

УДК 37:37.091.12:004

Вадим Кушнір, доктор філософії,
науковий співробітник відділу
цифрових освітніх ресурсів,
Інститут професійної освіти НАПН України,
м. Київ, Україна

ЗАПОБІГАННЯ ЗЛОВЖИВАННЮ ТА НЕКОРЕКТНОМУ ВИКОРИСТАННЮ ТЕХНОЛОГІЙ ШІ

Анотація. У матеріалах доповіді розглянуто проблему запобігання зловживанню та некоректному використанню технологій штучного інтелекту в сучасному освітньому середовищі. Проаналізовано ризики, пов'язані з використанням генеративних систем штучного інтелекту, зокрема порушення принципів академічної доброчесності, поширення недостовірного контенту, створення маніпулятивних цифрових матеріалів та використання штучного інтелекту без належного людського контролю. Акцентовано увагу на сучасних підходах Європейського Союзу щодо забезпечення прозорості застосування технологій штучного інтелекту, а також на необхідності формування цифрової відповідальності учасників освітнього процесу. Встановлено, що ефективне використання штучного інтелекту потребує поєднання технологічних, правових та педагогічних механізмів регулювання.

Ключові слова: штучний інтелект, генеративний штучний інтелект, цифрова безпека, академічна доброчесність, цифрова відповідальність, освітнє середовище.

Abstract. The paper examines the problem of preventing abuse and improper use of artificial intelligence technologies in the modern educational environment. Risks associated with the use of generative artificial intelligence systems are analyzed, including violations of academic integrity principles, dissemination of misleading content, creation of manipulative digital materials, and the use of artificial intelligence without adequate human oversight. Particular attention is given to contemporary European Union approaches aimed at ensuring transparency in the use of artificial intelligence technologies and to the importance of developing digital responsibility among participants in the educational process. It has been established that the effective use of artificial intelligence requires the integration of technological, legal, and pedagogical regulatory mechanisms.

Keywords: artificial intelligence, generative artificial intelligence, digital security, academic integrity, digital responsibility, educational environment.

Стрімке поширення технологій штучного інтелекту в освітньому середовищі сприяє трансформації традиційних підходів до навчання, створення навчального контенту, автоматизації оцінювання та підтримки індивідуальних освітніх траєкторій. Разом із позитивними можливостями виникають ризики їх некоректного використання та зловживання, що актуалізує проблему формування механізмів запобігання негативним наслідкам застосування таких технологій.

У сучасних наукових дослідженнях поняття «технологія штучного інтелекту» розглядається як система інтелектуальних агентів, здатних сприймати інформацію із навколишнього середовища, аналізувати її та здійснювати відповідні дії для досягнення поставлених цілей. Таке трактування запропоновано С. Расселом та П. Норвігом, праця яких вважається одним із найбільш авторитетних досліджень у галузі штучного інтелекту [1]. У сучасних умовах технології штучного інтелекту охоплюють алгоритми машинного навчання, нейронні мережі, генеративні моделі, системи обробки природної мови та інші цифрові інструменти, що використовуються в освітній діяльності.

Водночас використання штучного інтелекту в освітньому процесі супроводжується виникненням низки ризиків. Одним із найбільш поширених проявів некоректного використання є порушення принципів академічної доброчесності. Генеративні системи штучного інтелекту здатні автоматично створювати тексти, програмний код, презентації та інші навчальні матеріали, що може спричиняти заміну самостійної пізнавальної діяльності здобувачів освіти автоматизованою генерацією результатів. За таких умов виникає небезпека формального засвоєння знань та зниження якості освітньої діяльності [2].

Для систематизації основних проявів свідомого використання технологій штучного інтелекту з порушенням принципів доброчесності та безпечного цифрового середовища доцільно виокремити найбільш поширені види зловживань (табл. 1).

Наведені у табл. 1 ризики характеризуються свідомим використанням можливостей штучного інтелекту з метою отримання неправомірних переваг або створення негативного впливу на освітнє середовище.

Таблиця 1. Прояви зловживання технологіями штучного інтелекту в освітньому середовищі

Вид зловживання	Характеристика	Потенційні наслідки
Академічне шахрайство	Автоматизоване створення навчальних робіт без самостійного опрацювання матеріалу	Зниження якості знань, порушення принципів академічної доброчесності
Генерування неправдивого контенту	Створення фальсифікованих текстів, зображень або відповідей	Поширення дезінформації
Створення deepfake-матеріалів	Формування штучних аудіо-та відеоматеріалів	Маніпуляція свідомістю користувачів
Несанкціоноване використання персональних даних	Завантаження конфіденційної інформації у системи ШІ	Порушення конфіденційності та ризики витоку даних
Використання ШІ для маніпуляції рішеннями	Формування упереджених рекомендацій або інформаційного впливу	Спотворення результатів освітньої діяльності

Окрему групу ризиків становить створення недостовірного або маніпулятивного цифрового контенту. Генеративні моделі можуть продукувати неправдиву інформацію, помилкові твердження або штучно створені зображення й відеоматеріали, які зовні мають ознаки достовірності. Особливу небезпеку становлять технології deepfake, що створюють штучно згенеровані матеріали із високим рівнем реалістичності та можуть використовуватися для маніпуляції суспільною думкою чи введення користувачів в оману [3].

Поряд із навмисними порушеннями суттєвий вплив на якість освітнього процесу мають випадки некоректного використання технологій штучного

інтелекту, які переважно пов'язані з недостатнім рівнем цифрової грамотності та відсутністю критичного оцінювання результатів роботи таких систем (табл. 2).

На відміну від зловживань, наведені у таблиці 2 випадки переважно пов'язані не з умисними діями користувачів, а з недостатнім рівнем цифрової грамотності, критичного мислення та відсутністю належного контролю за використанням технологій штучного інтелекту.

Таблиця 2. Прояви некоректного використання технологій штучного інтелекту в освітньому середовищі

Вид некоректного використання	Характеристика	Потенційні наслідки
Некритичне використання результатів ШІ	Використання сформованої відповіді без перевірки її достовірності	Поява помилок у навчальних матеріалах
Надмірна залежність від ШІ	Заміна власного аналізу автоматично сформованими рішеннями	Зниження рівня критичного мислення
Відсутність перевірки джерел	Використання інформації без встановлення її походження	Використання недостовірних даних
Неправильне формулювання запитів	Створення нечітких або помилкових інструкцій для систем ШІ	Отримання некоректних результатів
Використання ШІ без людського контролю	Повне делегування процесу прийняття рішень цифровій системі	Помилкові управлінські та освітні рішення

Для наочного представлення найбільш поширених проявів зловживання та некоректного використання технологій штучного інтелекту в освітньому середовищі було здійснено їх узагальнення та систематизацію (рис. 1).

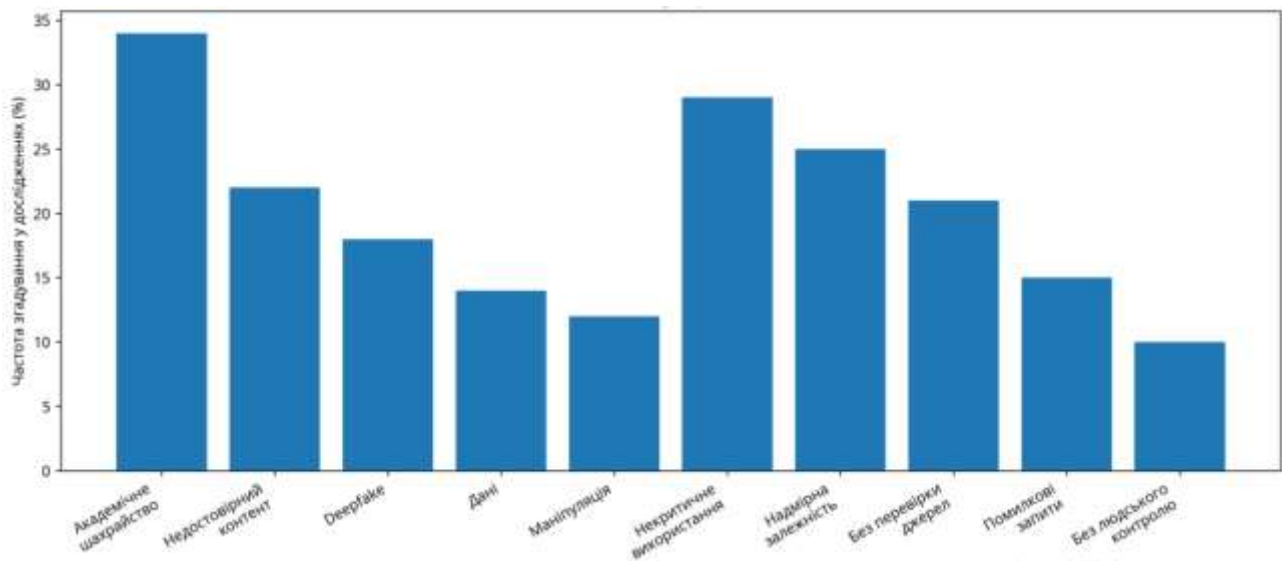


Рис. 1. Порівняльна характеристика проявів зловживання та некоректного використання технологій штучного інтелекту в освітньому середовищі

Примітка. Розроблено автором.

Як свідчать дані, наведені на рис. 1, серед проявів зловживання технологіями штучного інтелекту найбільш поширеним є використання генеративних систем для порушення принципів академічної доброчесності. Водночас серед випадків некоректного використання домінують некритичне використання результатів роботи систем штучного інтелекту та надмірна залежність від автоматизованих рішень. Це підтверджує необхідність формування цифрової культури, розвитку критичного мислення та посилення механізмів відповідального використання технологій штучного інтелекту.

Запобігання таким ризикам потребує комплексного підходу. Одним із пріоритетних напрямів виступає формування цифрової грамотності та культури відповідального використання технологій штучного інтелекту. Важливого значення набуває розвиток здатності оцінювати достовірність цифрової інформації, визначати ознаки автоматично створеного контенту та здійснювати критичний аналіз результатів роботи систем штучного інтелекту.

Важливим механізмом регулювання також виступають сучасні нормативні підходи Європейського Союзу. Оновлені рекомендації щодо застосування

положень Європейського акту про штучний інтелект (AI Act) передбачають забезпечення прозорості функціонування систем штучного інтелекту, обов'язкове позначення штучно створеного контенту, здійснення людського контролю над автоматизованими рішеннями та оцінювання потенційних ризиків використання таких технологій [4]. Реалізація зазначених підходів у сфері освіти сприятиме формуванню безпечного цифрового освітнього середовища.

Таким чином, запобігання зловживанню та некоректному використанню технологій штучного інтелекту в освітньому середовищі доцільно розглядати як сукупність правових, технологічних та педагогічних заходів, спрямованих на забезпечення академічної доброчесності, інформаційної безпеки та відповідального використання цифрових інструментів.

Список використаних джерел

1. Russell S., Norvig P. *Artificial Intelligence: A Modern Approach*. 3rd ed. Upper Saddle River : Pearson Education, 2010. 1132 p. URL: [Artificial Intelligence: A Modern Approach PDF](#) (дата звернення 19.05.2026).
2. Swanson J. *Perspectives on Artificial Intelligence in Education: A Conversation Between an Education Futurist and an Education Strategist* [Електронний ресурс]. KnowledgeWorks. 2024. URL: [KnowledgeWorks: Perspectives on Artificial Intelligence in Education](#) (дата звернення 19.05.2026).
3. Westerlund M. The Emergence of Deepfake Technology: A Review. *Technology Innovation Management Review*. 2019. Vol. 9 (11). P. 39–52.
4. European Commission. *Draft Guidelines on the Implementation of the Transparency Obligations for Certain AI Systems under Article 50 of the AI Act* [Електронний ресурс]. Brussels : European Commission, 2026. URL: [Draft Guidelines on the Implementation of the Transparency Obligations for Certain AI Systems under Article 50 of the AI Act](#) (дата звернення 19.05.2026).