

ШІ у наукових статтях 2026

Як уникнути відмови та відповідати новим вимогам
Scopus / Web of Science

AI-detection • нові editorial-політики • етика • робочі стратегії для авторів і журналів

Василь Малець

Експерт із наукових публікацій та штучного інтелекту.
Директор компанії «Академія статей».

- Понад **2000** опублікованих робіт у **Scopus / Web of Science**
- Консультую **міжнародні видавництва** щодо індексації журналів
- Особисто провів переговори з **2500+ редакціями Scopus/WoS**
- Автор аналітичної серії «**Претенденти на виліт зі Scopus**»
- Тестую **AI-detection системи** та відстежує політики
- **Elsevier, Springer, Wiley**
- Допоміг журналам **успішно пройти індексацію**

Практика > теорія. Робота з редакціями — щодня.





Тривалість вебінару

60–80 хвилин

✓ **Практичний чек-лист**

Як використовувати ШІ в наукових статтях без ризику відмови

✓ **Топ-7 AI-інструментів, що заощаджують до 100 годин роботи**

✓ **Сертифікат** учасника вебінару

✓ **Презентацію** спікера

✓ **Запис вебінару**

На вебінарі Ви дізнаєтесь

1. Вступ
2. Про лектора
3. Програма вебінару
4. Чому ця тема вирішальна у 2026 році
5. Новий світ публікацій: епоха генеративного ШІ
6. Політики 2025–2026: оновлені стандарти Scopus і великих видавців
7. Як виявляють AI-контент: інструменти Elsevier і редакцій
8. Етика використання ШІ: текст, дані, зображення
9. Декларації AI: що писати, щоб не відхилили (шаблони)
10. Реакція середніх і малих видавців: реальні кейси
11. Як редактори журналів використовують AI-інструменти у 2026
12. Типові помилки авторів — і як уникнути прапорця «AI-abuse»
13. Інструменти AI, які реально працюють для науковців
14. Топові інструменти для перетворення AI тексту на людський
15. Як автори використовують ШІ правильно: оптимальний workflow
16. AI і українська академічна система: стан та виклики
17. Чи замінить ШІ науковців і викладачів у 2030?
18. Як зміниться роль ученого, редактора та журналу найближчі 5–10 років
19. Як Академія статей використовує ШІ у роботі
20. Висновки + Q&A
21. Відповіді на питання

Чому саме 2024-2025 роки став переломним для наукових публікацій і ШІ



Різке зростання ретракцій через AI-пов'язані проблеми

- 2023 рік → рекорд ретракцій: 14 000+ відгуків.
- 764 AI-пов'язані статті проаналізовано: 667 ретракцій припали саме на 2023 рік.
- Основні причини: проблеми з peer review, помилки у даних, "unethical AI use".



Масове проникнення ШІ у тексти наукових статей

- 22% анотацій у Computer Science на arXiv містять LLM-фрагменти
- 10% biology-абстрактів
- 14% biomedical-журнальних абстрактів (2024)



Переломний момент 2024–2025: два тренди зішлись

- Автори **масово** використовують LLM для написання текстів, опису методів, структурування абстрактів.
- Видавці **масово** впровадили AI-детектори та системи integrity-перевірок.
- Це вперше створило **конфлікт між “генеративним AI” та “контрольним AI”**.
- Саме тому тема ШІ стає **критичною**, а не просто популярною.



Жорсткі нові політики та обмеження видавців

- AI не може бути автором
- AI дозволено лише як інструмент, за умови людського контролю.
- **Обов'язкове AI-disclosure**: недекларування = причина відмови.
- **Scopus оновив content policy**: журнали повинні мати власну GenAI-політику.
- **Рецензенти не мають права завантажувати рукописи у публічні LLM** — це порушення конфіденційності.



Нові редакційні інструменти детекції ШІ (2024–2026)

- **Springer Nature**: Geppetto, SnappShot, Snapp (14 автоматичних перевірок).
- **Wiley**: AI-signal detection (10–13% рукописів отримують AI-flag до рецензування).
- **Frontiers**: AIRA — система з 60 integrity-чеків.
- **STM Integrity Hub**: перевіряє 125 000 рукописів щомісяця; повна інтеграція AI-детекції.
- **iThenticate / Crossref**: новий модуль AI-writing detection.

У 2026 році успішна публікація залежить не тільки від наукової якості статті, а й від того, наскільки етично, прозоро та коректно автор використав ШІ.

Новий світ наукових публікацій: епоха генеративного ШІ



ШІ став масовим інструментом автора

- ChatGPT, Claude, Gemini стали стандартними інструментами написання текстів.
- Автори використовують ШІ для формулювання ідей, структурування статті, написання розділів, редагування.
- Час підготовки рукопису скоротився **в рази** — з місяців до днів.



Стиль наукових статей змінився

- Тексти стають «уніфікованими» та «згладженими» через використання LLM.
- З'явилися характерні патерни: чіткі переходи, повторювані структури, типовий «GPT-стиль».
- Рецензенти почали впізнавати AI-руки навіть без детекторів.



Зросла кількість “науково порожніх” текстів

- ШІ допомагає писати красиво, але **не додає наукової цінності**.
- Збільшилась кількість рукописів з ідеальною англійською, але слабкою методологією.
- Частина авторів замінює думання → генерацією тексту.



Генеративний ШІ змінив workflow редакторів і рецензентів

- Редакції впровадили LLM для первинного скринінгу (AI-assisted triage).
- Автоматичний аналіз відповідності темі, структури, методології.
- AI використовують для перевірки рецензентів (faux reviewers, аномалії в рецензіях).
- З'явилися системи «AI-підказок» для редакторів → Snapp, AIRA, Geppetto.



Зникла гра “приховай ШІ” — тепер головне прозорість

- У 2026 вже немає сенсу “маскувати” AI — детектори розвиваються швидше.
- Ключове питання журналів тепер: **чи відповідально використано ШІ?**
- Прозоре декларування + людський контроль = AI-safe workflow.

У 2026 році успішна публікація залежить не тільки від наукової якості статті, а й від того, наскільки етично, прозоро та коректно автор використав ШІ.

Політики 2025–2026: нові стандарти Scopus і провідних видавців



Спільні політики великих видавців (Elsevier, Springer Nature, Wiley, Taylor & Francis)

- **AI не може бути автором** (позиція COPE, ICMJE та всіх великих видавців).
- **AI допустимий лише як інструмент**, а не як джерело наукового змісту.
- **Автор несе повну відповідальність за текст**, навіть якщо його частково генерував ШІ.
- **Обов’язкове disclosure** будь-якого використання AI (у Methods, Acknowledgements або cover letter).



Заборони та обмеження 2025–2026

- **Заборонено завантажувати рукописи у публічні LLM** (конфіденційність, IP-ризиками).
- **Заборонено використовувати AI у peer review** — рецензенти не мають права передавати рукописи AI-моделям.
- **Заборонено використовувати ШІ для генерації даних, зображень, таблиць**, якщо це не задокументовано і не перевірено.
- **AI-generated rewriting / summarizing** → тільки з повним розкриттям.



Жорсткі нові політики та обмеження видавців

- З 2025 року **Scopus вимагає, щоб журнали мали власну GenAI-політику**.
- Очікуване: **прозора декларація використання AI** у контенті та в процесах рецензування.
- Elsevier вимагає:
 - AI → лише для мови/редагування
 - **обов’язкова декларація в окремому розділі**
 - автор має гарантувати відсутність фальсифікацій, вигаданих посилань, hallucinations
- Scopus AI (LLM-інструмент 2024–2026) підсилює роль перевірених джерел, а не “чистого” AI-тексту.



AI-детектори та інфраструктура контролю (2025–2026)

- **Wiley Papermill Detection:**
 - працює на ~270 журналах
 - **10–13%** рукописів отримують AI-signal до рецензії
 - **600–1000** відхилень щомісяця
- **Springer Nature:**
 - **Geppetto** — AI-детекція тексту
 - **SnappShot** — аналіз зображень
 - **Snapp** — 14 автоматичних перевірок на 100 000+ подань
- **Frontiers AIRA:**
 - 60 integrity-чеків до peer review
 - у 2024 р. половина відхилень приймалися на основі AIRA
- **STM Integrity Hub:**
 - 40 видавців
 - **125 000+** рукописів/міс
 - 1 000 підозрілих рукописів блокується щомісяця
 - з 2025 р. → інтегрована AI-детекція тексту
- **iThenticate (Crossref Similarity Check):**
 - модуль AI-writing detection використовується у Springer, Wiley, IEEE.

У 2026 році ключовою вимогою стає не заборона ШІ, а повна прозорість, відповідальність автора і жорсткий editorial-контроль за даними, стилем та оригінальністю.

Як виявляють AI контент: інструменти Elsevier і редакцій



Як видавці дозволяють (і забороняють) використовувати ШІ

- **Springer Nature:** ШІ дозволено для мови, але не для створення змісту; **disclosure** обов'язковий у Methods; рецензентам заборонено використовувати LLM.
- **Wiley:** "Transparent and detailed" **disclosure**; перевірка IP-умов; AI не може замінювати автора.
- **Elsevier:** ШІ можна застосовувати лише для читабельності; обов'язковий AI declaration statement; повна відповідальність автора.
- **Taylor & Francis:** ШІ дозволено для ідей та мови; **disclosure** з назвою інструмента; заборона AI у peer review.
- **SAGE:** AI-асистування дозволене, але не для створення даних; повна відповідальність автора; **disclosure** обов'язковий.



Які AI-детектори використовують великі видавці

- **Springer Nature:** Geppetto (детекція тексту), SnappShot (аналіз зображень), Snapp (14 автоматичних перевірок).
- **Elsevier:** Crossref AI-модуль + власні аномалічні перевірки стилю, посилань і структури.
- **Wiley:** Papermill/AI Detection Service (6 модулів), інтеграція у Research Exchange.
- **Taylor & Francis:** Внутрішня AI-детекція + ручне editorial screening.
- **SAGE:** GPTZero, Turnitin AI, Copyleaks; співпраця з Wiley.



Які AI-сигнали вважаються проблемними

- Неприродний «LLM-гладкий» стиль
- "Tortured phrases" (дивні синоніми)
- Вигадані DOI / неіснуючі джерела
- Логічні аномалії в структурі
- Різка невідповідність між «ідеальною англійською» та слабкою методологією



Найчастіші причини відмов у 2025–2026 роках

- Недеклароване використання ШІ
- Використання AI для створення основного змісту
- Вигадані або некоректні посилання
- IP-порушення: завантаження рукописів у публічні LLM
- Підозрілий AI-стиль, який суперечить методологічній частині

У 2026 році ключовий стандарт видавців: ШІ допустимий лише як допоміжний інструмент. Усе інше — **disclosure**, людський контроль і відповідальність автора.

Етика використання ШІ у наукових текстах, даних і зображеннях



ТЕКСТ — що дозволено, а що вважається порушенням

Дозволено:

- Редагування мови та стилю (language editing).
- Самаризація/чернетки під повним людським контролем.

Заборонено:

- Генерація великих частин статті без автора.
- Вигадані факти та джерела (“галюцинації”).
- Приховування використання ШІ.

Disclosure:

- Треба вказати який інструмент, яку версію, з якою метою.
- Дозволені місця: **Methods, Acknowledgements, AI use statement.**



ДАНІ — чітка межа між аналізом і фабрикацією

Заборонено:

- Генерація даних ШІ → прирівнюється до **falsification/fabrication**.
- “Заповнювання” прогалин у датасеті синтетичними точками.

Ризики:

- Дані виглядають правдоподібно, але не існують у природі.
- Неможливо відтворити результати.

Дозволено:

- Аналіз/класифікація реальних даних ШІ з обов’язковою людською перевіркою.



ЗОБРАЖЕННЯ — найсуворіша зона відповідальності

Заборонено:

- Генеративні зображення (Midjourney, DALL·E тощо).
- Домальовування деталей у гелях, мікроскопії, медичних знімках.

Дозволено:

- Легке глобальне покращення чіткості/контрасту без зміни змісту.

Виявлення маніпуляцій:

- SnappShot, ImageTwin, Proofig — автоматична детекція фальсифікацій.

Загальні правила етики ШІ

- ШІ не може бути автором.
- Автор несе **100% відповідальність** за текст і дані.
- Заборонено завантажувати рукописи у публічні LLM(конфіденційність!).

Текст:	✓ “Я використав ChatGPT для виправлення граматики.”	✗ “ChatGPT написав Discussion, я вставив без змін.”
Дані:	✓ “AI класифікував зображення, результати перевірів експерт.”	✗ “AI створив 50 додаткових точок даних.”
Література:	✓ “AI допоміг знайти статті, я їх прочитав.”	✗ “AI згенерував список літератури, я не перевірів DOI.”

Декларації AI у статтях: як правильно вказати використання ШІ (шаблони)



Коли декларація обов'язкова?

- Якщо ШІ використовувався для редагування тексту
- Якщо ШІ допомагав у структуруванні або самаризації
- Якщо ШІ застосовувався у будь-якому етапі аналізу, роботи з даними чи зображеннями

“Недеклароване використання AI = висока ймовірність відмови.”



Де ставити декларацію?

Methods (рекомендовано для технічного опису)

Acknowledgements (коли використання мінімальне)

Окремий розділ: “Declaration of AI use”
(вимагає Elsevier, Springer, Wiley)



Шаблон 1 — Мінімальне використання AI (language editing)

“The authors used ChatGPT-5.1 (OpenAI, 2025) solely for language editing and grammar correction. All scientific content, interpretations, and conclusions were generated by the authors. The final text was reviewed and verified by the authors.”

Підходить для: Introduction, Discussion, форматування тексту.



Шаблон 2 — AI допомагав структурувати або самаризувати текст

“Generative AI tools (ChatGPT-4, OpenAI, 2025) were used to summarize background literature and improve the organization of the manuscript. No AI-generated text was used without critical editing. The authors take full responsibility for the accuracy and integrity of the content.”

Підходить для: Methods, literature overview.



Шаблон 3 — AI у роботі з даними / аналізом

“AI tools (Gemini Advanced, 2025) were used for preliminary data classification. All datasets were generated by the authors, and all analyses were verified manually. No synthetic or AI-generated data were used.”

Підходить для: Data analysis, Results.



Шаблон 4 — універсальна коротка декларація

“AI tools were used in the preparation of this manuscript. All usage was limited to language editing and organizational support. The authors remain fully responsible for the content, data, and conclusions.”

Для журналів, які вимагають коротку формулу.



Чого НЕ МОЖНА писати у декларації:

- “AI wrote parts of the article.”
- “AI generated the data.”
- “The authors are not responsible for potential mistakes produced by AI.”

Це гарантує відмову — автор знімає з себе відповідальність.

Головне правило 2026 року:

AI можна використовувати — але потрібно чесно, чітко й конкретно це задекларувати.

Як малі та середні видавці реагують на ШІ у 2024–2026: реальні кейси

Найтиповіші дії редакцій (узагальнення 2024–2026)

- Автоматичний AI-скринінг із порогамі **10–20%**.
- Обов’язковий AI Usage Statement при будь-яких AI-сигналах.
- Ручний стильовий аналіз: відхилення при “відсутності авторського голосу”.
- Перевірка DOI → відмова при вигаданих посиланнях.
- Вимоги надати **raw data** для підтвердження не-AI походження.
- Відеоінтерв’ю з авторами для перевірки компетентності.
- Зростання desk-reject rate до **60–70%**, щоб відсіяти AI-контент.

4 ключові кейси (без назв журналів, тільки суть)

Desk rejection через AI-ідеї

Автор чесно зізнався, що ШІ допоміг з формуванням ідей.

Результат: Одразу відхилено: “журнал приймає лише оригінальний людський інтелектуальний внесок”.

Урок: У малих редакцій буквально тлумачення політик.

Відмова через “LLM-style writing”

Текст занадто гладкий, однакові структури речень, немає авторського стилю.

Результат: Desk rejection, навіть без AI-детектора.

Урок: Стилістичний аналіз стає важливішим за технології.

Перевірка первинних даних (raw data proof)

Дані виглядають «надто ідеально».

Результат: Журнал запросив сирі дані → автор не надав → рукопис відхилено.

Урок: Малі журнали активно запроваджують політику “data authenticity”.

Відеоінтерв’ю для підтвердження авторства

Стаття структурно бездоганна, але «без голосу автора».

Результат: Редакція вимагала 30-хвилинне відеоінтерв’ю.

Урок: Low-tech методи замінюють дорогі системи.

Інструменти, які реально використовують малі видавці

- **GPTZero** (основний безкоштовний AI-детектор)
- **ZeroGPT** (≈95% точності, великі ліміти)
- **Copyleaks AI** (API для невеликих журналів)
- **Turnitin AI-модуль** (через університетські ліцензії)
- **TinEye / Reverse Image Search** для перевірки зображень

Малі й середні журнали у 2026 працюють за принципом:

AI-фільтр → ручна перевірка → жорсткий disclosure → відмова при сумнівах.

Замість дорогих систем — пороги, інтерв’ю, перевірка стилю, raw data, прості детектори.

Як редактори журналів використовують AI-інструменти у 2026



AI на етапі Desk Review (перші 30 секунд)

AI виявляє AI-генерацію, paper mill-патерни, мовні та структурні аномалії.

Springer Geppetto · Frontiers AIRA · STM Integrity Hub · Turnitin / Copyleaks



Стиль, посилання, фейкові цитати

Перевірка вигаданих DOI, irrelevant citations, AI-шаблонів у бібліографії.



AI проти фальсифікації фігур

Пошук дублікацій, маніпуляцій, synthetic images у графіках і гел-блотах.

SnappShot · STM image modules · Reverse image search



AI як асистент редактора

Класифікація рукописів, пошук рецензентів, перевірка ORCID та афіліцій.



Що робить редактор після AI-скринінгу

✗ Desk reject — сильні paper mill сигнали

❓ Запит пояснень — LLM-стиль, підозрілі цитати

📁 Запит raw data — “ідеальні” або synthetic дані



AI = фільтр. Людина = рішення.

AI добре:

- виявляє патерни
- масштабується
- швидкий

AI погано:

- не оцінює новизну
- не розуміє сенс
- дає false positives

У 2026 році AI — це масовий скринінг.

Але фінальні рішення редактори залишають за собою.

Типові помилки авторів при використанні ШІ

✗ Використання ШІ для інтерпретації результатів

- ШІ пояснює значення результатів
- Відсутня авторська наукова позиція

👉 Інтерпретація = інтелектуальний внесок автора

✗ Генерація Abstract / Conclusions без зв'язки з текстом

- Abstract не відповідає тілу статті
- Conclusions обіцяють більше, ніж зроблено

👉 LLM-style overclaiming → desk reject

✗ Невідповідність між методологією та описом

- У методах одне, у результатах — інше
- Терміни змінюються після редагування ШІ

👉 LLM hallucination після "academic editing"

✗ Автоматичне academic polishing, яке ламає точність

- Узагальнення замість конкретики
- Strong claims замість обережних формулювань

👉 Style over substance

✗ Вигадані або «підозріло зручні» посилання

- Ідеально підігнані під аргумент
- Відсутня реальна наукова дискусія

👉 Миттєвий бан без права на правки

✗ Текст із ШІ без вичитки та авторського втручання

- Однаковий стиль по всій статті
- Відсутні сліди ручного редагування

👉 Підозра на AI-abuse

ШІ — це інструмент.

Порушення починаються там, де зникає автор.

AI-інструменти, які реально працюють для науковців

GPT — для структури й мислення

Спеціалізовані AI — для фактів, джерел і перевірки



CORE AI

GPT (ChatGPT) або Claude

Для чого: структура та мислення

- Побудова логіки статей і дисертацій
- Пояснення складних ідей «людською мовою»
- Формулювання research questions і гіпотез
- Робота з аргументацією і стилем

Perplexity

Для чого: актуальні джерела

- Пошук нових досліджень (2024–2026)
- Перевірка тверджень і фактів
- Посилання на реальні джерела
- Моніторинг AI-сервісів і оновлень



SPECIALIZED RESEARCH TOOLS

Elicit

Для чого: systematic review

- Пошук по 125+ млн статей
- AI-скринінг релевантності
- Автоматичне видобування даних
- Supporting quotes для перевірки

Scite.ai

Для чого: перевірка цитувань

- Smart Citations (support / contrast)
- Аналіз контексту цитування
- Виявлення спірних джерел
- Оцінка надійності літератури

SciSpace Copilot

Для чого: швидке читання

- Пояснення термінів і формул
- Аналіз таблиць і графіків
- Citation-backed відповіді
- Інтеграція з Zotero

Paperpal

Для чого: академічне редагування

- Human-like language polishing
- Перевірка логіки і стилю
- Journal readiness checks
- Plagiarism & consistency scan

AnswerThis

Для чого: research gaps

- Пошук відкритих питань у галузі
- Аналіз слабких місць літератури
- Формування напрямів дослідження
- Мислить як рецензент

Успішні публікації = правильний інструмент у правильному місці

Топові інструменти для перетворення AI тексту на природній

Як зробити текст живим, зрозумілим і без AI відбитків!



HumanizerPro

HumanizerPro вважається лідером галузі для масштабного перетворення AI-контенту.

- Підтримка до 3000 слів за один раз і до 1 мільйона слів на місяць
- Вбудований AI-детектор для отримання результатів у реальному часі
- Висока ефективність проти основних детекторів (Turnitin, GPTZero, Originality.ai)
- Ефективно зберігає первинне значення тексту



StealthWriter

Потужний інструмент з розширеними можливостями налаштування.

- 10 рівнів гуманізації (від 1 до 10) для точного контролю
- 8 різних стилів письма
- Вбудовані інструменти: генератор контенту, детектор AI, редактор
- Спеціально тренований проти нової моделі Turnitin 2026 року



WriteHuman AI

Спеціалізований інструмент, оптимізований саме для академічного письма

- Функція "Retry" для кількох спроб гуманізації
- Модель "Enhanced", спеціально налаштована для Turnitin
- Орієнтований на академічні стандарти
- Зберігає аргументи, факти та ключову інформацію без втрати змісту



Undetectable AI

Один з найпопулярніших інструментів з понад 17 мільйонами користувачів

- Комплексне рішення "все в одному": генератор, гуманізатор, детектор
- Перевірка через 8 різних AI-детекторів одночасно
- Гарантія повернення грошей, якщо гуманізований контент буде позначений
- Надійна робота з GPT-4 та Claude

Інструменти допомагають. Відповідальність завжди на автору!

Оптимальний workflow автора з ШІ у 2026



ЕТАП 1. ІДЕЯ, ТРЕНДИ, GAP

Ціль: зрозуміти що досліджувати і чому це важливо

Інструменти:

- **Perplexity** – швидкий ресерч, фактчек
- **AnswerThis** – відкриті research questions

Що робимо:

- шукаємо тренди, прогалини, необроблені проблеми
- формуємо 3–5 potential research questions

Порада: Perplexity = пошук, рішення приймає людина.



ЕТАП 2. ЗБІР ЛІТЕРАТУРИ (Literature Mapping)

Інструмент:

- **Elicit**

Що робимо:

- витягуємо з PDF ключові дані: автори, методи, sample, findings
- будуємо evidence table
- відсіюємо зайве

Порада: Elicit не читає за автора — тільки структурує.



ЕТАП 3. ГЛИБОКЕ ЧИТАННЯ (Deep Reading)

Інструменти:

- **SciSpace**

Що робимо:

- завантажуюмо PDF
- просимо пояснити складні параграфи, методи, обмеження
- переводимо статистику «на людську мову»

Результат: реальне розуміння, а не «начитаний GPT-виступ».



ЕТАП 4. НАПИСАННЯ ТЕКСТУ

Правило: текст пише автор. AI = допомога, не генератор статті.

Як можна використовувати GPT/CLAUDE:

- чернеткові формулювання
- покращення структури
- пояснення складних фрагментів

ВАЖЛИВО:

Якщо використовують GPT → «humanize» текст (перевірка стилю, логіки).



ЕТАП 5. АНГЛІЙСЬКА, ПОЛІРОВКА, ПЕРЕВІРКИ

Інструменти:

- **Paperpal** – мовний редактор
- **Zotero / Mendeley** – цитування
- **ChatGPT Pre-review** – технічний чек перед подачею

Що робимо:

- фінальна редактура мови
- перевірка правильності цитувань
- технічний аудит перед submission

Ціль: щоб у редактора не з'явилось «facepalm»

AI не замінює автора — він прискорює рутину. Ідеї, дані, аналіз і висновки мають бути створені людиною.

AI і Українська академічна система

Новий Закон України(10392) про академічну доброчесність і ШІ (2026)

Що вважається порушенням

Недоброчесне використання результатів ШІ — офіційний вид порушення.

До нього входить:

- генерація текстів без зазначення цього факту
- використання ШІ для написання робіт замість себе
- фальсифікація/фабрикація даних за допомогою ШІ
- використання матеріалів, створених ШІ, без перевірки та відповідальності
- недеклароване застосування будь-яких ШІ-інструментів

Нові правила

- 1. Прозоре декларування ШІ — обов'язкове
- 2. Результати ШІ не вважаються авторськими
- 3. Використання ШІ треба обґрунтувати
- 4. Обов'язкове посилання на програму ШІ

Санкції за порушення

ШІ як інструмент мислення, а не автоматизації життя

Для студентів:

попередження → повторна робота → втрата стипендії → відрахування

Для викладачів/науковців:

заборона керувати роботами
заборона брати участь у грантах
звільнення

позбавлення звання / ступеня

Для закладів:

ризик втрати акредитації або ліцензії

ШІ - це інструмент. Відповідальність - завжди на людині! Тепер це закон!

Чи замінить ШІ викладачів у 2030?



Прогноз 2030: автоматизація

- Персональні AI-тьютори (one-to-one)
- Адаптація пояснень під психотип і рівень знань
- Миттєвий фідбек і оцінка процесу мислення, а не лише відповіді
- Динамічні навчальні плани для кожного студента
- Автоматизація до 50% задач викладача
- (планування, оцінювання, адміністрування)

McKinsey Global Institute, 2024



Викладач ≠ лектор. Викладач = куратор і ментор

- Затвердження AI-згенерованих курсів
- Визначення навчальних цілей разом зі студентом (коучинг)
- Менторство та соціально-емоційна підтримка
- Фасилітація конфліктів і складних ситуацій
- Усні захисти, оцінювання в класі
- Роль судді й експерта, а не "перевірятьника робіт"



Гібридна освіта

- Елітарна освіта → живе спілкування як преміум-послуга
- Масова освіта → AI-курси + куратори
- Посилення розриву між масовою та елітарною освітою
- Університети тримають викладачів із власною цінністю

ЩО ТЕПЕР СТАЄ ЦІННІСТЮ ВИКЛАДАЧА

Наукова вага та публікації

Репутація та експертність

Вміння працювати з AI-інструментами

Менторські та соціальні навички

Здатність керувати AI-агентами, а не конкурувати з ними

ШІ не замінить викладачів.

Він замінить тих викладачів, які не вміють працювати з ШІ.

Як зміниться роль ученого, редактора та журналу найближчі 5–10 років



1. Учений

Від «письменника» → до архітектора знань

Що автоматизує ШІ:

- Чернетки статей і форматування (XML / LaTeX)
- Переклад на академічну англійську
- Літературні огляди, первинний ресерч
- Симуляції, автоматизовані експерименти
- Технічний скринінг (статистика, плагіат)

Human Core:

- Формулювання наукової проблеми
- Вибір методології
- Оцінка новизни та значущості
- Авторська й юридична відповідальність
- Етика дослідження



2. Редактор

Від «логіста» → до стратега і арбітра

Що автоматизує ШІ:

- Desk rejection за формальними критеріями
- Пошук і запрошення рецензентів
- Перевірка даних, графіків, статистики
- Виявлення AI-spam і фабрикацій

Human Core:

- Стратегічні рішення (що публікувати)
- Оцінка спірних і потенційно проривних робіт
- Управління конфліктами
- Формування наукового напрямку журналу



3. Журнал

Від «платформи» → до гаранта довіри

Що робить ШІ:

- AI-рецензування (швидкість, масштаб)
- Форматування, технічна перевірка
- Аудит даних і репродуктивності
- Скорочення часу від submit до publish

Нова роль журналу:

- Гарантія, що дані не сфабриковані
- Прозорий AI-audit
- Людська валідація сенсу
- Репутаційний фільтр якості

«Біполярна наука»

• Tier 1 (Еліта):

повільно, дорого, жорстка людська валідація

• Tier 2 (Мас-маркет):

швидко, автоматизовано, змішана якість

! Ризик: доступ до Tier 1 стає дорожчим для вчених з менших систем.

**ШІ забере 60–70% технічної роботи,
але 100% відповідальності залишиться за людиною.**

Як Академія статей використовує ШІ у роботі



Академія статей × ШІ

Наш підхід

Усе, що можна автоматизувати — автоматизується.

Усе, що потребує судження — залишається за людиною.

Як ми використовуємо ШІ:

- Оптимізація внутрішніх процесів і ланцюгів роботи
- Власні AI-агенти для допоміжних рішень
- Спеціалізований AI-експерт у сфері наукових публікацій
- Попереднє інтелектуальне рецензування рукописів
- Аналітична підтримка підбору авторів і рецензентів
- Перевірка гіпотез і стратегічних напрямів
- Підтримка маркетингових і контент-рішень

ШІ прискорює, але не приймає фінальних рішень.



Принципи, яких ми дотримуємось

- ШІ ≠ автор і ≠ рецензент
- ШІ — інструмент попереднього аналізу
- Відповідальність завжди на експерті
- Людська валідація — обов'язкова
- Етика і доброчесність > швидкість

Меседж:

Ми використовуємо ШІ не для заміни експертів, а для того, щоб експерти займались справжньою роботою.



Особисто: як я, Василь Малець, використовую ШІ

ШІ як інструмент мислення, а не автоматизації життя

- Моделюю сценарії рішень і наслідків
- Рефлекую власну поведінку та емоції
- Аналізую книги й складні ідеї
- Працюю з науковими статтями через браузерні AI-інструменти
- Створюю стратегії й перевіряю гіпотези
- Заглиблююсь в історичні та міждисциплінарні контексти

**Для мене ШІ — це не заміна мислення.
Це спосіб мислити глибше й чесніше.**

Майбутнє — не за тими, хто “користується ШІ”, а за тими, хто вміє з ним думати.

Висновки + Q&A

01

АІ у наукових публікаціях — це вже норма, а не експеримент

Журнали, редактори й видавці використовують ШІ системно з 2025–2026 рр.

02

Проблема не в АІ, а в неправильному використанні

Більшість відмов — через AI-abuse, hallucinations і відсутність авторського внеску.

03

Редактори бачать більше, ніж здається авторам

АІ-скринінг, citation-checks, image-forensics — працюють ДО читання статті.

04

Перемагають ті, хто використовує АІ як інструмент, а не як автора

Структура, ресерч, перевірка — так.

Інтерпретація, висновки, наукова позиція — лише людина.

05

Майбутнє — за гібридною моделлю

АІ автоматизує рутину,

людина відповідає за сенс, етику і новизну.

AI Research Assistant від «Академії статей»

Ваш персональний навігатор у світі наукових публікацій



Що може AI Research Assistant

-  Уся інформація про Академію статей
-  Консультація щодо ваших запитів
-  Підбірка видань категорії Б за вашим запитом
-  Аналіз ризиків журналів Scopus
-  Підбір теми для співавторства(Планується)
-  База знань Академії статей



Скануйте, щоб почати роботу з AI Research Assistant

Без реєстрації. Без спаму.
Доступ 24/7.

https://t.me/Articles_Academy_bot

Відповіді на питання



LinkedIn — особистий

LinkedIn — Василь Малець

Аналітика, публікації, кейси редакцій



Instagram — Академія статей

Instagram — Академія статей

Реальні кейси, помилки авторів, апдейти політик



Telegram — професійна спільнота

Telegram-група

Пропозиції публікацій, співавторства та знижки

