

всеосяжної безпеки є єдиним шляхом подолання деструктивних наслідків війни та зміцнення інтелектуальної конкурентоспроможності держави.

Література

1. Жижко Т.А. Деформація персоналу: виклики системи управління персоналом закладів освіти в умовах воєнного стану. 2024.
2. Залознова Ю.С., Азьмук Н.А. Людський капітал України в умовах війни: втрати та здобутки. *Економіка та суспільство* . 2022. № 38. DOI: 10.32782/2524-0072/2022-38-59.
3. Корнілова О.В. Антикризове управління людським капіталом у період воєнних викликів. *Матеріали Міжнар. наук.-практ. конф.* . Одеса, 2026.
4. Сідлецький С.І., Шандар А.М. Вища освіта в Україні в умовах воєнного стану: виклики та пріоритети. *Економіка та суспільство* . 2024. № 64. DOI: 10.32782/2524-0072/2024-64-151.

DOI <https://doi.org/10.36059/978-966-397-651-8-74>

ДИСКУРСИВНИЙ ПСЕВДОАВТОР: ВИКЛИКИ НАУКОВОГО КЕРІВНИЦТВА В ЕПОХУ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ

Ягодзінський Сергій Миколайович
*доктор філософських наук, професор,
проректор з навчально-методичної роботи,
ПВНЗ «Європейський університет»,
м. Київ, Україна
ORCID: 0000-0001-8755-2235*

Ягодзінська Валерія Євгенівна
*аспірантка
Київський національний університет імені Тараса Шевченка
м. Київ, Україна
ORCID: 0000-0002-0015-7767*

Проблема. Генеративні мовні моделі по-новому ставлять питання про те, хто говорить у тексті. Але на відміну від підходів структуралістів чи навіть представників постмодерну [1; 2], поширення штучного інтелекту змушує переформулювати як питання, так і методологію пошуку відповіді на нього.

Читаючи текст, ми схильні впізнавати у ньому голос (автора, позицію, думку, ставлення тощо). Цей голос щось стверджує, визнає межі власної обізнаності, змінює тональність, регістр, емоцію, рівень аргументованості. Ефект присутності когось у тексті та за ним є первинною даністю читання. Але за текстом, згенерованим нейронними мережами, не передбачається суб'єкта, якому можна приписати свідомість, намір, знання чи відповідальність у людському розумінні [3]. Наявні категорії описують цю ситуацію по-різному, але жодна з концепцій не стверджує і не заперечує, що феноменологічний та епістемологічний рівні згенерованого тексту не просто співіснують, але й конституують один одного. Саме ця невідповідність складає проблему даного дослідження.

Завдання. Дана публікація має на меті поставити такі завдання: (1) показати, що всі наявні категорії суб'єкта мовлення в AI-тексті асиметричні і розмежують феноменологічний та епістемологічний рівні; (2) обґрунтувати гіпотезу взаємної конституції зазначених рівнів; (3) запропонувати поняття, яке цю конституцію зафіксує; (4) окреслити практичні наслідки приписування авторства й детекції AI-текстів [4; 5].

Пропозиції. Ми обґрунтовуємо гіпотезу взаємної конституції: ефект присутності суб'єкта в AI-тексті є функцією епістемічної форми цього тексту, а не її передумовою. Під епістемічною формою ми розуміємо набір структурних ознак, за якими текст упізнається як претензія на знання: несуперечливість, зв'язність аргументації, збереження референції, розмежування впевненого і гіпотетичного.

Для позначення суб'єкта, що існує винятково як функція епістемічної форми, ми запроваджуємо поняття дискурсивного псевдоавтора. Це суб'єкт мовлення, який конститується винятково епістемічною формою тексту, не має носія поза текстом, не зводиться до психологічної ілюзії читача, оскільки виникає як стійкий ефект епістемічної організації тексту.

Головним результатом вважаємо не сам концепт, а перевірюване передбачення, яке з нього випливає: цілеспрямоване руйнування епістемічної форми знищує ефект присутності суб'єкта в AI-тексті, тоді як у тексті людського автора присутність зберігається і перетлумачується читачем як іронія, помилка або патологія. Це передбачення фальсифіковане: його можна спростувати серією досліджень із контрольованим руйнуванням форми текстів обох типів та подальшим оцінюванням читацького сприйняття.

Практичним наслідком запропонованої концепції є можливість перегляду моделі наукового керівництва аспірантами (загалом – кваліфікаційними роботами випускників та науковими роботами здобувачів і вчених). В умовах генеративного ШІ текст більше не може автоматично розглядатися як достатній доказ того, що здобувач сформував відповідне судження, розуміє аргументацію та здатний її захистити [6]. Тому *контроль* доцільно переносити з *питання «хто*

написав текст?» на питання «чи існує за цим текстом дослідник, здатний відтворити, обґрунтувати й переглянути власне судження?». Для цього пропонуємо чотири підходи.

1. Процесуальний підхід: фіксувати судження до генерації. Критичним є не сам факт використання генеративної моделі, а послідовність дослідницьких операцій. *До звернення до ШІ аспірант має письмово зафіксувати власну тезу, основний аргумент, дослідницьке питання або попередню інтерпретацію матеріалу.* Лише після цього модель може використовуватися для структурування, пошуку альтернативних формулювань, критики чи підготовки чернетки. Такий порядок не усуває ШІ з дослідницького процесу, а зміщує його після формування первинного судження, зменшуючи ризик когнітивного аутсорсингу [7], коли текст виникає раніше за власну позицію дослідника і веде дослідника, а не навпаки.

2. Діалогічний підхід: замінити контроль тексту перевіркою здатності його захищати. Замість одного підсумкового захисту доцільно запроваджувати короткі мікрозахисти окремих підрозділів роботи. Здобувач має впродовж 10–15 хвилин пояснити написаний текст, назвати головну тезу, показати логіку переходу до висновку та відповісти на контраргументи. Особливо продуктивними є питання, що навмисно «ламають» готову (згенеровану) форму: «Чому?», «Що могло б спростувати цей висновок?», «З яким автором Ви тут не погоджуєтесь?», «Як зміниться висновок, якщо прибрати цей аргумент?». Перевіряється не пам'ять про текст, а здатність дослідника продовжувати міркування поза ним і без нього.

3. Верифікаційний підхід: оцінювати не лише прийняте, а й відхилене. Корисним інструментом може стати «журнал відхиленого», у якому аспірант коротко фіксує не промпти до генеративної моделі, а запропоновані нею тези, джерела, аргументи чи формулювання, які він вирішив не використовувати, із поясненням причин. Здатність погодитися з переконливим текстом майже нічого не говорить про сформованість дослідницького судження. *Здатність аргументовано відкинути його є значно сильнішим індикатором.* Аналогічно кожне використане джерело має бути особисто відкритим, прочитаним і співвіднесеним здобувачем із власним аргументом, а не просто перенесеним із машинної відповіді.

4. Просектувальний підхід: створювати завдання, для яких готової генерації недостатньо. Найнадійнішою стратегією є не виявлення використання ШІ після написання тексту, а така організація дослідження, за якої машинна генерація не може замінити дослідницьку роботу [8]. До таких завдань належать збір емпіричних даних, робота з локальними або неоцифрованими джерелами, побудова авторської типології, порівняння суперечливих позицій у літературі, аналіз матеріалу конкретної установи чи спільноти. Доцільним є також

«питання про дельту»: що саме в цьому підрозділі створено дослідником і не могло бути отримане простим запитом до генеративної моделі? Такою дельтою можуть бути нові дані, авторська класифікація, виявлена суперечність, новий зв'язок між відомими положеннями або аргументований вибір між альтернативними поясненнями.

Отже, *ефективне наукове керівництво в умовах генеративного ШІ* має зміщуватися **від контролю походження тексту до контролю походження судження**. Завдання наукового керівника полягає не в тому, щоб домогтися тексту «без ШІ», а в тому, щоб забезпечити існування дослідника поза текстом – суб'єкта, який може пояснити, захистити, змінити або відкинути сформульоване положення і взяти на себе відповідальність за нього.

Література

1. Barthes R. The Death of the Author. Aspen. 1967. No. 5-6. P. 2-6.
2. Foucault M. Qu'est-ce qu'un auteur? Bulletin de la Société française de Philosophie. 1969. Vol. 63, No. 3. P. 73-104.
3. 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). New York: Association for Computing Machinery, 2021. P. 610-623. DOI: 10.1145/3442188.3445922.
4. Levy N. Responsibility is not required for authorship. Journal of Medical Ethics. 2025. Vol. 51, No. 4. P. 230-232. DOI: 10.1136/jme-2024-109912.
5. Helgesson G., Bülow W. Responsibility is an adequate requirement for authorship: a reply to Levy. Journal of Medical Ethics. 2025. Vol. 51, No. 4. P. 295-296. DOI: 10.1136/jme-2024-110245.
6. Киричук Б. С., Гришко В. І. Академічна доброчесність і штучний інтелект: подолання викликів у освітньо-науковій діяльності України та зарубіжних держав. Науковий вісник Ужгородського національного університету. Серія: Право. 2024. Т. 2, № 85. DOI: 10.24144/2307-3322.2024.85.2.45.
7. Gerlich M. AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking. Societies. 2025. Vol. 15, No. 1. Art. 6. DOI: 10.3390/soc15010006.
8. Паламар С., Науменко М. Штучний інтелект в освіті: використання без порушення принципів академічної чесності. Освітологічний дискурс. 2024. № 1(44). С. 68-83. DOI: 10.28925/2312-5829.2024.15.