

**Міністерство освіти і науки України
Національний університет «Києво-Могилянська академія»
Факультет соціальних наук і соціальних технологій
Кафедра зв'язків з громадськістю**

**IV науково-практична конференція
«ОСОБЛИВОСТІ ТРАНСФОРМАЦІЇ КОМУНІКАЦІЙ В УМОВАХ
НОВІТНІХ СУСПІЛЬНИХ ВИКЛИКІВ»
(11 квітня 2025 року, м. Київ)**

Київ – 2025

Особливості трансформації комунікацій в умовах новітніх суспільних викликів. Матеріали IV науково-практичної конференції «Особливості трансформації комунікацій в умовах новітніх суспільних викликів», 11 квітня 2025р.,(м. Київ) / Наукове видання / Національний університет «Києво-Могилянська академія» [За ред. докт. соціол.н. Сусської О.О., докт. соціол.н. Щербини В.М., канд. соціол.н. Коника Д.Л., канд. соціол. н. Левцуна О.Г.]. Київ, НаУКМА, 2025. 74 с.

У збірнику представлені наукові доповіді учасників IV науково-практичної конференції «Особливості трансформації комунікацій в умовах новітніх суспільних викликів» (м. Київ, 11 квітня 2025 року).

Для наукових працівників, викладачів, аспірантів і студентів, які навчаються за спеціальністю 061. Журналістика (освітньо-наукова програма: «Зв'язки з громадськістю»), а також широкого кола читачів, хто цікавиться питаннями трансформації та удосконалення зв'язків з громадськістю, розвитком сучасної науки в галузі соціальних комунікацій.

Статті, включені до збірника, друкуються в авторській редакції.

відсутність стратегічного бачення в системі урядових комунікацій, надмірна централізація комунікації та наслідки обмеження свободи слова в довгостроковій перспективі тощо.

Джерела:

1. PERFORMANCE PURPOSE. Government Communication Service Strategy 2022–2025. URL: <https://strategy.gcs.civilservice.gov.uk/wp-content/uploads/2022/05/gcs-strategy-2022-25.pdf> (дата звернення: 2.04.2025)
2. Дзюбань О. П. Стратегічні комунікації: до проблеми осмислення сутності. *Міжнародні відносини: теоретико-практичні аспекти*. 2018. № 2. С. 254–264. URL: https://www.researchgate.net/publication/327922969_STRATEGICNI_KOMUNIKACII_DO_PROBLEMI_OSMISLENNIA_SUTNOSTI (дата звернення: 29.05.2025).
3. Дубов Д. В. Стратегічні комунікації: проблеми концептуалізації та практичної реалізації. *Політика. СТРАТЕГІЧНІ КОМУНІКАЦІЇ*. 2016. № 4 (41). С. 9-23. URL: <https://ippi.org.ua/sites/default/files/dubov.pdf>. (дата звернення: 29.05.2025).
4. Горбик Р., Дуцик Д., Шалайський С. Ефективність протидії російській дезінформації в Україні в умовах повномасштабної війни. Аналітичний звіт. – ГО «Український інститут медіа та комунікації». 2023. 66 с. URL: https://www.jta.com.ua/wp-content/uploads/2023/08/UMCI_Effectiveness-of-Russian-Disinformation-Counteration_UA.pdf?utm_source=chatgpt.com (дата звернення: 29.05.2025).
5. Соціологічний моніторинг «Українське суспільство» після 16 місяців війни. Інститут соціології НАН України. 2023. URL: https://www.kiis.com.ua/materials/pr/20230829_d/%D0%9F%D1%80%D0%B5%D0%B7%D0%B5%D0%BD%D1%82%D0%B0%D1%86%D1%96%D1%8F-%D0%BC%D0%BE%D0%BD%D1%96%D1%82%D0%BE%D1%80%D0%B8%D0%BD%D0%B3%D1%83-2023.pdf (дата звернення: 29.05.2025).
6. Оцінка громадянами ситуації в країні, довіра до соціальних інститутів, політико-ідеологічні орієнтації громадян України в умовах російської агресії (вересень–жовтень 2022р.). Центр Розумкова. URL: <https://razumkov.org.ua/napriamky/sotsiologichni-doslidzhennia/otsinka-gromadianamy-sytuatsii-v-kraini-dovira-do-sotsialnykh-instytutiv-politykoideologichni-orientatsii-gromadian-ukrainy-v-umovakh-rosiiskoi-agresii-veresen-zhovten-2022r> (дата звернення: 29.05.2025).
7. Оцінка ситуації в країні та діяльності влади, довіра до соціальних інститутів, політиків, посадовців та громадських діячів, віра в перемогу (вересень 2024р.) Центр Разумкова. URL: razumkov.org.ua/napriamky/sotsiologichni-doslidzhennia/otsinka-sytuatsii-v-kraini-ta-diialnosti-vlady-dovira-do-sotsialnykh-instytutiv-politykiv-posadovtsiv-ta-gromadskykh-diiachiv-vira-v-peremogu-veresen-2024r?utm_source=chatgpt.com (Дата звернення: 29.05.2025).

ІВАНОВ ВАЛЕРІЙ ФЕЛІКСОВИЧ,

докт. філол. н., проф., професор кафедри соціальних комунікацій Інституту журналістики
Київського національного університету ім. Тараса Шевченка,

ІВАНОВА ТЕТЯНА ВІКТОРІВНА,

докт. пед. н., професор, завідувачка кафедри соціальних комунікацій Маріупольського
державного університету,

ЄФРЕМОВА ОКСАНА ВОЛОДИМИРІВНА,

старший викладач кафедри соціальних комунікацій
Маріупольського державного університету

ШТУЧНИЙ ІНТЕЛЕКТ: МІФ ЧИ РЕАЛЬНІСТЬ АБСОЛЮТНОЇ ДОСТОВІРНОСТІ?

Штучний інтелект (ШІ) стрімко розвивається та проникає в усі сфери життя – від медицини до фінансів, від науки до творчості. Водночас постає питання: наскільки можна довіряти інформації, яку генерує ШІ? Суспільство часто сприймає ці технології як абсолютний авторитет, але чи є вони насправді непогрішними?

З одного боку, алгоритми обробляють величезні масиви даних, усуваючи людський фактор і мінімізуючи помилки. Проте вони залежать від якості вхідної інформації та налаштувань, створених людиною. Дезінформація, упередженість даних і технічні обмеження можуть впливати на точність результатів, що ставить під сумнів абсолютну достовірність ШІ.

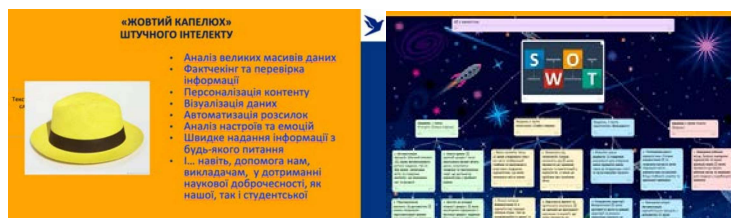
Актуальність цієї теми зростає, оскільки суспільство дедалі більше покладається на ШІ в ухваленні рішень. Важливо розуміти його можливості й обмеження, критично оцінювати отримані дані та працювати над удосконаленням алгоритмів для зменшення похибок і підвищення надійності штучного інтелекту.

Для нас, працівників галузі соціальних комунікацій та викладачів, які займаються підготовкою майбутніх журналістів, взаємодія із ШІ є дуже важливою з декількох причин:

- Етичне використання штучного інтелекту в журналістиці є критично важливим для збереження довіри аудиторії та забезпечення об'єктивності висвітлення подій.
- ШІ може допомагати у зборі та аналізі інформації, проте важливо уникати маніпуляцій, поширення дезінформації та автоматизованого створення матеріалів без людської перевірки.
- Журналісти мають відповідально ставитися до алгоритмів, забезпечуючи прозорість їхньої роботи та зберігаючи баланс між швидкістю обробки даних і достовірністю контенту.

І саме головне, використання ШІ не повинно підміняти журналістську етику, а навпаки – сприяти об'єктивності, фактчекінгу та захисту прав людини.

На кафедрі соціальних комунікацій Маріупольського державного університету у рамках навчальної дисципліни “Сучасні стандарти журналістики та імперативи у соціальних комунікаціях” приділяється велика увага питанням усвідомлення наслідків неетичного використання ШІ для журналістів. Студенти у процесі семінарських занять виконують різноманітні вправи на системне розуміння функціонування штучного інтелекту. Вкрай популярними у студентів є виконання завдань за методом “Шість капелюхів” Едварда де Боно, а також за методикою SWOT-аналіз. Наведені скріншоти показують результат виконання цих вправ.



Тобто є зрозумілим, що інтеграція штучного інтелекту в освітні процеси відкриває значні можливості для дотримання норм академічної доброчесності, автоматизації рутинних завдань, створення персоналізованих навчальних траєкторій, а також спрощення оцінювання результатів навчання. Однак разом із перевагами виникають ризики, пов'язані з використанням текстів, згенерованих ШІ, які важко відрізнити від авторських. Тому виникає нагальна потреба у розробці ефективних методів верифікації таких текстів. Це питання є актуальним особливо для освітніх установ, де контроль академічної доброчесності є першочерговим завданням. Головними питаннями залишаються: *«А чи дійсно є можливим верифікувати тексти, згенеровані ШІ?»*, *«Чи може процес верифікації досягти абсолютної достовірності?»*.

Оцінити роботу перевірки текстів згенерованих ШІ авторами цих тез вдалося під час експерименту з використанням безкоштовних сервісів перевірки.

Хід експерименту

1. Створено три тексти (авторський, згенерований та скопійований).

Авторський текст, написаний доктором педагогічних наук, професором Тетяною Івановою, без використання ШІ.

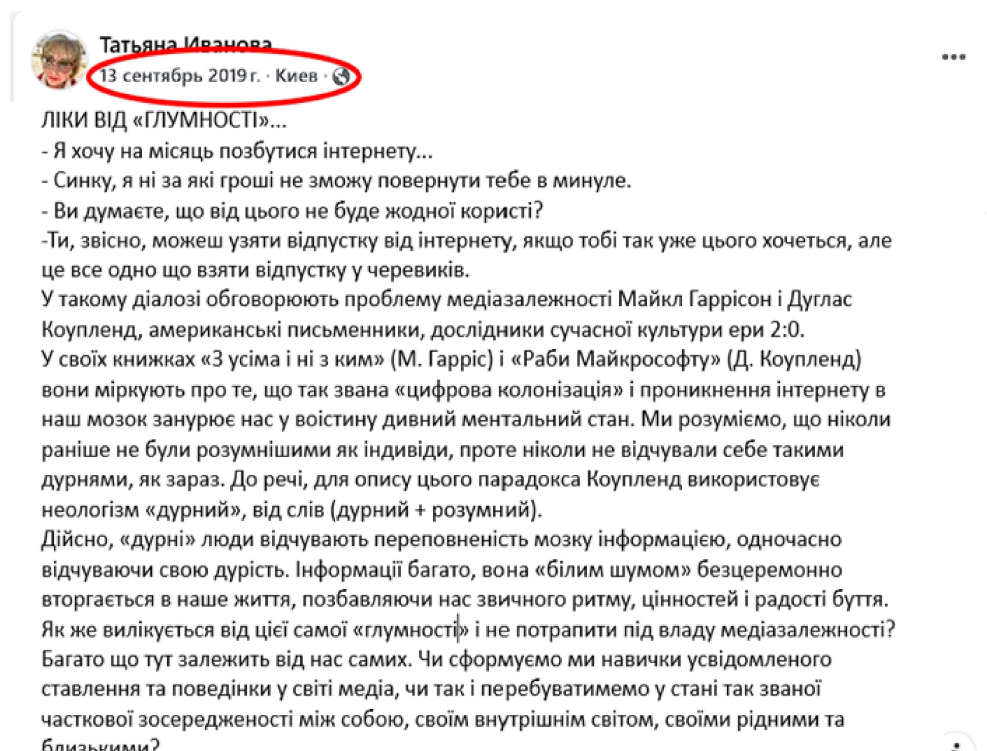


Рисунок 1. Авторський текст

Текст, згенерований штучним інтелектом (на основі сучасних генеративних моделей

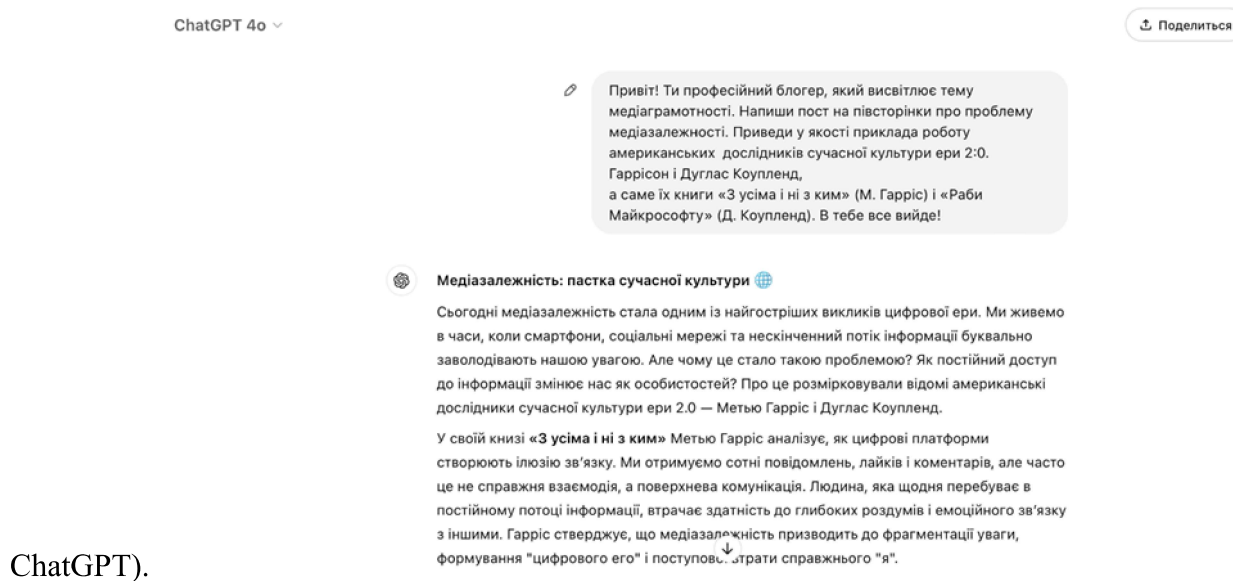


Рисунок 2. Текст згенерований ChatGPT

Скомпільований текст мав комбіновану структуру, що складалася з двох частин: перша половина тексту була взята з оригінального авторського матеріалу проф. Тетяни Іванової, а друга – згенерована за допомогою штучного інтелекту.

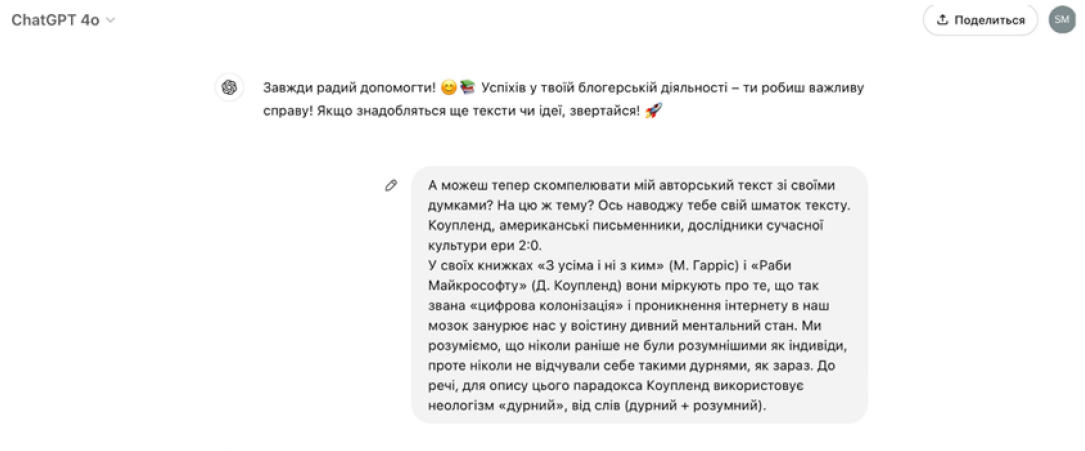
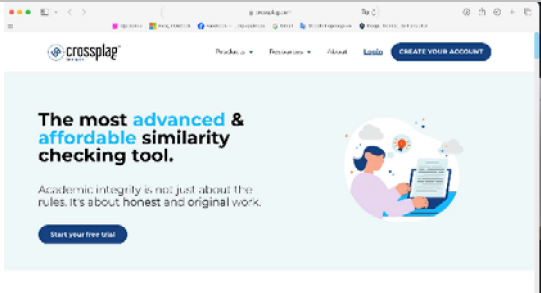
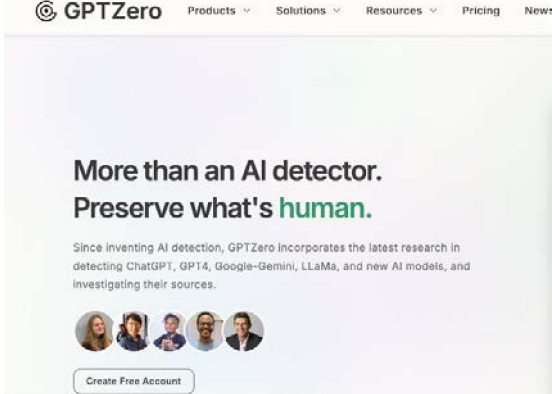


Рисунок 3. Скомпільований текст

2. Для перевірки застосовано три сервіси (безкоштовних):

1. Crossplag (http://crossplag.com): інструмент для перевірки унікальності текстів.	
2. GPTZero (http://gptzero.me): сервіс для виявлення текстів, створених ШІ.	
3. Writer AI Content Detector (https://writer.com/ai-content-detector/): інструмент для визначення ймовірності штучного походження тексту.	

1. Crossplag (<http://crossplag.com>): інструмент для перевірки унікальності текстів.
2. GPTZero (<http://gptzero.me>): сервіс для виявлення текстів, створених ШІ.
3. Writer AI Content Detector (<https://writer.com/ai-content-detector/>): інструмент для визначення ймовірності штучного походження тексту.

3. Результати перевірки

Жоден із протестованих сервісів не продемонстрував високої точності чи здатності надати коректну відповідь щодо визначення авторства текстів. Результати перевірок виявили суттєві розбіжності залежно від використаного інструмента, а в окремих випадках надана інформація була суперечливою або недостатньо чіткою. Сервіси часто демонстрували збої в роботі, що підкреслює їхню неточність і ненадійність (усі мають точність нижче 80%), також їхня ефективність ще більше погіршується при використанні методів додаткової обробки, таких як ручне редагування або машинне перефразування.

Також було виявлено, що вони діагностують документи, написані автором, як створені штучним інтелектом (див. рис. 4) і часто діагностують тексти, створені ШІ, як написані людиною (див. рис. 5). Результати перевірки скопійованого тексту за допомогою сервісу GPTZero

показали 50% людський, що є коректним результатом, оскільки текст насправді складався на 50% з авторського матеріалу (див. рис. 6).

Інструменти виявлення здебільшого схильні класифікувати результати як написані людиною, а не виявляти вміст, створений штучним інтелектом. Загалом, приблизно 50% текстів, згенерованих ШІ, які зазнали певної обробки, ймовірно, будуть помилково приписані людині [7, с. 26]. Оскільки інструменти не надають жодних доказів, ймовірність того, що навчальний заклад зможе довести факт порушення академічної доброчесності, вкрай низька.

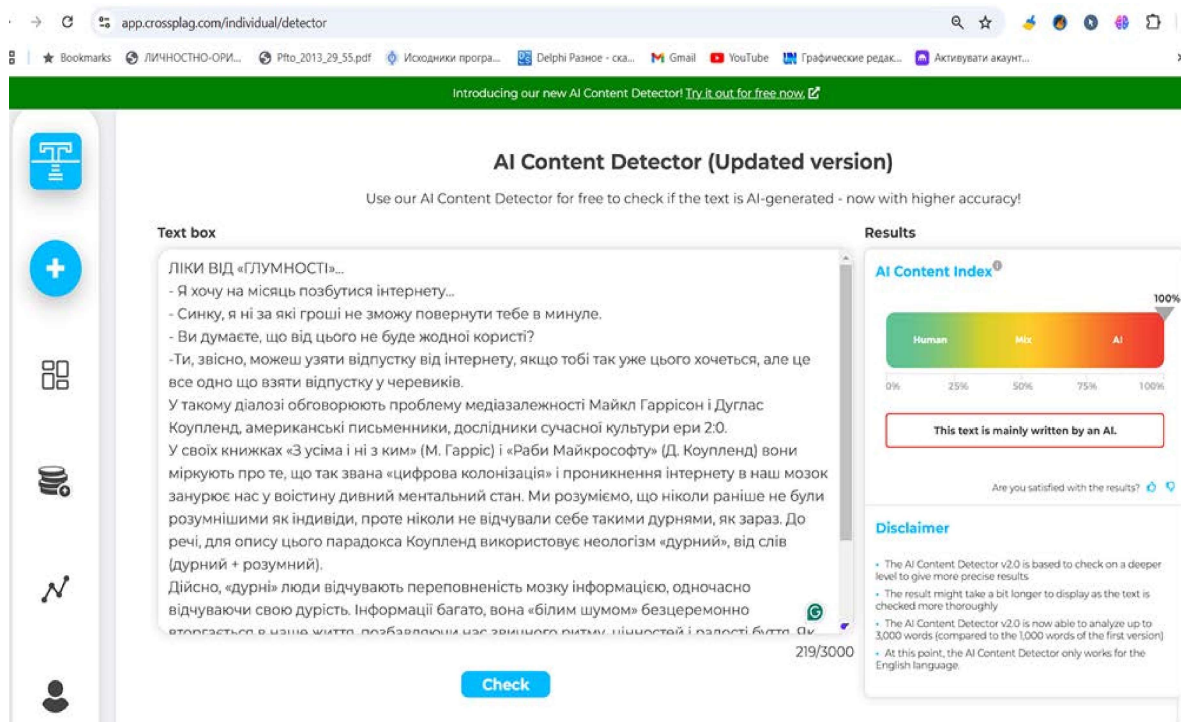


Рисунок 4. Результати перевірки авторського тексту сервісом Crossplag: помилково визначено як 100% ШІ.

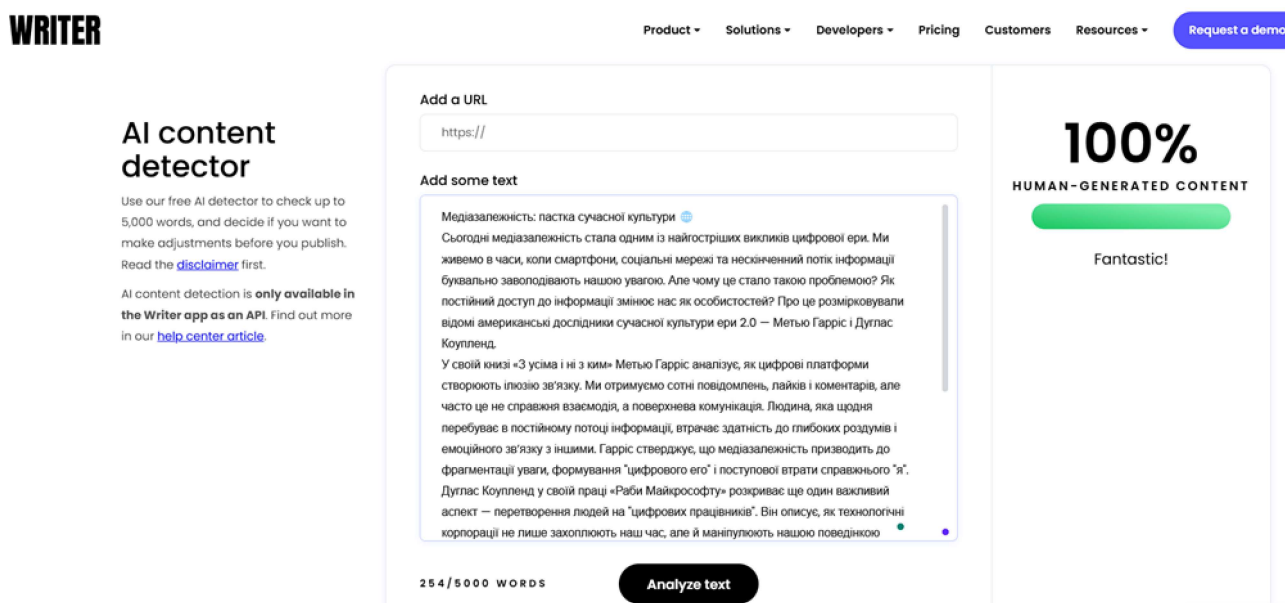


Рисунок 5. Результати перевірки тексту, згенерованого ChatGPT, сервісом Writer: помилково визначено як 100% людський

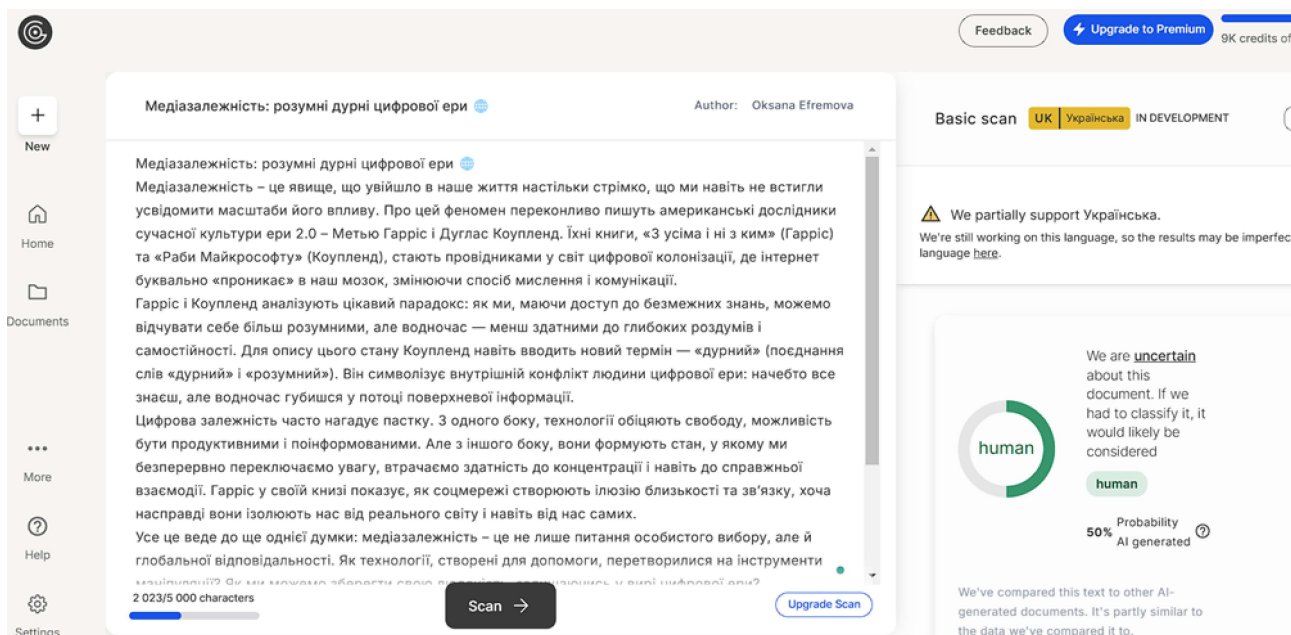


Рисунок 6. Результати перевірки скопійованого тексту сервісом GPTZero: точно визначено як 50% людський

4. Аналіз результатів

Результати дослідження показали, що існуючі моделі перевірки текстів перебувають на стадії розвитку та вдосконалення.

Основними причинами помилок є:

- недостатня точність алгоритмів розпізнавання;
- використання різних наборів даних для навчання моделей
- відсутність єдиних критеріїв для класифікації текстів.

Загальні висновки.

Штучний інтелект, незважаючи на свою здатність до швидкої обробки великих обсягів даних, не гарантує абсолютної достовірності інформації. Він залишається залежним від джерел даних, алгоритмічних рішень та людського втручання. Отже, повна довіра автоматизованим системам без перевірки фактів є ризикованим як для журналістів, так і для всіх споживачів інформації.

Прозорість алгоритмів та відповідальність. Для забезпечення етичного використання ШІ у журналістиці необхідно гарантувати прозорість алгоритмів та їх вплив на кінцевий зміст. Розробники та журналісти повинні нести відповідальність за застосування ШІ у процесі підготовки матеріалів. ШІ може сприяти як ефективному збору та аналізу інформації, так і її маніпуляції. Використання технологій глибокого навчання для створення дезінформації (наприклад, дипфейків) є серйозною загрозою для довіри до медіа.

На основі проведеного експерименту, який був спрямований на здатність ШІ перевіряти авторство текстів, було також встановлено, що жоден із обраних інструментів не продемонстрував ВИСОКОЇ ТОЧНОСТІ У ВИЗНАЧЕННІ АВТОРСТВА ТЕКСТУ. Це свідчить про те, що технології перевірки ще перебувають на стадії розвитку та мають суттєві обмеження. Результати експерименту засвідчили необхідність критичного підходу до використання даних технологій у викладацькій практиці.

Таким чином, дослідження підкреслює, що використання ШІ у журналістиці та навчальному процесі підготовки журналістів потребує обережного підходу, який враховує етичні стандарти, прозорість та відповідальність. Лише за таких умов можливо забезпечити надійність інформації та зберегти довіру аудиторії.

Джерела:

1. Антіпова К. О. Виклики розпізнавання текстів, згенерованих штучним інтелектом Штучний інтелект у вищій освіті: ризики та перспективи інтеграції: матеріали всеукраїнського науково-педагогічного підвищення кваліфікації, 1 липня – 11 серпня 2024 року. – Львів – Торунь : Liha-Pres, 2024. С. 14-19.
2. Jawahar G., Abdul-Mageed M., Lakshmanan, L. V. Automatic Detection of Machine Generated Text: A Critical Survey. The 28th International Conference on Computational Linguistics (COLING). 2020. DOI: <https://doi.org/10.48550/arXiv.2011.01314>
3. Bakhtin A., Gross S., Ott M., Deng Y. et al. Real or Fake? Learning to Discriminate Machine from Human Generated Text. 2019. DOI: <https://doi.org/10.48550/arXiv.1906.03351>
4. Ippolito D., Duckworth D., Callison-Burch C., Eck D. et al. Human and Automatic Detection of Generated Text. 2019. DOI: <https://doi.org/10.48550/arXiv.1911.00650>
5. Krishna K., Song Y., Karpinska M., Wieting J. et al. Paraphrasing evades detectors of AI-generated text, but retrieval is an effective defense. Advances in Neural Information Processing Systems, 36. 2023. DOI: <https://doi.org/10.48550/arXiv.2303.13408>
6. Wolff M. Attacking Neural Text Detectors. International Conference on Learning Representations (ICLR). 2020. DOI: <https://doi.org/10.48550/arXiv.2002.11768>
7. Weber-Wulff D., Anohina-Naumeca A., Bjelobaba S., Foltýnek T. et al. Testing of detection tools for AI-generated text. International Journal for Educational Integrity, 19. 2023. DOI: <https://doi.org/10.1007/s40979-023-00146-z>