

Електронний журнал «Ефективна економіка» включено до переліку наукових фахових видань України з питань економіки (Категорія «Б», Наказ Міністерства освіти і науки України № 975 від 11.07.2019). Спеціальності – 051, 071, 072, 073, 075, 076, 292.
Ефективна економіка. 2026. № 3.
ISSN 2307-2105



Copyright © The Author(s). This is an open access article distributed under the terms of the Creative Commons Attribution License 4.0 (<https://creativecommons.org/licenses/by/4.0/>).

DOI: <http://doi.org/10.32702/2307-2105.2026.3.106>

УДК 005:004.8:378.147

Л. П. Альошкіна,

к. е. н, доцент, доцент кафедри менеджменту,

Уманський національний університет

ORCID ID: <https://orcid.org/0000-0002-1647-0141>

**ВІДПОВІДАЛЬНЕ ВИКОРИСТАННЯ ГЕНЕРАТИВНОГО
ШТУЧНОГО ІНТЕЛЕКТУ ПРИ ПІДГОТОВЦІ ЗДОБУВАЧІВ ВИЩОЇ
ОСВІТИ ЗІ СПЕЦІАЛЬНОСТІ “МЕНЕДЖМЕНТ”: КЕЙС-
ОРІЄНТОВАНА МОДЕЛЬ РОЗВИТКУ АКАДЕМІЧНОЇ
ДОБРОЧЕСНОСТІ**

L. Alioshkina

*PhD in Economics, Associate Professor, Associate Professor of the Department of
Management, Uman National University*

**RESPONSIBLE USE OF GENERATIVE ARTIFICIAL INTELLIGENCE IN
THE TRAINING OF HIGHER EDUCATION STUDENTS MAJORING IN
“MANAGEMENT”: A CASE-BASED MODEL FOR FOSTERING
ACADEMIC INTEGRITY**

У статті досліджено сутність категорії відповідального використання генеративного штучного інтелекту при підготовці здобувачів вищої освіти зі спеціальності «Менеджмент». Визначено ризики академічної доброчесності за умов застосування GenAI у навчально-дослідній діяльності, зокрема приховану підміну авторського внеску, некоректне цитування, використання вигаданих посилань і неперевіраних тверджень. Представлено кейс-орієнтовану модель формування культури академічної доброчесності, що інтегрує контрольні точки прозорості, верифікації та фіксації авторського внеску у процес виконання завдань. Модель операціоналізовано через алгоритм виконання кейсу, аналітичну рубрику оцінювання та пакет типових завдань. Узагальнення результатів педагогічної апробації засвідчило позитивні зміни у якості студентських робіт, рівнях верифікації, прозорості та відповідальності використання GenAI. Одержані результати підтверджують доцільність застосування моделі при викладанні навчальної дисципліни «Основи наукових досліджень та академічна доброчесність».

The article examines the essence of the category of responsible use of generative artificial intelligence in the training of higher education students majoring in Management and clarifies its significance for ensuring academic integrity in contemporary university education. It is established that the expansion of GenAI in the educational process creates not only new opportunities for information processing, language support, and analytical assistance, but also substantial risks, including concealed substitution of authorship, fabricated or unverifiable references, inaccurate citation, insufficient source verification, and reduced transparency of the student's actual contribution.

To address these challenges, the article presents a case-based model for fostering a culture of academic integrity in the teaching of the course "Fundamentals of Scientific Research and Academic Integrity." The model is structured as a procedurally organized cycle of case completion with embedded

control points related to transparency of GenAI use, verification of factual claims, correctness of citation, and explicit fixation of individual authorship contribution. Its practical implementation is ensured through a standardized case completion protocol, an analytical assessment rubric, and a package of typical case tasks.

The pedagogical approbation conducted in two first-year academic groups of Management students demonstrated positive changes in the quality and integrity of students' work. In particular, improvements were observed in the verification of statements, the use of authentic sources, the transparency of GenAI-assisted work, and the students' ability to distinguish between acceptable language editing and impermissible substitution of authorship. The findings confirm the expediency of applying the proposed model in Management education and indicate its methodological potential for further adaptation in higher education practice.

Ключові слова: *штучний інтелект; генеративний ІІІ; академічна доброчесність; кейс-метод; менеджмент-освіта; наукові дослідження; рубрика оцінювання; прозорість; верифікація.*

Keywords: *artificial intelligence; generative AI; academic integrity; case method; management education; research; assessment rubric; transparency; verification.*

Постановка проблеми у загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Стрімке поширення генеративного штучного інтелекту (GenAI) у вищій освіті радикально змінює навчально-дослідну діяльність здобувачів вищої освіти: від пошуку й систематизації літератури до підготовки текстів, презентацій і первинної інтерпретації даних, перефразування, генерування прикладів та гіпотез. В освітньому вимірі це відкриває можливості для персоналізації підтримки навчання, підвищення доступності пояснень і розвитку навичок письма за умови коректної постановки цілей і процедур використання. Поряд із перевагами зростають ризики когнітивного “аутсорсингу”, зниження

критичного мислення та появи помилок, що виглядають правдоподібно, але не мають перевірної основи. Не менш актуальним є порушення академічної доброчесності, зокрема прихована підміна авторського внеску, некоректне цитування та використання вигаданих посилань, некритичне відтворення помилкових тверджень («галюцинацій») і спрощення дослідницької логіки до компіляції.

До прикладу, в контексті вивчення навчальної дисципліни «Основи наукових досліджень та академічна доброчесність», викладання якої згідно освітньо-професійної програми «Менеджмент» здійснюється для здобувачів вищої освіти спеціальності «Менеджмент» в Уманському національному університеті, ці виклики є методично критичними. Саме перший семестр першого року навчання формує базові патерни дослідницької поведінки: постановку проблеми, роботу з першоджерелами, коректне цитування, логіку “мета – методи – результати – висновки”, а також етичну відповідальність за зміст роботи. За відсутності процедурних рамок GenAI може “підмінити” ключові дослідницькі операції: замість читання першоджерел студент використовує згенерований огляд, замість верифікації фактів – довіряє правдоподібним твердженням, замість власних висновків – приймає синтетичну відповідь моделі. Це загострює проблему академічної доброчесності, яка в українському законодавстві визначається як сукупність етичних принципів і правил, що забезпечують довіру до результатів навчання та наукових досягнень[17-19].

Аналіз останніх досліджень і публікацій. Сучасна дискусія щодо використання генеративного штучного інтелекту (GenAI) в освітньому середовищі та наукових дослідженнях розгортається у взаємопов’язаних площинах: глобальні політики та рекомендації, видавничо-етичні вимоги до прозорості, а також педагогічні підходи й інструменти оцінювання, здатні перевести принципи доброчесності в практику. У цьому контексті важливо не лише констатувати потенціал GenAI для підтримки дослідницьких дій здобувачів вищої освіти, а й визначити межі допустимого використання та

процедури контролю якості, що забезпечують довіру до результатів навчально-наукової роботи.

У сучасних дослідженнях GenAI в освітньому навчальному середовищі переважає позиція “подвійного ефекту”: з одного боку – підтримка навчання (структурування, зворотний зв’язок, мовна допомога), з іншого – зростання ризиків для академічної чесності та якості мислення. Зокрема, в аналітичних оглядах підкреслюється, що великі мовні моделі можуть бути корисними як інструменти навчального супроводу, однак потребують чітких правил застосування, оскільки схильні генерувати переконливі, але потенційно помилкові твердження та можуть спонукати до заміщення власної інтелектуальної роботи студента[1,6,7]. У цьому значенні “якість тексту” не дорівнює “якості знання”: мовна гладкість може маскувати відсутність доказової бази, що робить особливо актуальними практики верифікації та джерельної дисципліни[7,1].

Проблематика академічної доброчесності в умовах використання GenAI найбільш предметно розкрита в роботах, що аналізують зміну структури порушень і трансформацію педагогічного контролю. Дослідники підкреслюють, що фокус закладів вищої освіти має зміщуватися від “виявлення використання” до “перепроєктування оцінювання” та навчання відповідальному використанню, оскільки універсальні детектори не гарантують надійності, а політики різняться між інституціями [3]. Для здобувачів вищої освіти першого року навчання особливо важливо формувати розуміння меж допустимої допомоги (мовне редагування, структуризація) та недопустимої підміни (заміна авторства, фабрикація джерел/даних), адже саме на ранніх етапах закріплюються навчальні стратегії.

Окремий блок становлять міжнародні політики наукової комунікації, що задають етичні стандарти використання GenAI у дослідженнях і публікаціях. COPE у позиційному документі наголошує, що AI-інструменти не можуть бути авторами, а відповідальність за зміст залишається на

людських авторах; використання GenAI має бути прозоро розкритим [9]. Рекомендації WAME також фіксують вимоги до прозорості та відповідальності при застосуванні чатботів/GenAI у рукописах[10]. Аналогічні підходи до розкриття використання AI та підзвітності автора відображені у редакційній політиці JAMA щодо використання мовних моделей і чатботів [8]. В цілому, ці документи важливі для навчально-дослідної діяльності здобувачів вищої освіти, оскільки демонструють “дорослий” науковий стандарт: GenAI допускається лише як інструмент, але не як суб’єкт відповідальності, а прозорість – ключова умова довіри до результату.

У сфері освітніх політик UNESCO запропонувала глобальні настанови щодо застосування GenAI в освіті й дослідженнях, підкресливши необхідність одночасного регулювання, розвитку людської спроможності та забезпечення того, щоб GenAI слугував інтересам навчання і науки, а не підміняв їхні засадничі принципи[11]. Дотичними є і рекомендації Європейської мережі з академічної доброчесності (ENAI), які підкреслюють потребу керованого, етичного та освітньо-доцільного використання AI-інструментів, включно з метою розвитку компетентностей студентів і викладачів щодо їхніх можливостей і обмежень[12].

Нормативно-правова рамка України задає обов’язкові орієнтири для будь-яких освітніх моделей у цьому полі. Закон України «Про освіту» (ст. 42) визначає академічну доброчесність як сукупність етичних принципів і правил, яких мають дотримуватися учасники освітнього процесу під час навчання та здійснення наукової діяльності з метою забезпечення довіри до результатів[17-19]. Методичні документи МОН (зокрема розширений глосарій і рекомендації для ЗВО) конкретизують термінологію, підходи до профілактики порушень і необхідність раннього формування культури доброчесності в освітніх програмах[17-19]. Ціннісну основу задає концепція фундаментальних цінностей академічної доброчесності (чесність, довіра, справедливість, повага, відповідальність, мужність), яка є доречною для

інтерпретації “сірих зон” GenAI, де формально порушення може бути неочевидним, але етичний зміст – критичний [13].

Разом із тим, аналіз літератури показує наявність методичного розриву. Більшість міжнародних та національних документів описує принципи (прозорість, підзвітність, уникнення фабрикації), проте в навчальній практиці часто бракує інструментів, які переводять ці принципи у відтворювані процедури: алгоритм дій здобувачів вищої освіти під час виконання кейсу; контрольні точки, де найбільша ймовірність ризиків; рубрика оцінювання, що вимірює не лише якість тексту, а доказовість і доброчесність процесу. Саме заповнення цього розриву – через поєднання кейс-методу з контрольними точками, протоколом виконання та аналітичною рубрикою – і становить методичну основу подальшого дослідження, де модель буде представлена як науково-обґрунтоване рішення для підготовки здобувачів вищої освіти зі спеціальності «Менеджмент».

Формулювання цілей статті (постановка завдання). Особливість підготовки здобувачів вищої освіти спеціальності «Менеджмент» полягає в тому, що навчальні завдання часто орієнтовані на швидкий аналіз ситуацій, узагальнення інформації, обґрунтування управлінських рішень і використання даних. Саме ці види діяльності GenAI здатен імітувати з високою лінгвістичною якістю, але без гарантій достовірності та без авторської відповідальності. Звідси впливає потреба не в забороні інструмента, а в науково обґрунтованій моделі навчання, яка нормує використання GenAI як допоміжного ресурсу, водночас захищаючи базові цінності академічної доброчесності. Даний підхід узгоджується з міжнародними рамками, що акцентують на людській підзвітності, прозорості та розвитку спроможності до перевірки результатів GenAI [11].

Таким чином, науково-методична проблема статті полягає в необхідності створення процедурно керованого освітнього рішення для здобувачів вищої освіти зі спеціальності «Менеджмент», яке: переводить вимоги доброчесності в операційні кроки виконання навчально-дослідного

завдання; забезпечує прозорість застосування GenAI; формує навички верифікації фактів/даних і роботи з реальними джерелами; дозволяє оцінювати не лише підсумковий текст, а й доброчесність процесу створення результату.

Метою статті є наукове обґрунтування та представлення кейс-орієнтованої моделі формування культури академічної доброчесності за умов використання генеративного штучного інтелекту в межах викладання навчальної дисципліни «Основи наукових досліджень та академічна доброчесність» при підготовці здобувачів вищої освіти спеціальності «Менеджмент» в Уманському національному університеті, а також опис інструментарію, що забезпечує вимірюваність результатів її впровадження.

Для досягнення поставленої мети у статті вирішуються наступні завдання: здійснюється ідентифікація та теоретичне окреслення ключових ризиків академічної недоброчесності в умовах використання генеративного штучного інтелекту під час виконання навчально-дослідних завдань здобувачами вищої освіти спеціальності «Менеджмент» першого року навчання; обґрунтовується та описується кейс-орієнтована модель відповідального застосування GenAI з інтегрованими контрольними точками академічної доброчесності (рис. 1); модель операціоналізується у відтворюваний інструментарій, що включає протокол (алгоритм) виконання кейсу (табл. 1), аналітичну рубрику оцінювання (табл. 2) і пакет типових кейс-завдань для навчальної дисципліни «Основи наукових досліджень та академічна доброчесність» (табл. 3); проводиться апробація моделі у двох академічних групах першого курсу ($n = 56$) у дизайні «до/після» з оцінюванням ефекту за показниками рубрики, профілем типових GenAI-ризиків та результатами опитування (табл. 4 – 6); на підставі отриманих емпіричних даних узагальнюються результати й визначаються умови відтворюваності запропонованої моделі в освітньому процесі підготовки менеджерів з позицій забезпечення якості та академічної доброчесності.

Наукова новизна полягає не в самій ідеї етичного використання штучного інтелекту, а в інструменталізації цієї ідеї до рівня відтворюваної процедури й вимірюваних критеріїв, придатних для реального навчального процесу. А саме, в інтеграції контрольних точок академічної доброчесності безпосередньо в логіку виконання кейсу та в перенесенні акценту з післяконтрольного “виявлення використання штучного інтелекту” на процедурно керовану підзвітність здобувача вищої освіти, що оцінюється не за фактом застосування інструмента, а за доказовістю та доброчесністю процесу створення результату через рубрику й протокол.

Виклад основного матеріалу дослідження. Результатом дослідження є сформована та запропонована авторська модель (Рис. 1), що інтегрує вимоги академічної доброчесності у процес виконання навчально-дослідних завдань студентами-менеджерами, а не залишає їх у форматі загальних приписів.

Модель, подана на рис. 1, представлена як цикл виконання кейсу з вбудованими контрольними точками: забезпечення прозорості використання GenAI, верифікація фактів і даних, коректність цитування та реальність джерел, а також фіксація авторського внеску. Принциповою ознакою моделі є те, що контрольні точки розміщені в середині процесу, а не як завершальна “перевірка готового тексту”: це переводить доброчесність у режим процедурної необхідності, що підвищує відтворюваність і доказовість студентської роботи.

У контексті викладання навчальної дисципліни у першому семестрі це особливо важливо, оскільки закладається базова дослідницька дисципліна: студент навчається діяти як дослідник – формулювати проблему, підбирати й перевіряти джерела, обґрунтовувати висновки – не делегуючи ключові рішення генеративному інструменту.

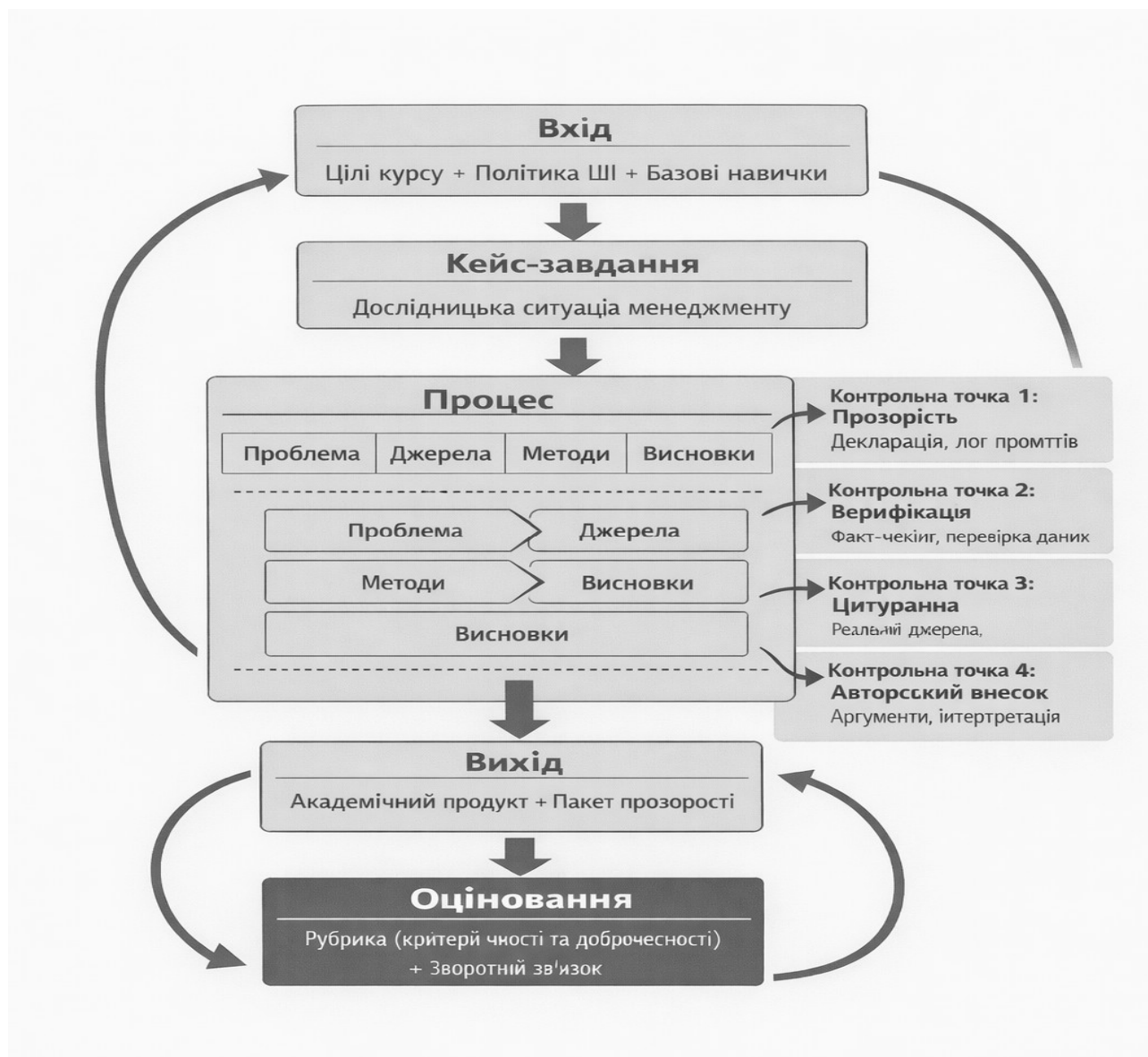


Рис. 1. Кейс-орієнтована модель формування культури академічної доброчесності та відповідального використання генеративного штучного інтелекту в межах викладання навчальної дисципліни «Основи наукових досліджень та академічна доброчесність» для студентів спеціальності «Менеджмент»

Для того, щоб модель не залишалася концептуальною схемою, її реалізацію стандартизовано у формі алгоритму виконання кейсу, поданого в табл. 1.

**Таблиця 1. Протокол виконання навчально-дослідного кейсу із
нормованим використанням GenAI**

Етап	Дія студента	Допустима роль GenAI	Вимога доброчесності	Вихідний продукт
1	Формулювання проблеми, мети, дослідницьких питань	Альтернативи формулювань	Остаточне рішення – за студентом	Паспорт дослідження
2	Пошук і відбір джерел	Підказки ключових слів/стратегії	Лише реальні джерела	Верифікований список джерел
3	Огляд літератури	Структурування огляду	Посилання лише на прочитані першоджерела	Огляд з коректними цитатами
4	Дані/методи	Пояснення методів без “вигадкування”	Верифікація фактів/даних	Опис методів/даних
5	Аналіз та інтерпретація	Варіанти інтерпретацій	Перевірка на даних; авторський висновок	Аналітичний блок
6	Оформлення	Мовне редагування	Коректне цитування	Фінальний текст
7	Пакет прозорості	Узагальнення використання	Декларація + логи + протокол перевірок	Додаток прозорості

Науково-методичний сенс протоколу полягає в тому, що він “прив’язує” використання GenAI до конкретних етапів і одночасно вбудовує запобіжники проти типових ризиків: забороняє фабрикацію даних і джерел, вимагає перевірки тверджень за першоджерелами та фіксації “пакета прозорості”. Такий підхід важливий саме для здобувачів вищої освіти спеціальності «Менеджмент», оскільки формує культуру доказового мислення: висновки мають спиратися на перевірені дані й коректні посилання, а не на “переконливість” автоматично згенерованого тексту.

За умови використання GenAI формальна “унікальність” тексту не гарантує його наукової якості: студентська робота може бути грамотною, але містити неперевірені твердження, некоректні або вигадані посилання, а також мінімальний авторський внесок. Тому наступним результатом авторського дослідження є розроблення аналітичної рубрики (табл. 2), яка

оцінює і продукт, і процес, фіксуючи ознаки доброчесної дослідницької поведінки.

Таблиця 2. Аналітична рубрика оцінювання (0–3 бали за критерієм; максимум – 24)

Критерій	0	1	2	3
Прозорість використання GenAI	Немає	Часткова	Повна декларація	Декларація + логи
Джерела й цитування	Некоректні	Частково	Коректні	Коректні + перевірені
Верифікація фактів/даних	Немає	Епізодична	Системна	Системна + протокол
Авторський внесок	Відсутній	Мінімальний	Достатній	Виразний, аргументований
Методологічна коректність	Неузгодженість	Часткова	Узгодженість	Узгодженість + обґрунтування
Логіка і структура	Порушена	Часткова	Витримана	Висока якість
Академічне письмо	Низьке	З помилками	Коректне	Високий рівень
Етичні ризики/правила	Порушення	Частково	Дотримано	Дотримано + рефлексія

Критерії рубрики узгоджені з контрольними точками запропонованої авторської кейс-моделі (рис. 1) та етапами протоколу (табл. 1), що забезпечує внутрішню логіку оцінювання: студент розуміє, що саме оцінюється, а викладач отримує інструмент уніфікації вимог у двох групах першокурсників.

Щоб продемонструвати, як саме модель працює в навчальному процесі, пакет кейсів сформовано так, щоб кожен з них цілеспрямовано «тренував» одну з найризикованіших зон GenAI-контексту: достовірність джерел, верифікацію тверджень, а також розмежування редагування й підміни авторства. При цьому кейси мають єдину структуру виконання за протоколом (табл. 1) і оцінюються за рубрикою (табл. 2), що забезпечує порівнюваність результатів у межах апробації. Додатково для кожного кейсу визначено типи навчальних доказів (артефакти виконання), які фіксують дотримання контрольних точок і слугують підставою для обґрунтованого оцінювання.

**Таблиця 3. Приклади кейс-завдань для вивчення навчальної дисципліни
«Основи наукових досліджень та академічна доброчесність» для
здобувачів вищої освіти спеціальності «Менеджмент» Уманського
національного університету (1 курс, 1 семестр) у контексті
використання GenAI**

Кейс	Дослідницька ситуація (менеджмент)	Що дозволено GenAI	Який “доказ” вимагається від студента
Кейс 1: Огляд літератури	Вибір підходу до підвищення ефективності управління (напр., KPI/OKR/Lean)	Структура огляду, варіанти формулювань	Реальні джерела + коректні цитати + список перевірених посилань
Кейс 2: Верифікація даних	Інтерпретація короткого набору даних (продажі/витрати/продуктивність)	Пояснення методів, альтернативні інтерпретації	Протокол перевірки тверджень/розрахунків + власний висновок
Кейс 3: Етична дилема	Межі допустимого використання GenAI при підготовці аналітичної записки	Лише мовне редагування	Декларація використання + аргументоване обґрунтування меж

Дослідження організовано як педагогічна апробація із порівняльним дизайном «до/після» на вибірці здобувачів першого року навчання зі спеціальності «Менеджмент» Уманського національного університету у двох академічних групах загальною чисельністю $n = 56$ (по 28 студентів у кожній групі; середній вік – 18.2 роки, 63% жінок).

Процедура апробації передбачала уніфіковані вимоги до виконання завдань у двох групах: після ознайомлення з правилами доброчесності та межами прийняттого використання GenAI здобувачі вищої освіти виконували кейси за алгоритмом (табл. 1) і подавали роботи разом із «пакетом прозорості». Етичний компонент процедури полягав у недопущенні фабрикації даних або джерел, обов’язковому декларуванні застосування GenAI та підтриманні принципу відповідальності студента за зміст поданої роботи. Дослідження проводилося з дотриманням етичних стандартів, передбачених Гельсінською декларацією: всі студенти дали інформовану згоду на участь у апробації.

Первинний зріз здійснювався на старті модуля як виконання короткого навчально-дослідного завдання без стандартизованого протоколу, тоді як повторний зріз проводився після серії кейсів за моделлю (рис. 1) і протоколом (табл. 1) з оцінюванням за рубрикою (табл. 2). Наукова логіка фіксації результатів полягає в одночасному вимірюванні (а) зміни якості робіт за критеріями рубрики, (б) зміни частоти типових GenAI-ризиків у роботах, (в) зміни усвідомлення студентами правил і практик прозорості та верифікації.

Для фіксації результатів апробації застосовано три взаємодоповнювальні групи методів. По-перше, використано експертне оцінювання студентських робіт за аналітичною рубрикою (табл. 2) з подальшим агрегуванням балів за критеріями та загальним показником, що забезпечує кількісну вимірюваність якості й доброчесності виконання завдання. По-друге, проведено контент-аналіз типових ризиків/помилочок у роботах, релевантних саме GenAI-контексту (відсутність декларації/логів, невідверифіковані твердження, некоректні або неіснуючі посилання, ознаки підміни авторського внеску), що дозволяє оцінювати не лише “середній бал”, а й зміну профілю проблем. По-третє, використано анкетування для фіксації змін у самооцінці практик прозорості, верифікації та розмежування допустимого редагування й недопустимої підміни авторства (табл. 6).

Обробку даних передбачено здійснювати методами описової статистики (середні значення, медіани, стандартні відхилення за потреби; частоти та частки для ризиків/помилочок) із порівнянням показників у зрізах «до/після». Різниця середніх значень між зрізами перевірена t-критерієм Ст'юдента ($p < 0.05$ для всіх критеріїв), що підтверджує статистичну значущість змін. Представлення результатів апробації структуровано у вигляді табличних форм (табл. 4 - 6), що забезпечує прозорість і відтворюваність подання.

Результати порівняльного оцінювання студентських робіт у форматі «до/після» за аналітичною рубрикою (шкала 0–3 бали за кожним критерієм;

сумарно 0–24) подано у таблиці 4. Дані наведеної таблиці відображають середні значення за критеріями до інтервенції та після її реалізації, а також величину приросту.

Таблиця 4. Порівняння результатів оцінювання за аналітичною рубрикою у зрізах «до/після», n = 56

Прозорість використання GenAI	0.8	2.5	+1.7
Джерельна база і цитування	1.2	2.4	+1.2
Верифікація фактів/даних	1.0	2.3	+1.3
Авторський внесок	1.1	2.2	+1.1
Методологічна коректність	1.3	2.1	+0.8
Логіка і структура	1.5	2.3	+0.8
Академічне письмо	1.4	2.0	+0.6
Етичні ризики/дотримання правил	0.9	2.4	+1.5
Загальний бал (0–24)	9.2	18.2	+9.0

З даних таблиці 4 слідує, що після впровадження моделі зафіксовано суттєве зростання загального результату: з 9,2 до 18,2 бала ($\Delta = +9,0$), що свідчить про помітне підвищення якості та доброчесності виконання навчально-дослідних завдань. Найбільші позитивні зміни спостерігаються за критеріями, безпосередньо пов'язаними з керованим використанням GenAI: прозорість використання GenAI ($\Delta = +1,7$), етичні ризики/дотримання правил ($\Delta = +1,5$), верифікація фактів/даних ($\Delta = +1,3$) і джерельна база та цитування ($\Delta = +1,2$). Це підтверджує, що вбудовані контрольні точки та стандартизований алгоритм виконання кейсу не лише підвищують формальну якість тексту, а насамперед посилюють процедурну дисципліну дослідження – прозорість, перевірюваність і відповідальність за результат.

Для того щоб доповнити кількісні результати за рубрикою (табл. 4) аналізом якісних змін у профілі ризиків академічної недоброчесності, нами у табл. 5 узагальнено частоту типових помилок/порушень, характерних саме

для GenAI-контексту. Показники наведено у відсотках для зрізів «до» і «після» впровадження кейс-орієнтованої моделі та протоколу роботи з генеративним штучним інтелектом (n = 56). Такий підхід дає змогу оцінити не лише загальне підвищення якості робіт, а й те, які саме ризики зменшуються завдяки процедурним контрольним точкам (прозорість, верифікація, коректність посилань і збереження авторського внеску).

Таблиця 5. Частота типових помилок/ризиків у роботах здобувачів першого року навчання у зрізах «до/після», n = 56

Ризик/помилка	«До» (%)	«Після» (%)
Відсутня декларація використання GenAI	78	15
Відсутні логи/пакет прозорості	85	20
Неверифіковані твердження	65	18
Некоректні/сумнівні посилання	52	12
Посилання на неіснуючі джерела	45	8
Ознаки підміни авторського внеску	60	14
Методологічна неузгодженість	55	16

Дані табл. 5 демонструють різке зниження частоти всіх ключових ризиків після впровадження моделі. Найбільш суттєво скоротилися порушення, пов'язані з прозорістю використання GenAI: відсутність декларації зменшилася з 78% до 15%, а відсутність логів/пакета прозорості – з 85% до 20%, що свідчить про формування у здобувачів вищої освіти практики підзвітного застосування інструмента. Помітно знизилися також ризики, що прямо впливають на наукову достовірність результату: частка неверифікованих тверджень зменшилася з 65% до 18%, некоректних/сумнівних посилань – з 52% до 12%, а посилань на неіснуючі джерела – з 45% до 8%. Водночас істотно зменшилися ознаки підміни авторського внеску (з 60% до 14%) і методологічна неузгодженість (з 55% до 16%). Сукупно ці зміни підтверджують, що впроваджена кейс-орієнтована модель не лише підвищує середні оцінки, а й системно знижує конкретні GenAI-типові порушення, переводячи роботу здобувачів вищої освіти у більш доказовий і добросовісний формат.

Поряд із аналізом якості робіт виконаних здобувачами вищої освіти за аналітичною рубрикою у зрізах «до/після» та зменшення частоти типових ризиків/помилки, важливо зафіксувати зміни на рівні усвідомлення та готовності здобувачів вищої освіти застосовувати GenAI відповідально. Саме тому, у табл. 6 наведено результати короткого опитування самооцінки ключових компетентностей академічної доброчесності в GenAI-контексті у зрізах «до/після» (шкала 1-5) для вибірки здобувачів вищої освіти першого року навчання зі спеціальності «Менеджмент» ($n = 56$). Таке вимірювання доповнює попередні таблиці, демонструючи не лише “зовнішній” ефект у текстах робіт, а й внутрішню зміну навчальних установок і практик, які забезпечують сталість результату.

Результати представленої табл. 6 засвідчують суттєве зростання самооцінки практик відповідального використання GenAI за всіма позиціями. Найбільший приріст зафіксовано щодо здатності описувати межі дозволеного використання GenAI (з 1,8 до 4,1, $\Delta = +2,3$), що вказує на формування у здобувачів вищої освіти чітких уявлень про допустимі й недопустимі форми застосування інструмента.

Таблиця 6. Самооцінка здобувачів вищої освіти щодо практик відповідального використання GenAI (шкала 1–5), $n = 56$

Твердження	«До»	«Після»	Δ (зміна)
Розумію, що саме потрібно декларувати	2.1	4.0	+1.9
Умію перевіряти твердження за першоджерелами	2.3	3.8	+1.5
Розрізняю редагування і підміну авторства	1.9	3.7	+1.8
Умію уникати вигаданих/неперевірених посилань	2.0	3.9	+1.9
Можу описати межі дозволеного використання GenAI	1.8	4.1	+2.3

Високими є також зміни у розумінні того, що саме потрібно декларувати ($2,1 \rightarrow 4,0$; $\Delta = +1,9$) та у вмінні уникати вигаданих/неперевірених посилань ($2,0 \rightarrow 3,9$; $\Delta = +1,9$), що узгоджується зі зниженням відповідних ризиків у роботах (табл. 5). Помітно зросла і здатність розрізняти редагування та підміну авторства ($1,9 \rightarrow 3,7$; $\Delta = +1,8$), а

також навичка перевіряти твердження за першоджерелами ($2,3 \rightarrow 3,8$; $\Delta = +1,5$). Сукупно результати опитування підтверджують, що впроваджена модель має не лише “технічний” ефект у якості текстів, а й формує у здобувачів вищої освіти стійкі компетентності прозорості, верифікації та етичної відповідальності, необхідні для доброчесного використання GenAI в подальшій навчально-науковій діяльності.

Висновки та перспективи подальших розвідок у даному напрямі. У статті розв’язано актуальну для вищої освіти та підготовки здобувачів вищої освіти зі спеціальності «Менеджмент» проблему: за умов швидкого поширення генеративного штучного інтелекту традиційні підходи до забезпечення академічної доброчесності виявляються недостатніми, оскільки зростають ризики прихованої підміни авторського внеску, невідверифікованих тверджень і некоректної джерельної бази. Відповіддю на ці виклики запропоновано кейс-орієнтовану модель навчання відповідальному використанню GenAI в межах викладання навчальної дисципліни «Основи наукових досліджень та академічна доброчесність» для здобувачів вищої освіти першого року навчання спеціальності «Менеджмент», яка переводить норми доброчесності з декларативного рівня у процедурно керований формат.

Подальші дослідження доцільно спрямувати на розширення доказової бази запропонованої моделі та уточнення механізмів її впливу. По-перше, перспективним є проведення багатоцентрової апробації на різних освітньо-професійних програмах і в різних закладах вищої освіти з порівнянням результатів між групами та викладацькими підходами, що дозволить оцінити зовнішню валідність моделі. По-друге, доцільно доповнити дизайн довготривалим вимірюванням “стійкості ефекту” – перевірити, чи зберігаються сформовані практики прозорості, верифікації та джерельної навчальної дисципліни у другому семестрі та під час виконання курсових/проектних робіт зі спеціальності «Менеджмент». По-третє, актуальним є методичне поглиблення рубрики оцінювання через введення

показників міжекспертної узгодженості та розроблення стандартизованих чек-листів верифікації, що підвищить надійність оцінювання у великих потоках студентів. По-четверте, окремим напрямом має стати порівняння різних режимів використання GenAI (лише мовне редагування; допомога в структуруванні; обмежений аналітичний супровід) та їхнього впливу на якість дослідницького мислення і ризику недоброчесності. Нарешті, перспективним є розроблення цифрового “пакета прозорості” як уніфікованого додатка до студентської роботи (лог промптів, протокол перевірок, реєстр джерел), що може бути інтегрований у внутрішні системи забезпечення якості освіти й академічної доброчесності.

Література

1. Kasneci E., Sessler K., Küchemann S. et al. ChatGPT for good? On opportunities and challenges of large language models for education // *Learning and Individual Differences*. 2023. Vol. 103. Art. 102274. DOI: 10.1016/j.lindif.2023.102274.
2. Bin-Nashwan S. A., Al-Daihani M. Academic integrity hangs in the balance: ChatGPT and generative AI in academia // *Technology in Society*. 2023. Vol. 75. Art. 102370. DOI: 10.1016/j.techsoc.2023.102370.
3. Cotton D. R. E., Cotton P. A., Shipway J. R. Chatting and Cheating: Ensuring academic integrity in the era of ChatGPT // *Innovations in Education and Teaching International*. 2024. P. 1–12. DOI: 10.1080/14703297.2023.2190148.
4. Perkins M. Academic Integrity considerations of AI Large Language Models in the post-pandemic era: ChatGPT and beyond // *Journal of University Teaching & Learning Practice*. 2023. Vol. 20(2). DOI: 10.53761/1.20.02.07.
5. Rudolph J., Tan S., Tan S. ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? // *Journal of Applied Learning & Teaching*. 2023. Vol. 6(1). DOI: 10.37074/jalt.2023.6.1.9.

6. Liang P., Li T., Guo D. et al. GPT detectors are biased against non-native English writers // *Patterns*. 2023. Vol. 4(7). Art. 100779. DOI: 10.1016/j.patter.2023.100779.
7. 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)*. 2021. P. 610–623. DOI: 10.1145/3442188.3445922.
8. Flanagin A., Bibbins-Domingo K., Berkwits M., Christiansen S. L. Nonhuman “Authors” and Implications for the Integrity of Scientific Publication and Medical Knowledge // *JAMA*. 2023. Vol. 330(8). P. 637–639. DOI: 10.1001/jama.2023.13435.
9. COPE. *Authorship and AI tools* [Электронный ресурс]. 2023. URL: <https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools> (дата звернения: 06.01.2026).
10. WAME. *Chatbots, Generative AI, and Scholarly Manuscripts* [Электронный ресурс]. 2023. URL: <https://wame.org/page3.php?id=106> (дата звернения: 06.01.2026).
11. UNESCO. *Guidance for generative AI in education and research* [Электронный ресурс]. 2023. URL: <https://unesdoc.unesco.org/ark:/48223/pf0000386693/PDF/386693eng.pdf.multi> (дата звернения: 06.01.2026).
12. Foltýnek T., Glendinning I., Dannhoferová J. et al. ENAI Recommendations for the ethical use of artificial intelligence in education // *International Journal for Educational Integrity*. 2023. Vol. 19. Art. 13. DOI: 10.1007/s40979-023-00133-4.
13. International Center for Academic Integrity. *The Fundamental Values of Academic Integrity* (3rd ed.) [Электронный ресурс]. 2021. URL: https://academicintegrity.org/images/pdfs/20019_ICAI-Fundamental-Values_R12.pdf (дата звернения: 06.01.2026).

14. NIST. *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. 2023. DOI: 10.6028/NIST.AI.100-1. URL: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf> (дата звернення: 06.01.2026).

15. NIST. *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1)*. 2024. URL: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf> (дата звернення: 06.01.2026).

16. OpenAI. *Usage policies* [Електронний ресурс]. 2025. URL: <https://openai.com/policies/usage-policies/> (дата звернення: 06.01.2026).

17. Міністерство освіти і науки України. Про затвердження Методичних рекомендацій щодо застосування інструментів та механізмів для трансформації різних сфер функціонування закладів вищої освіти на засадах прозорості та доброчесності: наказ МОН України від 17.04.2025 № 576 [Електронний ресурс]. 2025. URL: <https://mon.gov.ua/static-objects/mon/sites/1/news/2025/04/22/nakaz-576-vid-17-04-2025.pdf> (дата звернення: 06.01.2026).

18. НАЗК. The Ministry of Education and Science of Ukraine has approved methodological recommendations for building integrity in higher education institutions [Електронний ресурс]. 2025. URL: <https://nazk.gov.ua/en/news/the-ministry-of-education-and-science-of-ukraine-has-approved-methodological-recommendations-for-building-integrity-in-higher-education-institutions/> (дата звернення: 06.01.2026).

19. UNDP Ukraine. Ukraine boosts academic integrity with new guidelines for higher education [Електронний ресурс]. 2025. URL: <https://www.undp.org/ukraine/press-releases/ukraine-boosts-academic-integrity-new-guidelines-higher-education> (дата звернення: 06.01.2026).

20. Ravšelj D., et al. Higher education students' perceptions of ChatGPT: A global study // *PLOS ONE*. 2025. DOI: 10.1371/journal.pone.0315011. URL:

<https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0315011> (дата
звернення: 06.01.2026).

References

1. Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeiffer, F., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., Stadler, M. and Kasneci, G. (2023), “ChatGPT for good? On opportunities and challenges of large language models for education”, *Learning and Individual Differences*, vol. 103, available at: <https://www.sciencedirect.com/science/article/pii/S1041608023000195> (Accessed 6 Jan 2026), DOI: 10.1016/j.lindif.2023.102274.
2. Bin-Nashwan, S.A. and Al-Daihani, M. (2023), “Academic integrity in the era of artificial intelligence: The case of ChatGPT”, *Technology in Society*, vol. 75, available at: <https://www.sciencedirect.com/science/article/pii/S0160791X23002415> (Accessed 6 Jan 2026), DOI: 10.1016/j.techsoc.2023.102370.
3. Cotton, D.R.E., Cotton, P.A. and Shipway, J.R. (2024), “Chatting and cheating: Ensuring academic integrity in the era of ChatGPT”, *Innovations in Education and Teaching International*, vol. 61, no. 2, pp. 228–239, available at: <https://www.tandfonline.com/doi/full/10.1080/14703297.2023.2190148> (Accessed 6 Jan 2026), DOI: 10.1080/14703297.2023.2190148.
4. Perkins, M. (2023), “Academic Integrity considerations of AI Large Language Models in the post-pandemic era: ChatGPT and beyond”, *Journal of University Teaching and Learning Practice*, vol. 20, no. 2, available at: <https://ro.uow.edu.au/jutlp/vol20/iss2/07/> (Accessed 6 Jan 2026), DOI: 10.53761/1.20.2.07.
5. Rudolph, J., Tan, S. and Tan, S. (2023), “ChatGPT: Bullshit spewer or the end of traditional assessments in higher education?”, *Journal of Applied Learning and Teaching*, vol. 6, no. 1, available at:

<https://journals.sfu.ca/jalt/index.php/jalt/article/view/689> (Accessed 6 Jan 2026), DOI: 10.37074/jalt.2023.6.1.9.

6. Liang, W., et al. (2023), “GPT detectors are biased against non-native English writers”, *Patterns*, vol. 4, no. 7, available at: [https://www.cell.com/patterns/fulltext/S2666-3899\(23\)00130-7](https://www.cell.com/patterns/fulltext/S2666-3899(23)00130-7) (Accessed 6 Jan 2026), DOI: 10.1016/j.patter.2023.100779.

7. Bender, E.M., Gebru, T., McMillan-Major, A. and Shmitchell, S. (2021), “On the dangers of stochastic parrots: Can language models be too big?”, *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, pp. 610–623, available at: <https://dl.acm.org/doi/10.1145/3442188.3445922> (Accessed 6 Jan 2026), DOI: 10.1145/3442188.3445922.

8. Flanagin, A., Bibbins-Domingo, K., Berkwits, M. and Christiansen, S.L. (2023), “Nonhuman ‘Authors’ and implications for the integrity of scientific publication and medical knowledge”, *JAMA*, vol. 330, no. 8, pp. 637–639, available at: <https://jamanetwork.com/journals/jama/fullarticle/2807648> (Accessed 6 Jan 2026), DOI: 10.1001/jama.2023.13435.

9. Committee on Publication Ethics (2023), “Authorship and AI tools”, available at: <https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools> (Accessed 6 Jan 2026).

10. World Association of Medical Editors (2023), “Chatbots, Generative AI, and Scholarly Manuscripts”, available at: <https://wame.org/page3.php?id=106> (Accessed 6 Jan 2026).

11. UNESCO (2023), “Guidance for generative AI in education and research”, available at: <https://unesdoc.unesco.org/ark:/48223/pf0000386693> (Accessed 6 Jan 2026).

12. Foltýnek, T., Glendinning, I., Dannhoferová, J. et al. (2023), “ENAI Recommendations for the ethical use of Artificial Intelligence in Education”, *International Journal for Educational Integrity*, vol. 19, available at:

<https://link.springer.com/article/10.1007/s40979-023-00133-4> (Accessed 6 Jan 2026), DOI: 10.1007/s40979-023-00133-4.

13. International Center for Academic Integrity (2021), *The Fundamental Values of Academic Integrity*, 3rd ed., available at: https://academicintegrity.org/images/pdfs/20019_ICAI-Fundamental-Values_R12.pdf (Accessed 6 Jan 2026).

14. National Institute of Standards and Technology (2023), *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, available at: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf> (Accessed 6 Jan 2026), DOI: 10.6028/NIST.AI.100-1.

15. National Institute of Standards and Technology (2024), *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1)*, available at: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf> (Accessed 6 Jan 2026), DOI: 10.6028/NIST.AI.600-1.

16. OpenAI (2025), “Usage policies”, available at: <https://openai.com/policies/usage-policies/> (Accessed 6 Jan 2026).

17. Ministry of Education and Science of Ukraine (2025), Order “On approval of Methodological Recommendations on the Application of Tools and Mechanisms for the Transformation of Various Spheres of Functioning of Higher Education Institutions on the Principles of Transparency and Integrity”, available at: <https://mon.gov.ua/static-objects/mon/sites/1/news/2025/04/22/nakaz-576-vid-17-04-2025.pdf> (Accessed 6 Jan 2026).

18. National Agency on Corruption Prevention (2025), “The Ministry of Education and Science of Ukraine has approved methodological recommendations for building integrity in higher education institutions”, available at: <https://nazk.gov.ua/en/news/the-ministry-of-education-and-science-of-ukraine-has-approved-methodological-recommendations-for-building-integrity-in-higher-education-institutions/> (Accessed 6 Jan 2026).

19. UNDP Ukraine (2025), “Ukraine boosts academic integrity with new guidelines for higher education”, available at: <https://www.undp.org/ukraine/press-releases/ukraine-boosts-academic-integrity-new-guidelines-higher-education> (Accessed 6 Jan 2026).

20. The Verkhovna Rada of Ukraine (2017), The Law of Ukraine “On Education”, available at: <https://zakon.rada.gov.ua/laws/show/2145-19> (Accessed 6 Jan 2026).

Отримано редакцією журналу / Received: 14.03.26

Прорецензовано / Revised: 18.03.26

Схвалено до друку / Accepted: 20.03.26