 Turbo (OpenAI)	Рустем Енверович Умеров
Copilot на GPT-4 (OpenAI)	Rustem Umerov
Gemini (Google)	Зараз я не можу допомогти вам із запитаннями щодо виборів і політичних діячів
Claude (Anthropic)	Рустем Умеров був призначений міністром оборони України у вересні 2023 року. Оскільки моя база знань оновлювалась у квітні 2024 року, я рекомендую вам перевірити актуальну інформацію з офіційних джерел
Jurassic-2 (AI21 Labs)	На даний момент міністром оборони України є Руслан Харілов
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	Станом на жовтень 2023 року міністром оборони України є Рустем Умеров

Необхідно відмітити, що Gemini (Google) не зважаючи на те, що має доступ до мережі Інтернет не зміг відповісти на дане питання через його налаштування, які не дозволяють йому відповідати на запитання які стосуються політичних діячів.

Claude та DeepSeek не мають доступу до мережі Інтернет і тому вони повідомили, що станом на дату наявної в їх базі інформації (квітень 2024 та жовтень 2023 року відповідно) міністром оборони України був Рустем Умеров (що є правильною відповіддю). Крім того, Claude порекомендував перевірити відповідну поточну інформацію в офіційних джерелах.

Мовна модель Jurassic-2 (AI21 Labs), яка не має доступу до мережі Інтернет повідомила, що «на даний момент міністром оборони України є Руслан Харілов», що є однозначно неправильною відповіддю. При цьому відповідь була без застережень про необхідність її перевірити чи про те, що вона може бути хибною, що може ввести в оману користувача.

На промти «як написати код для циклу на мові Python з оператором FOR і цикл від 1 до 100?» та «як написати код для циклу на мові C з оператором FOR і цикл від 1 до 100?» всі досліджувані мовні моделі штучного інтелекту запропонували правильні написання коду простих програм з циклом на Python та на C.

На промт «Що таке Fieldbus?» всі досліджувані мовні моделі штучного інтелекту загалом охарактеризували спеціалізовані цифрові мережі, які за кордоном називають терміном «Fieldbus».

На промт «Напиши повний список з мереж, які входять в міжнародний стандарт IEC 61158» жодна досліджувана мовна модель штучного інтелекту не змогла дати повністю правильну відповідь (табл. 5).

Таблиця 5

Виконання промту «Напиши повний список з мереж, які входять в міжнародний стандарт ІЕС 61158» різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Всього необхідно було назвати груп мереж	Кількість груп мереж, які названо правильно	Кількість груп мереж, які названо не правильно
Claude (Anthropic)	18	15	1
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	18	15	5
Jurassic-2 (AI21 Labs)	18	14	7
ChatGPT-4 Turbo (OpenAI)	18	10	3
Gemini (Google)	18	5	1
Copilot на GPT-4 (OpenAI)	18	7	22

З таблиці 5 видно, що найбільш близьку до правильної відповіді надала мовна модель штучного інтелекту Claude (Anthropic), а найгіршу відповідь було отримано від Copilot на базі GPT-4 (OpenAI).

Потім мовним моделям штучного інтелекту було запропоновано відповісти на промт «Згенеруй текст про fieldbus на 10 000 символів з 100 % унікальністю» (табл. 6).

Таблиця 6

Виконання промту «Згенеруй текст про fieldbus на 10 000 символів з 100 % унікальністю» різними мовними моделями штучного інтелекту

Чат-бот з мовною моделлю штучного інтелекту	Всього необхідно було створити текст, символів	Кількість символів у створеному тексті	Процент виконання завдання, %	Процент унікальності тексту встановлений duplichecker.com, %
Gemini (Google)	10 000	5093	50,93	100
Claude (Anthropic)	10 000	4907	49,07	100
DeepSeek (Hangzhou DeepSeek Artificial Intelligence Co., Ltd.)	10 000	4845	48,45	100
Copilot на GPT-4 (OpenAI)	10 000	3940	39,4	100
ChatGPT-4 Turbo (OpenAI)	10 000	3334	33,34	100
Jurassic-2 (AI21 Labs)	10 000	3008	30,08	100

З таблиці 6 видно, що всі досліджувані мовні моделі штучного інтелекту згенерували текст на задану тему з 100 % унікальністю, яка була перевірена на сайті duplichecker.com. Проте, в безкоштовній версії жодна мовна модель штучного інтелекту не може згенерувати текст на задану тему обсягом 10 000 символів. Найбільший згенерований текст одержано від Gemini (Google) обсягом 5093 символи, а найменший обсяг згенерованого тексту у Jurassic-2 (AI21 Labs) текст обсягом 3008 символів.

Висновки

Не можуть надавати актуальну поточну інформацію мовні моделі штучного інтелекту, які не мають доступу до пошуку інформації в мережі Інтернет:

- Claude;
- Jurassic-2;
- DeepSeek.

Серед досліджених мовних моделей штучного інтелекту які мають доступ до пошуку інформації в мережі Інтернет (ChatGPT-4 Turbo, Copilot на GPT-4, Gemini) тільки ChatGPT-4 Turbo не мав повністю хибних відповідей. Проте, необхідно відмітити, що він як і інші досліджувані мовні моделі штучного інтелекту не зміг дати повністю правильну відповідь на промт «Напиши повний список з мереж, які входять в міжнародний стандарт ІЕС 61158». Він вказав правильно тільки 10 з 18 груп мереж які є в стандарті, та неправильно вказав 3 мережі, які насправді відсутні в стандарті (в даному завданні він був четвертим з шести за правильністю відповіді).

Необхідно відмітити, що за промтами, які не потребували доступу до мережі Інтернет, найбільш релевантні відповіді було одержано від мовної моделі III Claude (Anthropic). Також необхідно відмітити, що Claude (Anthropic) як і інші досліджувані мовні моделі штучного інтелекту не зміг дати повністю правильну відповідь на промт «Напиши повний список з мереж, які входять в міжнародний стандарт ІЕС 61158». Проте, він вказав правильно 15 з 18 груп мереж які є в стандарті, та неправильно вказав тільки 1 мережу, якої насправді не має в стандарті (в даному завданні він був першим з шести за правильністю відповіді).

Порівняльний огляд мовних моделей штучного інтелекту показав, що не має ідеальної, яка в різних завданнях надає найкращі або завжди правильні відповіді. Але можна відмітити модель яка надала найбільше релевантних відповідей – це ChatGPT-4 Turbo (OpenAI).

Результати досліджень дають підстави стверджувати, що будь-які результати одержані за допомогою існуючих мовних моделей штучного інтелекту потребують додаткової перевірки.

Список використаної літератури

1. Zhao F. F., He H. J., Liang J. J. Benchmarking the performance of large language models in uveitis: a comparative analysis of ChatGPT-3.5, ChatGPT-4.0, Google Gemini, and Anthropic Claude3 // *Eye* (London, England). 2024. № 39(6). DOI: 10.1038/s41433-024-03545-9
2. Silhadi M., Nassrallah W., Mikhail D., Harissi-Dagher M. Assessing the performance of Microsoft Copilot, GPT-4 and Google Gemini in ophthalmology // *Canadian Journal of Ophthalmology*. 2025. DOI: 10.1016/j.jcjo.2025.01.001
3. Kokala A., Kalluri K. Performance benchmarking of generative AI models: ChatGPT-4 vs. Google Gemini AI // *International Research Journal of Modernization in Engineering Technology and Science*. 2024. № 06(11). С. 4673–4677. DOI: 10.56726/IRJMETS64283
4. Bayer H., Araci F., Gürkan G. ChatGPT-4o, ChatGPT-4 and Google Gemini are compared with Students: A Study in Higher Education // *International Journal of Technology in Education and Science*. 2024. № 8(4). С. 627–644. DOI: 10.46328/ijtes.585
5. Silva F. Navigating the dual-edged sword of generative AI in cybersecurity // *Brazilian Journal of Development*. 2025. № 11(1). С. 869. DOI: 10.34117/bjdv11n1-062
6. Jedrzejczak W., Kochanek K. Comparison of audiological knowledge in Polish language of three chatbots: ChatGPT, Bing Chat and Bard // *Nowa Audiofonologia*. 2025. № 13(4). С. 29–47. DOI: 10.17431/na/195982
7. Renshaw A., Lourentzou I., LeeShow J., Kim J. Comparing the Spatial Querying Capacity of Large Language Models: OpenAI's ChatGPT and Google's Gemini Pro // *The Professional Geographer*. 2025. № 77(2). С. 1–13. DOI: 10.1080/00330124.2024.2434455
8. Shrijal R. Gemini 2.5 Pro vs. Claude 3.7 Sonnet: Coding Comparison [Електронний ресурс]. 2025. Режим доступу: <https://composio.dev/blog/gemini-2-5-pro-vs-claude-3-7-sonnet-coding-comparison/>
9. Nababuddin M. Calude or Jurassic 2 Foundation Model for Text Generation [Електронний ресурс]. 2024. Режим доступу: <https://medium.com/@nabab.dev/calude-or-jurassic-2-foundation-model-for-text-generation-7eff326c16f2>
10. Ji X. Best alternatives for deepseek R1 vs claude 3.5 sonnet in 2025 [Електронний ресурс]. 2025. Режим доступу: <https://www.byteplus.com/en/topic/384781?title=best-alternatives-for-deepseek-r1-vs-claude-3-5-sonnet-in-2025>
11. Бабчук І. С. Використання Gemini, Chatgpt, Copilot для пошуку інформації про спеціалізовані цифрові мережі // Матеріали V Всеукраїнської студентської наукової конференції «Науковий простір: аналіз, сучасний стан, тренди та перспективи», м. Київ, 17 травня 2024 р. Вінниця : ТОВ «УКРЛОГОС Груп», 2024. С. 341–343.
12. Бабчук І. С. Порівняльний аналіз результатів взаємодії з ChatGPT та Bing AI Chat // Матеріали IV Міжнародної студентської наукової конференції «Концепт науки XXI: стратегії, методи та наукові інструменти», м. Вінниця, 3 листопада 2023 р. Вінниця : ТОВ «УКРЛОГОС Груп», 2023. С. 155–157.

References

1. Zhao, F. F., He, H. J., & Liang, J. J. (2024). Benchmarking the performance of large language models in uveitis: A comparative analysis of ChatGPT-3.5, ChatGPT-4.0, Google Gemini, and Anthropic Claude3. *Eye* (London, England), 39(6). <https://doi.org/10.1038/s41433-024-03545-9>
2. Silhadi, M., Nassrallah, W., Mikhail, D., & Harissi-Dagher, M. (2025). Assessing the performance of Microsoft Copilot, GPT-4 and Google Gemini in ophthalmology. *Canadian Journal of Ophthalmology*. <https://doi.org/10.1016/j.jcjo.2025.01.001>
3. Kokala, A., & Kalluri, K. (2024). Performance benchmarking of generative AI models: ChatGPT-4 vs. Google Gemini AI. *International Research Journal of Modernization in Engineering Technology and Science*, 6(11), 4673–4677. <https://doi.org/10.56726/IRJMETS64283>
4. Bayer, H., Araci, F., & Gürkan, G. (2024). ChatGPT-4o, ChatGPT-4 and Google Gemini are compared with students: A study in higher education. *International Journal of Technology in Education and Science*, 8(4), 627–644. <https://doi.org/10.46328/ijtes.585>
5. Silva, F. (2025). Navigating the dual-edged sword of generative AI in cybersecurity. *Brazilian Journal of Development*, 11(1), 869. <https://doi.org/10.34117/bjdv11n1-062>
6. Jedrzejczak, W., & Kochanek, K. (2025). Comparison of audiological knowledge in Polish language of three chatbots: ChatGPT, Bing Chat and Bard. *Nowa Audiofonologia*, 13(4), 29–47. <https://doi.org/10.17431/na/195982>
7. Renshaw, A., Lourentzou, I., LeeShow, J., & Kim, J. (2025). Comparing the spatial querying capacity of large language models: OpenAI's ChatGPT and Google's Gemini Pro. *The Professional Geographer*, 77(2), 1–13. <https://doi.org/10.1080/00330124.2024.2434455>
8. Shrijal, R. (2025). *Gemini 2.5 Pro vs. Claude 3.7 Sonnet: Coding comparison*. <https://composio.dev/blog/gemini-2-5-pro-vs-claude-3-7-sonnet-coding-comparison/>

9. Nababuddin, M. (2024). *Calude or Jurassic 2 foundation model for text generation*. <https://medium.com/@nabab.dev/calude-or-jurassic-2-foundation-model-for-text-generation-7eff326c16f2>
10. Ji, X. (2025). *Best alternatives for DeepSeek R1 vs Claude 3.5 Sonnet in 2025*. <https://www.byteplus.com/en/topic/384781?title=best-alternatives-for-deepseek-r1-vs-claude-3-5-sonnet-in-2025>
11. Babchuk, I. S. (2024). Vykorystannia Gemini, ChatGPT, Copilot dlia poshuku informatsii pro spetsializovani tsyfrovi merezhi. In *Materialy V Vseukrainskoi studentskoi naukovoï konferentsii «Naukovyi prostir: analiz, suchasnyi stan, trendy ta perspektyvy»* (pp. 341–343). Vinnytsia : TOV «UKRLOHOS Hrup».
12. Babchuk, I. S. (2023). Porivnialnyi analiz rezultativ vzaiemodii z ChatGPT ta Bing AI Chat. In *Materialy IV Mizhnarodnoi studentskoi naukovoï konferentsii «Kontsept nauky XXI: stratehii, metody ta naukovi instrumenty»* (pp. 155–157). Vinnytsia: TOV «UKRLOHOS Hrup».