

Хавроненко В. Д.,

кандидат філософських наук,

доцент кафедри філософії,

Київський Національний університет будівництва і архітектури

Гетун С. Ю.,

аспірант кафедри будівельної механіки,

Київський Національний університет будівництва і архітектури

Кобзар А. Ю.,

доцент кафедри мовної підготовки і комунікації,

Київський Національний університет будівництва і архітектури

ПОРІВНЯЛЬНА ШКАЛА КОГНІТИВНИХ ЗДІБНОСТЕЙ: ВІД РЕПТИЛІЙ ДО СУЧАСНОГО ШІ ТА ЛЮДИНИ

Стрімкий розвиток штучного інтелекту (ШІ), особливо великих мовних моделей, актуалізує питання порівняння його можливостей з біологічним, зокрема людським, інтелектом. Розуміння поточного стану ШІ є критичним для оцінки прогресу до Штучного Загального Інтелекту (ШЗІ), його потенціалу та ризиків [1, 3]. Однак таке порівняння стикається з методологічними викликами через багатогранність поняття «інтелект» та різницю метрик оцінки у тварин, людей та ШІ [2, 5, 10].

Мета та завдання дослідження: проаналізувати когнітивні профілі рептилій, мишей, людей з низьким та високим IQ, та сучасного передового ШІ; розробити концептуальну шкалу складності вирішення завдань для цих сутностей; оцінити позицію та швидкість руху ШІ на цій шкалі; сформулювати прогнози щодо майбутніх можливостей ШІ.

Запропоновано концептуальну якісну шкалу, що впорядковує сутності (рептилія, миша, людина з низьким IQ, людина з середнім IQ, людина з високим IQ, сучасний передовий ШІ станом на 2024-2025 рр.) за максимальною продемонстрованою здатністю вирішувати проблеми зростаючої складності. Ключові виміри шкали:

Складність: кількість кроків міркування, обсяг інформації, рівень невизначеності.

Загальність: широта доменів застосування навичок, адаптація до нових завдань.

Абстрактність: оперування символами, концепціями, метакогнітивне мислення.

Горизонт планування: здатність до цілеспрямованої поведінки на тривалих часових проміжках. Аналіз базується на синтезі наукової літератури з порівняльної психології [10], нейронауки, психометрії, досліджень ШІ (бенчмарки, технічні звіти) [1, 9].

Результати дослідження та їх обговорення. На основі аналізу сутності розміщено на шкалі в такому порядку зростання можливостей:

Рептилія: базове асоціативне навчання, проста просторова орієнтація, обмежена поведінкова гнучкість [10].

Миша: складна просторова когніція, гнучка навігація, але обмежене абстрактне міркування.

Людина з низьким IQ: якісний стрибок завдяки мові та символічному мисленню; вирішення простих конкретних проблем, обмеження у складному міркуванні.

Людина з середнім IQ: типовий рівень; вирішення широкого кола повсякденних завдань, помірна абстракція.

Сучасний передовий ШІ (класу GPT-4o, Claude 3 Opus, Gemini 1.5 Pro): унікальний «нерівний» профіль. Перевершує людей у доступі до інформації, розпізнаванні патернів, продуктивності у специфічних складних завданнях (код, переклад) [9]. Однак фундаментально обмежений у здоровому глузді, каузальному міркуванні [11], надійному узагальненні поза тренувальними даними та адаптивності до якісно нових ситуацій [6, 7, 8].

Людина з високим IQ: висока здатність до абстрактного міркування, складного планування, креативність, метапізнання, гнучка адаптація [4].

Сучасний ШІ демонструє «крихку надінтелектуальність» [6]: потужний, але ненадійний поза межами тренувального розподілу. Рух ШІ до вищих рівнів шкали (ШЗІ) стримується фундаментальними проблемами: відсутність справжнього розуміння світу та здорового глузду, дефіцит каузального міркування, проблеми з узагальненням та адаптивністю, обмеження міркування, ненадійність (галюцинації), відсутність втілення (embodiment). Подолання цих проблем ставить під сумнів достатність простого масштабування поточних архітектур і може вимагати фундаментальних зрушень [6, 7]. Дослідження має обмеження: концептуальність шкали, суб'єктивність оцінок, недоліки бенчмарків [2], швидкий розвиток ШІ. Позиціонування ШІ порушує етичні та соціальні питання: надійність, безпека [12], вплив на зайнятість, упередженість, автономія, дезінформація. Прогнози вказують на значний прогрес ШІ у вузьких сферах (5-10 років), але досягнення ШЗІ залишається невизначеним і залежить від наукових проривів [1, 3].

Запропонована концептуальна шкала візуалізує якісні стрибки у розвитку когнітивних здібностей. Сучасний ШІ посідає унікальну

позицію, демонструючи надлюдські здібності в одних сферах та глибокі обмеження в інших, критичних для загального інтелекту. Його профіль суттєво відрізняється від збалансованого людського інтелекту. Шлях до ШЗІ вимагає подолання фундаментальних наукових викликів. Необхідна розробка кращих методів оцінки ШІ та механізмів для його безпечного й етичного розвитку.

Список використаних джерел:

1. AI Index Steering Committee. The AI Index 2025 Annual Report. Stanford Institute for Human-Centered AI, Stanford University, 2025. URL: https://hai-production.s3.amazonaws.com/files/hai_ai_index_report_2025.pdf (дата звернення: 28.04.2025)
2. Bommasani R. та ін. Can we trust AI benchmarks? An interdisciplinary review of current issues in AI evaluation. arXiv preprint arXiv:2502.06559. 2025. [Preprint]. DOI: 10.48550/arXiv.2502.06559.
3. Grace K., Stewart H. та ін. Thousands of AI authors on the future of AI. arXiv preprint arXiv:2401.02843. 2024. [Preprint]. DOI: 10.48550/arXiv.2401.02843.
4. Haier R. J. The neuroscience of intelligence. Cambridge : Cambridge University Press, 2017. DOI: 10.1017/9781316105771.
5. Leffer L. As AI advances, the meaning of artificial general intelligence remains murky // Science News. URL: <https://www.sciencenews.org/article/artificial-general-intelligence-ai-unclear> (дата звернення: 28.04.2025)
6. Maes S. H. The Trouble with GenAI: LLMs are still not any close to AGI. They will never be. Zenodo. 2024. [Preprint/Report]. DOI: 10.5281/zenodo.14567206.

7. Marcus G. The next decade in AI: Four steps towards robust artificial intelligence. arXiv preprint arXiv:2002.06177. 2020. [Preprint]. DOI: 10.48550/arXiv.2002.06177.
8. Morris M. R., Sohl-Dickstein J., Fiedel N. та ін. Levels of AGI for operationalizing progress on the path to AGI. arXiv preprint arXiv:2311.02462v4. 2024. [Preprint]. DOI: 10.48550/arXiv.2311.02462.
9. OpenAI. GPT-4 Technical Report. arXiv preprint arXiv:2303.08774. 2023. [Preprint]. DOI: 10.48550/arXiv.2303.08774.
10. Schubiger M. N., Fichtel C., Burkart J. M. Validity of cognitive tests for non-human animals: Pitfalls and prospects // Frontiers in Psychology. 2020. Vol. 11. Art. 1835. DOI: 10.3389/fpsyg.2020.01835.
11. Schölkopf B., Locatello F., Bauer S. та ін. Towards causal representation learning. Proceedings of the IEEE. 2021. Vol. 109, № 5. P. 612–634. DOI: 10.1109/JPROC.2021.3058954.
12. UK AI Safety Institute. International AI Safety Report. DSIT 2025/001. 2025. URL: <https://www.gov.uk/government/publications/international-ai-safety-report-2025> (дата звернення: 28.04.2025)