

<https://doi.org/10.5281/zenodo.17534593>

Валін Максим,

Чернівецький національний університет імені Юрія Федьковича,
ORCID: 0009-0002-2966-2887, e-mail: maks21787@gmail.com

Орищук Василь,

доктор філософії у галузі публічного управління та адміністрування,
Київський національний університет імені Тараса Шевченка,
ORCID: 0000-0003-4362-2785, e-mail: voryshchuk@ukr.net

LLM У НАУЦІ: АТРИБУЦІЯ ТА ВІДТВОРЮВАНІСТЬ

АНОТАЦІЯ. Великі мовні моделі (LLM) трансформують дослідження та освіту, дозволяючи створювати нові форми генерації, аналізу та автоматизації текстів. У дослідженні розглядається поточна політика та етичні стандарти щодо використання генеративного штучного інтелекту в наукових дослідженнях, зосереджуючись на атрибуції авторства, прозорості та відтворюваності. Аналізуються політики провідних міжнародних журналів (Science, Nature, Springer Nature, Elsevier) та рекомендації фінансуючих агентств (NIH, NSF), а також рекомендації Міністерства освіти і науки України (2025). Результати дослідження демонструють глобальний консенсус: ШІ не може бути автором, і його використання має бути розкрито. Для забезпечення відтворюваності дослідникам рекомендується надавати набори даних, код, параметри моделі та підказки. Такі ризики, як галюцинації, фальшиві цитування та дрейф моделі, потребують зменшення шляхом відстеження походження, перевірки DOI, водяних знаків та етичного розкриття інформації. Дотримання стандартів COPE, ICMJE, NISO, FAIR/RO-Crate, ORCID та ROR забезпечує підзвітність та чесність у дослідженнях за допомогою ШІ. Зроблено висновок, що відповідальна інтеграція штучного інтелекту в науку та освіту може сприяти інноваціям, зберігаючи при цьому довіру та відтворюваність.

КЛЮЧОВІ СЛОВА: великі мовні моделі, генеративний штучний інтелект, наукові дослідження, авторство, прозорість, відтворюваність, академічна доброчесність, етика штучного інтелекту.

ВСТУП. Великі мовні моделі (LLM) дедалі активніше застосовуються в науці, освіті та публікаціях. Вони дозволяють створювати тексти, аналізувати великі обсяги даних і навіть формулювати гіпотези. Проте використання генеративного ШІ створює ризики для академічної доброчесності — від неправильної атрибуції авторства до проблем із відтворюваністю результатів. Світові наукові організації, журнали та фонди вже запровадили низку політик, спрямованих на забезпечення прозорості, достовірності та етичного використання ШІ. Метою дослідження є узагальнення таких політик і визначення спільних стандартів, що формують основу для відповідального використання LLM у науці.

МАТЕРІАЛИ І МЕТОДИ. Було проведено аналіз редакційних політик журналів Science, Nature, Springer Nature, Elsevier, документів міжнародних комітетів COPE, ICMJE, а також рекомендацій NIH і NSF щодо використання генеративного ШІ. Дослідження включало контент-аналіз офіційних заяв, настанов і меморандумів, а також українського нормативного документа — листа МОН №1/779-23 від 24.04.2025. Основна увага приділялася вимогам до авторства, прозорості, відкритості даних і ризикам, пов'язаним із генерацією контенту за допомогою ШІ.

РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ. Провідні журнали та фонди демонструють схожі підходи до використання ШІ. Science і Nature забороняють вказувати ШІ як співавтора та зобов'язують авторів декларувати будь-яке використання генеративних інструментів. Springer Nature та Elsevier у 2023 році оновили політики, вимагаючи зазначати назву, версію моделі, характер використання та рівень людського контролю. ICMJE наголошує, що ШІ не може відповідати критеріям авторства, тому його внесок може бути лише технічним і має бути описаний у подяках. COPE підтверджує, що відповідальність за точність і доброчесність завжди несе людина.

Грантові агентства також активно реагують: NIH у 2023 році заборонив рецензентам користуватися ChatGPT для оцінювання заявок, а у 2025 — розширив заборону на самих заявників, якщо до створення їхніх текстів великою мірою долучався ШІ. Важливим викликом є відтворюваність результатів, отриманих із застосуванням LLM. Мінімальний набір для перевірки включає: відкриті дані, код, опис конфігурацій, версію моделі та точні промпти, що дозволяють незалежним дослідникам повторити експеримент. Принципи FAIR (Findable, Accessible, Interoperable, Reusable) та формат RO-Crate допомагають забезпечити структуровану подачу даних, кодів і метаданих. Elsevier і Springer рекомендують авторам включати опис використання ШІ у розділ «Методи» та подавати приклади промптів як додатки до статті.

Ризики застосування LLM залишаються значними. Серед найпоширеніших — галюцинації (створення неправдивих фактів), вигадані цитати або джерела, дрейф моделей (зміна поведінки після оновлень), а також порушення конфіденційності. Дослідження 2023 року виявили, що до 30% бібліографій, створених ChatGPT без перевірки, містять вигадані DOI. Такі викривлення здатні знизити достовірність наукових робіт, якщо не застосовуються механізми перевірки. У відповідь пропонуються рішення: DOI-check (автоматична перевірка посилань), provenance-технології для відстеження походження контенту, а також watermarking — цифрове маркування тексту для виявлення вмісту машинного походження.

Українське Міністерство освіти і науки у 2025 році офіційно визначило правила використання ШІ в закладах вищої освіти. МОН дозволяє застосування LLM для навчальних і дослідницьких цілей (генерація завдань, переклади, допомога в підготовці текстів), але категорично забороняє автоматизоване створення курсових, кваліфікаційних і дипломних робіт без розкриття цього факту. Викладачам рекомендовано модифікувати завдання, щоб запобігти неконтрольованому використанню ШІ, зберігаючи роль людини як суб'єкта пізнання. Цей документ узгоджується з позицією ЄС і міжнародними етичними межами.

Міжнародні стандарти — COPE, ICMJE, NISO, FAIR/RO-Crate, ORCID, ROR — утворюють комплексну систему управління науковою доброчесністю. COPE встановлює правила авторства, ICMJE — медичні критерії прозорості, NISO — стандарти метаданих для ШІ, FAIR/RO-Crate — технічні вимоги до відтворюваності, а ORCID і ROR забезпечують ідентифікацію авторів і установ. Сукупно ці елементи формують екосистему довіри, де кожен науковий результат має чітке авторство, дані й цифровий слід.

ВИСНОВКИ. Використання LLM у науці відкриває шлях до нових форм співпраці між людиною й технологією, але потребує дотримання принципів прозорості, відповідальності та відтворюваності. Жодна система ШІ не може бути автором; її роль — допоміжна, а рішення й відповідальність залишаються за дослідником. Міжнародний і український досвід демонструє єдність підходів: відкриті дані, чітка атрибуція, освітня етика та технічні інструменти контролю. Формування культури відповідального використання ШІ — ключ до збереження довіри в науці та освіті.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Nature Editorial Board. (2023). *Tools such as ChatGPT threaten transparent science; here are our ground rules for their use*. Nature, 613(7945), 612. <https://doi.org/10.1038/d41586-023-00191-1>
2. Springer Nature. (2023). Statement on the use of generative artificial intelligence in manuscripts. Springer Nature. <https://group.springernature.com/gp/group/media/press-releases/archive-2023/first-ai-generated-book/26189712>
3. Elsevier. (2023). Principles for use of generative artificial intelligence tools in the publication process. <https://www.elsevier.com/about/policies-and-standards/the-use-of-generative-ai-and-ai-assisted-technologies-in-writing-for-elsevier>

4. COPE Council. (2023). Authorship and AI tools – COPE position statement. Eastleigh: Committee on Publication Ethics (COPE).; 2023.<https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools>
5. International Committee Of Medical Journal Editors. (2024). Recommendations for the conduct, reporting, editing and publication of scholarly work in medical journals. *Ewha medical journal*, 47(4), e48. <https://doi.org/10.12771/emj.2024.e48>
6. National Institutes of Health (NIH), Office of Extramural Research. (2024). Guidance on the use of generative AI in peer review and grant applications.
7. National Science Foundation (NSF), Directorate for Technology, Innovation and Partnerships. (2024). Responsible use of generative artificial intelligence in research and education. <https://www.nsf.gov/focus-areas/ai#:~:text=Share,Find%20funding%20in%20AI>
8. Міністерство освіти і науки України. (2025). Щодо використання систем штучного інтелекту у закладах освіти: лист №1/779-23 24.04.2025. <https://osvita.ua/legislation/list/340/>
9. Wilkinson, M. D., Sansone, S. A., Schultes, E., Doorn, P., Bonino da Silva Santos, L. O., & Dumontier, M. (2024). FAIR principles: Reproducibility and transparency in data and AI research. *Sci Data*. 11(1), 145.
10. Thorp, H. H. (2023). ChatGPT is fun, but not an author. *Science*. 379(6630), 313. <https://doi.org/10.1126/science.adg7879>
11. Resnik, D. B., Hosseini, M., & Rasmussen, L. M. Using AI to write scholarly publications. (2023). *Accountability in Research*, 31(7):1–9. <https://doi.org/10.1080/08989621.2023.2168535>
12. Міністерство освіти і науки України & Міністерство цифрової трансформації України. (2025). Рекомендації щодо відповідального впровадження та використання технологій штучного інтелекту в закладах вищої освіти. <https://www.kmu.gov.ua/news/mintsyfry-ta-mon-razom-z-ekspertamy-rozrobyly-rekomendatsii-shchodo-vidpovidalnoho-vykorystannia-shi-u-vuzakh>

Valin Maxym,

Yuriy Fedkovych Chernivtsi National University,
ORCID: 0009-0002-2966-2887, e-mail: maks21787@gmail.com

Oryshchuk Vasyl,

PhD, Doctor of Philosophy in Public Administration and Management,
Taras Shevchenko National University of Kyiv,
ORCID: 0000-0003-4362-2785, e-mail: voryshchuk@ukr.net

LLM IN SCIENCE: ATTRIBUTION AND REPRODUCIBILITY

ABSTRACT. *Large Language Models (LLMs) are transforming research and education by enabling new forms of text generation, analysis, and automation. This study reviews current policies and ethical standards regarding the use of generative AI in scientific research, focusing on authorship attribution, transparency, and reproducibility. It analyzes leading international journal policies (Science, Nature, Springer Nature, Elsevier) and funding agency guidelines (NIH, NSF), as well as the Ukrainian Ministry of Education and Science (2025) recommendations. The findings show a global consensus: AI cannot be an author, and its use must be disclosed. To ensure reproducibility, researchers are encouraged to provide datasets, code, model parameters, and prompts. Risks such as hallucinations, fake citations, and model drift require mitigation via provenance tracking, DOI verification, watermarking, and ethical disclosure. Adherence to COPE, ICMJE, NISO, FAIR/RO-Crate, ORCID, and ROR standards ensures accountability and integrity in AI-assisted research. The paper concludes that responsible integration of AI into science and education can enhance innovation while preserving trust and reproducibility.*

KEYWORDS: *large language models, generative artificial intelligence, scientific research, authorship, transparency, reproducibility, academic integrity, AI ethics.*