

Відомості Верховної Ради України. 2026. № 6. Ст. 42.

3. Про затвердження плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2025–2026 роки : розпорядження Кабінету Міністрів України від 09.05.2025 № 457-р // Урядовий кур'єр. 2025. № 92.

4. Берназюк Я. Проєкт Закону України «Про штучний інтелект»: концептуальні та правові підходи. Constitutionalist. 2026. URL: <http://constitutionalist.com.ua> (дата звернення: 04.06.2026).

5. Харитонов Є. О., Харитонova О. І. Цивільно-правове регулювання відносин у сфері використання штучного інтелекту: виклики та перспективи 2025–2026 рр. Часопис цивілістики. 2025. Т. 54. С. 12–19.

Козак Валентина Андріївна,

кандидат філологічних наук, доцент,

Київський інститут Національної гвардії України

ORCID: 0009-0000-0121-4949

ШТУЧНИЙ ІНТЕЛЕКТ ЯК ВИКЛИК АКАДЕМІЧНИЙ ДОБРОЧЕСНОСТІ: РИЗИКИ «ЦИФРОВОГО ПЛАГІАТУ» У НАУКОВИХ РОБОТАХ ЗДОБУВАЧІВ ВИЩОЇ ОСВІТИ

Генеративний штучний інтелект – зокрема великі мовні моделі типу ChatGPT – стрімко увійшов в освітній простір і змінив усталені уявлення про авторство наукового тексту. Інструменти, покликані полегшити дослідницьку роботу, на практиці розмили межу між власним інтелектуальним зусиллям та машинною генерацією. Це ставить під загрозу не лише якість навчальних робіт здобувачів вищої освіти, а й принципи академічної доброчесності.

Університетська спільнота зіткнулася з явищем «цифрового плагіату» – прихованого академічного шахрайства, за якого здобувач вищої освіти передає інтелектуальну функцію алгоритму, не осмислюючи й не перевіряючи результат. Без сформованих навичок наукового письма та критичного мислення здобувачі вищої освіти легко потрапляють у пастку ілюзорної простоти генерованого тексту.

Широкий контекст етичних та безпекових викликів, пов'язаних із ШІ, сьогодні активно порушується у працях багатьох українських та зарубіжних науковців. Зокрема, стаття О. Петасюк «Штучний інтелект: контекст та осмислення» досліджує феномен ШІ в цифрову епоху, наголошуючи на його перевагах у медицині, науці, освіті та військовій сфері. Авторка переконує, що

страхи щодо «повстання машин» чи знищення людства є спекулятивними та передчасними, оскільки наразі вся машинна творчість керується людським кодом, а дії нейромереж залежать від ступеня інтеграції в суспільне життя [2]. Розвиток таких технологій, як ChatGPT чи машинне навчання, розглядається не як загроза, що спричинить безробіття, а як інструмент, який стимулює людську творчість і стає логічним еволюційним продовженням людського інтелекту.

Дослідження Ю. Ваврика та І. Опірського «Штучний інтелект: кібербезпека нового покоління» зосереджується на подвійній ролі ШІ у сучасному цифровому захисті. З одного боку, ШІ революціонізує традиційні системи кібербезпеки, дозволяючи ефективно аналізувати величезні масиви даних, швидко виявляти аномалії, прогнозувати загрози та автоматизувати реагування на інциденти [1]. З іншого боку, стаття звертає увагу на «темний бік» ШІ – його активне використання зловмисниками для створення шкідливого ПЗ, реалістичного фішингу, дипфейків та автоматизації кібератак, що вимагає комплексного і постійного вдосконалення систем захисту.

Окреслені дослідниками етичні та алгоритмічні обмеження ШІ яскраво виявляються в освітньому процесі. Якщо традиційний плагіат передбачає привласнення чужого тексту, то «цифровий плагіат» має іншу природу: автор не копіює роботу іншої людини, а використовує генеративну модель для створення нового тексту. Формально текстових збігів немає, але зникає головне – власна думка, аналіз і персональна відповідальність за зміст. Це прихована форма академічної нечесності, яку традиційні антиплагіатні системи не виявляють.

Генеративні моделі видають стилістично переконливі відповіді: граматично правильні, структуровані, насичені науковою лексикою. Саме тому здобувач вищої освіти, отримавши такий текст, щиро сприймає його як якісне дослідження. Однак через брак власного досвіду він не помічає логічних суперечностей, поверхневих висновків і фактичних помилок. Зрештою, він підписує роботу, суті якої сам до кінця не розуміє, і перетворюється з дослідника на оператора ШІ-інструменту. Це веде до руйнування дослідницьких навичок ще на етапі їх формування.

Найнебезпечніша пастка для початківців – беззастережна довіра до бібліографічних списків, згенерованих нейромережею. Великі мовні моделі не оперують реальними фактами, а відтворюють імовірні послідовності слів. Наслідок – так звані «галюцинації»: переконливо оформлені, але цілком вигадані назви статей, прізвища науковців і вихідні дані журналів. Здобувач вищої освіти, який не перевіряє кожне джерело, вносить фікцію до списку літератури й руйнує наукову достовірність роботи.

Генеративні моделі орієнтовані на виважений, нейтральний виклад. Передаючи написання розділів нейромережі, дослідник отримує текст без дослідницької позиції: критичний аналіз підмінено загальновідомими констатаціями, авторське «Я» зникає, а наукова новизна розчиняється у кліше. Робота втрачає ознаки дослідження і стає рефератною компіляцією.

Традиційні антиплагіатні системи орієнтовані на пошук текстових збігів і не виявляють генерованих парафразів. Спеціалізовані детектори ШІ-тексту також ненадійні: їхні алгоритми легко обійти завдяки мінімальному ручному редагуванню. Грамотно структурований текст більше не є свідченням компетентності здобувача вищої освіти. Це вимагає перегляду самої моделі оцінювання: акцент має перейти з готового продукту на процес – на динаміку роботи, усні захисти, дискусію з науковим керівником.

Генеративний штучний інтелект не зникне з освітнього простору – і заборона тут не є ані дієвим, ані доцільним інструментом. Проблема не в самій технології, а в тому, як її використовують. Здобувач вищої освіти, який делегує алгоритму не лише форматування, а й мислення, позбавляє себе головного – досвіду наукового пошуку.

Відповідь на виклик «цифрового плагіату» має бути системною. По-перше, необхідно чітко визначити допустимі межі використання ШІ та запровадити обов'язкове маркування ШІ-згенерованого контенту. По-друге, слід переосмислити саму модель оцінювання: перенести акцент з письмового продукту на живу наукову дискусію – усні захисти, поетапні консультації, аргументований захист власної позиції. По-третє, з перших тижнів навчання необхідно формувати у здобувачів навички інформаційної грамотності: вміння перевіряти джерела, розпізнавати галюцинації нейромереж і критично оцінювати будь-який контент – незалежно від його походження.

Таким чином, ключовим завданням закладу вищої освіти є не обмеження доступу до інструментів ШІ, а формування у здобувачів критичного ставлення до їх результатів. Дослідник, який розуміє принципи роботи генеративних моделей і їхні обмеження, здатен використовувати ці інструменти продуктивно – без шкоди для якості наукової роботи та власного професійного розвитку.

Список використаних джерел:

1. Ваврик Ю., Опірський І. Штучний інтелект: кібербезпека нового покоління. *Ukrainian Scientific Journal of Information Security*. 2024. Т. 30, № 2. С. 244–255. URL: <https://doi.org/10.18372/2225-5036.30.19235> (дата звернення: 08.06.2026).