

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В МЕДІА

Козир О.О.,
аспірант кафедри права інтелектуальної власності
та патентної юстиції,
Національний університет «Одеська юридична академія»

Козир О.О. Використання штучного інтелекту в медіа.

Стаття присвячена дослідженню використання штучного інтелекту в медіа. Серед чисельних викликів та загроз, які несе штучний інтелект в медіа, виділено такі. По-перше, поширення недостовірної інформації. Штучний інтелект може генерувати інформацію, новини, різні інформаційні повідомлення на підставі аналізу інформації, розміщеної у вільному доступі в мережі Інтернет. Така інформація може бути як результатом сформульованого запиту щодо написання фейку, так і згенерованим текстом на поставлені питання. По-друге, ризики проявів дискримінації в інформаційних повідомленнях, згенерованих штучним інтелектом, через відсутність встановлених фільтрів чи обмежень або некоректності сформованого запиту. По-третє, зниження професійних стандартів та практик формування журналістських матеріалів, через заміну комунікації та пошуку інформації на її генерування ШІ. По-четверте, ризики несанкціонованого доступу до конфіденційної інформації, персональних даних, які можуть завантажуватись штучним інтелектом під час підготовки різних матеріалів для медіа.

Зазначено, що впровадження стандартів етичного використання штучного інтелекту дозволить сучасним медіа досягти якісних результатів при мінімізації негативних наслідків. Натомість звернення до штучного інтелекту з метою створення та поширення діпфейків може нести ризики як для окремих користувачів, так і національної безпеки в цілому. Наведено результати досліджень, які свідчать з однієї сторони про важливість позначення матеріалу, згенерованого ШІ, а також вплив таких позначень на сприйняття матеріалу користувачами.

Обґрунтовано, що відповідальне використання штучного інтелекту в медіа має передбачати розробку і впровадження етичних принципів кожним суб'єктом медіа; відкриту політику щодо позначення матеріалів, згенерованих ШІ; перевірку матеріалів, підготовлених ШІ на їх достовірність та відповідність політикам суспільної моралі, недискримінації, нешкідливості для суспільства та окремих категорій користувачів; мінімізація завантаження конфіденційної інформації для подальшого опрацювання ШІ.

Ключові слова: медіа, інформація, штучний інтелект, генерування тексту, етичні принципи, діпфейк, безпека, новини, персональні дані, захист.

Kozyr O.O. Artificial intelligence in media.

This article is devoted to the study of the use of artificial intelligence (AI) in the media. Among the numerous challenges and threats posed by AI in this field, several key issues stand out. First, the spread of misinformation. AI systems are capable of generating information, news stories, and various informational messages based on the analysis of publicly available data from the internet. Such content may result from deliberately crafted prompts to produce disinformation (i.e., fake news), or it may be generated in response to neutral queries but still contain false or misleading information. Second, the risk of discriminatory content in AI-generated messages due to the absence of established filters, safeguards, or insufficiently precise user input. This may lead to biased or harmful narratives being produced and disseminated. Third, the erosion of professional standards and journalistic practices. As AI increasingly replaces traditional journalistic processes such as communication, fact-checking, and field investigation with automated content generation, there is a danger of diminishing the integrity and quality of media work. Fourth, the risk of unauthorized access to confidential or personal data, which AI systems may inadvertently upload or process while preparing materials for media use.

The article emphasizes that the implementation of ethical standards for the use of AI will enable modern media organizations to achieve high-quality outcomes while minimizing harmful consequences. Conversely, the use of AI for generating and distributing deepfakes poses significant risks both to individual users and to national security at large. The article presents research findings that highlight, on the one hand, the importance of labeling AI-generated content, and on the other hand, the impact of such labeling on user perception.

It is argued that responsible use of AI in media should include the development and adoption of ethical guidelines by each media entity; a transparent policy for identifying AI-generated content; fact-checking of AI-produced materials to ensure their accuracy and alignment with societal moral standards, non-discrimination principles, and non-harm to the public or specific user groups; the minimization of uploading confidential information for AI processing.

Key words: media, information, artificial intelligence, text generation, ethical principles, deepfake, security, news, personal data, protection.

Постановка проблеми. Штучний інтелект (ШІ) широко використовують у всіх сферах суспільного життя. Не є виключенням сучасні медіа, в яких ШІ може нести як свої переваги, так і чисельні ризики. Використання ШІ дозволяє швидше знаходити інформацію, робити публікації більш яскравими та творчими. Однак здатність ШІ робити зображення вкрай реалістичними породжує ризики поширення через медіа недостовірної інформації, що може вводити в оману користувачів інформаційних продуктів.

За даними опитування Інституту масової інформації (ІМІ), лише 22% представників українських медіа повідомили, що їхня редакція на постійній основі використовує хоча б один інструмент ШІ, 30% опитаних зазначили, що в їхній редакції зрідка залучають ШІ, водночас 20% журналістів повідомили, що не використовують інструменти ШІ у своїй роботі. Крім того, 16% опитаних зазначили, що раніше користувалися інструментами ШІ, але наразі припинили [1; 2].

Тому дослідження сучасних практик використання ШІ в медіа має важливе практичне значення, оскільки дозволить визначити ризики, окреслити напрями правового регулювання відносин у зазначеній сфері.

Метою статті є визначення основних напрямів, переваг та ризиків використання ШІ у сфері медіа.

Стан опрацювання проблематики. Слід відзначити, що питанням використання ШІ присвячені чисельні наукові дослідження. Зокрема питання визначення правової природи ШІ, в тому числі у сфері права інтелектуальної власності, досліджено у працях Г. Андрощука, О. Заярного, О. Тараненко, О. Харитонової, Є. Харитонова, Є. Якубівського.

Загрози та ризики використання ШІ досліджуються такими вченими, як М. Ворохоб, М. Гуменюк, О. Скіцько, П. Складанний, Р. Ширшов. До ризиків віднесено посилення кібератак, утворення каналів витоку інформації, атаки отруєння інформації, містифікація даних [3]. У дослідженнях також виділяють ступені ризику ШІ, які впливають на можливість встановлення заборони на його використання, а саме: неприпустимий, високий, обмежений ризик [4].

Чимало досліджень присвячено питанням використання ШІ в різних галузях – освіті, медицині, юриспруденції. Не є виключенням і сфера медіа, де ШІ отримує широке застосування, може надавати як переваги, так і нести чисельні ризики, що обумовлює актуальність проведення дослідження.

Виклад основного матеріалу. Як слушно зазначає Ситник О., штучний інтелект стає все більш актуальним у медіаіндустрії, оскільки він пропонує значні переваги, зокрема, підвищена ефективність, персоналізація, аналіз даних, створення контенту та контроль якості. Оскільки технологія штучного інтелекту продовжує розвиватися, вона, ймовірно, матиме ще більший вплив на медіаіндустрію та спосіб споживання та взаємодії з медіаконтентом [5, с. 263].

Враховуючи потенційні ризики, які можуть бути обумовлені використанням ШІ, на міжнародному та національному рівні розробляються принципи та політики відповідального використання ШІ в медіа.

У Керівних принципах щодо відповідального використання ШІ в журналістиці, прийнятих 30 листопада 2023 року Керівним комітетом з питань медіа та інформаційного суспільства (CDMSI), зокрема визначено, що система штучного інтелекту в журналістиці – це система, що безпосередньо використовується в бізнесі або практиці регулярного виробництва інформації про сучасний

стан справ, що становлять суспільний інтерес, і покладена в основу журналістських матеріалів. Журналістські системи штучного інтелекту – це не єдина технологія, а низка різних, часто взаємопов'язаних інструментів для автоматизації конкретних завдань. Рішення про застосування ШІ має бути не просто технологічним чи комерційним, а повинно відповідати місії медіа. Це рішення є редакційним, оскільки воно має вирішальне значення для реалізації редакційної місії та професійних цінностей медіа. Крім цього, в медіа має бути фахівець, який відповідає за впровадження та результати використання журналістського ШІ [6].

В Рекомендаціях щодо відповідального використання ШІ в медіа виділено важливі принципи такого використання, серед яких відповідальне редакційне рішення; законність; регулярне проведення оцінки ризиків, пов'язаних з використанням систем ШІ; прозорість і зрозумілість; поінформованість аудиторії про використання ШІ та про природу поширеного контенту; чітке розмежування автентичного та згенерованого ШІ контенту тощо); конфіденційність і захист даних; різноманітність; фаховий людський контроль; відповідальність; адаптивність [7].

Медіа-організації, як правило, залишаються стриманими й уникають надмірної конкретики щодо використання технологій ШІ, зокрема технологій на основі штучного інтелекту. Деякі з них оприлюднили спеціальні настанови щодо використання соціальних мереж (наприклад, *The Guardian*, *The Washington Post* або *BBC*). Інші – розкрили редакційні стандарти, які регулюють використання цифрового контенту (наприклад, *Al Jazeera* та *Associated Press*) [8].

На підставі дослідження 40 документів, серед яких 35 – це спеціальні рекомендації щодо використання ШІ та 5 – загальні етичні кодекси, які містять хоча б один розділ, присвячений ШІ, створених 84 медіаорганізаціями з усього світу, встановлено наступне. Європейські країни мають найбільшу кількість медіаорганізацій, які впровадили етичні рекомендації щодо ШІ, за ними йде Північна Америка. Серед окремих країн (без урахування міжнародних організацій), США та Велика Британія лідирують за кількістю редакцій із такими кодексами. Наступними йдуть Канада, Німеччина та Іспанія [9].

Виклики, які постають у сфері використання ШІ, є спільними для всіх медіа, тому розробка певних рекомендацій та настанов мають загальні засади. З урахуванням позитивних прикладів управління використанням ШІ визначено принципи формування етичної стратегії використання ШІ в медіа індустрії.

По-перше, формулювання чітких цілей та гнучких підходів щодо різних практик впровадження ШІ. Залучення представників різних структурних підрозділів медіа до обговорення політик впровадження ШІ.

По-друге, забезпечення системи підтримки та контролю за впровадженням етики використання ШІ, залучення керівників структурних підрозділів до реалізації політик ШІ.

По-третє, розвиток культурного та інтелектуального розмаїття, включення до груп розробки етичних політик ШІ широкого кола учасників (дослідників аудиторії, продакт-менеджерів, юристів).

По-четверте, широка комунікація на всіх рівнях щодо впровадження етичного використання ШІ.

По-п'яте, прозорість та оцінка досягнень у сфері впровадження етичного використання ШІ, відкрите обговорення помилок та невдач, складнощів та шляхів їх подолання [10].

Розробка політик та етичних стандартів використання ШІ має важливе практичне значення, оскільки безконтрольне використання ШІ може мати негативні наслідки як для медіа, так і користувачів.

Серед чисельних викликів та загроз, які несе ШІ в медіа, слід виділити такі.

По-перше, поширення недостовірної інформації. ШІ може генерувати інформацію, новини, різні інформаційні повідомлення на підставі аналізу інформації, розміщеної у вільному доступі в мережі Інтернет. Така інформація може бути як результатом сформульованого запиту щодо написання фейку, так і згенерованим текстом у відповідь на поставлені питання.

Маніпулятивні методи на основі ШІ можуть використовуватися для того, щоб переконати людей вдаватися до небажаної поведінки або обдурити їх, підштовхуючи до рішень таким чином, що це підриває та порушує їхню автономію, прийняття рішень та вільний вибір. Розміщення на ринку, введення в експлуатацію або використання певних систем ШІ з метою суттєвого впливу на людську поведінку, в результаті чого можливо завдати значної шкоди, зокрема, мати достатньо суттєвий негативний вплив на фізичне, психологічне здоров'я або фінансові інтереси, є особливо небезпечними і тому повинні бути заборонені. Такі системи ШІ використовують підсвідомі компоненти, зокрема, аудіо, зображення, відеостимулятори, які люди не можуть сприймати, оскільки

ці стимули знаходяться поза людським сприйняттям, або інші маніпулятивні чи оманливі методи, які впливають на самостійність прийняття рішень або вільний вибір людини таким чином, що люди не усвідомлюють цих методів, або, якщо вони усвідомлюють їх, все одно можуть бути введені в оману або не можуть контролювати їх [11].

У дослідженнях, присвячених потенційному впливу дипфейків на медіа, виділяють можливі позитивні сторони та негативні наслідки їх поширення. Загрози дипфейків розглядаються через їх негативну психологічну, фінансову та соціальну шкоду, яка може завдаватись у випадку поширення недостовірної згенерованої інформації. Психологічна шкода може завдаватись через використання дипфейків з метою цькування, маніпулювання, шантажу, залякування, наклепу. Фінансова шкода полягає в тому, що дипфейки можуть використовуватись під час шахрайства з особистими даними, проведення несанкціонованих транзакцій, отримання доступу до банківських рахунків, поширення фейкових повідомлень про банкрутство або злиття компаній. Соціальна шкода може завдаватись через маніпулювання громадською думкою, фейкові відео або заяви від імені відомих публічних осіб, фальсифікацію доказів у судових справах та поширення недостовірної інформації про діяльність судів, що може породжувати сумніви щодо справедливості правосуддя [12 с. 2163-2164].

По-друге, ризики проявів дискримінації в інформаційних повідомленнях, згенерованих ШІ, через відсутність встановлених фільтрів чи обмежень або недостатню коректність сформованого запиту.

По-третє, зниження професійних стандартів та практик формування журналістських матеріалів, через заміну комунікації та пошук інформації на її генерування ШІ.

По-четверте, ризики несанкціонованого доступу до конфіденційної інформації, персональних даних, які можуть завантажуватись ШІ під час підготовки різних матеріалів для медіа.

В Рекомендаціях Комісії з журналістської етики щодо використання штучного інтелекту для створення журналістських матеріалів зазначається про те, що існує небезпека витоку чутливих персональних даних під час діалогу журналіста та генеративної моделі. Журналісти не можуть знати, як обробляється їхній запит, хто має доступ до даних, які містяться в запиті, і як ці дані зберігатимуться в майбутньому. Тому Комісія не радить ділитися зі штучним інтелектом будь-якою інформацією, яка містить персональні дані [13].

Заярний О., Деркаченко Ю. сформулювали такі висновки та рекомендації щодо обробки персональних даних з використанням алгоритмічних систем, зокрема, ChatGPT: 1. Обробка персональних даних повинна здійснюватись з дотриманням підстав її проведення та на основі добровільної, поінформованої згоди суб'єкта персональних даних за умови додержання легітимної та пропорційної мети обробки. 2. Важливою умовою забезпечення правомірної обробки персональних даних з використанням ChatGPT є дотримання розробниками проєктних рішень, де використовується відповідна технологія безпеки персональних даних за проєктуванням і за замовчуванням. 3. Обґрунтованим вбачається прийняття нової редакції Закону України «Про захист персональних даних» та імплементація Модернізованої редакції Конвенції Ради Європи «Про захист осіб у зв'язку з автоматизованою обробкою персональних даних» (Конвенції 108+). 4. Доцільним також вбачається прийняття окремого ДСТУ «Технічний захист інформації. Забезпечення права на невтручання у приватне і сімейне життя при застосуванні алгоритмічних систем та ботів» [14, с. 61].

Враховуючи широке використання ШІ в медіа, для забезпечення його доброчесного використання важливо забезпечувати належне маркування згенерованого матеріалу. Позначення згенерованого матеріалу дозволить забезпечити відкриті політики використання ШІ в медіа, уникнути введення користувачів в оману та протидіяти поширенню фейкової інформації.

Маркування контенту, згенерованого ШІ, на сьогодні не має єдиних підходів. Серед основних питань, які постають під час позначення згенерованого матеріалу ШІ, слід виділити, по-перше, який саме матеріал слід позначати – лише той, що повністю згенерований ШІ, чи будь-який матеріал, під час створення якого використовувались технології ШІ; по-друге, коли і в якій формі має зазначатись інформація про використання ШІ (на початку матеріалу, наприкінці, чи виділяти лише згенеровані фрагменти); по-третє, використовувати текстові позначення чи спеціальні знаки; визначати сам факт використання ШІ чи конкретизувати який саме.

Trusting News було проведено опитування щодо використання ШІ в медіа у липні та серпні 2024 року за участю 10 партнерських новинних організацій, які зібрали понад 6000 відповідей на опитування від своєї аудиторії. Партнерськими організаціями були 100 Days in Appalachia, American City Business Journals, Houston Landing, Iowa Public Radio, KXAN -TV з Остіна, штат Техас, Science News, The Associated Press, The Texas Tribune, USA TODAY Network та Voice Media Group.

Серед респондентів 94% сказали, що хочуть, щоб редакції загалом розкривали інформацію про використання ШІ, тоді як 87% зазначили що розкриття інформації про ШІ повинно містити причини, чому журналісти використовували ШІ, 92% сказали, що хотіли б знати, що до перевірки інформації, створеної ШІ, була залучена людина [15].

Мета такого маркування полягає не лише у забезпеченні відкритої політики використання ШІ в медіа, а формуванні довіри користувачів до медіа. Натомість дослідження свідчать, що наявність або відсутність позначення не завжди суттєво впливає на сприйняття інформації, поширеної в медіа. В США було проведено дослідження за участю 1500 респондентів, які формували свої погляди щодо запропонованих політичних повідомлень за 4 темами. При цьому використовувались три підходи щодо надання інформації: з маркуванням, що текст згенеровано ШІ; зазначенням, що текст створено людиною; та без будь-яких позначень. За результатами дослідження встановлено, що додавання позначок змінює уявлення людей про те, чи був контент створений ШІ чи людиною – що свідчить про підвищення прозорості. Проте немає доказів, що ці позначки істотно впливають на переконливість самого контенту. Серед учасників, яким вказали, що повідомлення створене ШІ, 94,6% повірили в це. Серед тих, кому сказали, що текст написала людина, 89,3% повірили. Серед учасників, які не бачили жодної позначки, 39% вважали, що повідомлення створено людиною, 31% – ШІ, а 30% не були впевнені [16].

Висновки. ШІ став невід’ємною частиною сучасних медіа. Ефективний помічник – аналітик, перекладач, генератор творчих робіт та цікавих ідей, ШІ несе чималі ризики як для збереження якісного медіа, так і безпеки користувачів. Ризики чи переваги від використання ШІ, в першу чергу, залежать від мети та політики етичності та професійності медіа. Впровадження стандартів етичного використання ШІ дозволить сучасним медіа досягти якісних результатів при мінімізації негативних наслідків. Натомість звернення до ШІ з метою створення та поширення дипфейків може нести ризики як для окремих користувачів, так і національної безпеки в цілому.

Відповідальне використання ШІ в медіа має передбачати розробку і впровадження етичних принципів кожним суб’єктом медіа; відкрити політику щодо позначення матеріалів, згенерованих ШІ; перевірку матеріалів, підготовлених ШІ на їх достовірність та відповідність політикам суспільної моралі, недискримінації, нешкідливості для суспільства та окремих категорій користувачів; мінімізацію завантаження конфіденційної інформації для подальшого опрацювання ШІ.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ:

1. Мажкова Яна Медіа та штучний інтелект: Як українські медіа використовують штучний інтелект. 01 липня 2024 року. URL: <https://imi.org.ua/monitorings/ukrayinski-media-ta-shtuchnyj-intelekt-yak-redaktsiyi-zaluchayut-shidlya-stvorenniya-kontentu-i62217>.
2. Інформаційна довідка щодо використання штучного інтелекту в медіа: потенціал і ризики. *Дослідницька служба Верховної Ради України*. <https://research.rada.gov.ua/uploads/documents/33465.pdf>.
3. Скіцько О., Складанний П., Ширшов Р., Гуменюк М., Ворохоб М. Загрози та ризики використання штучного інтелекту. *Кібербезпека: освіта, наука, техніка*. 2023. № 2 (22). С. 6–18.
4. Sláva Gracová, Martin Graca Threats to the use of artificial intelligence and its legislative. *Вісник Національного університету «Львівська політехніка»: журналістика*. 2024. № 1 (7). С. 72–78.
5. Ситник О. Проблематика впровадження штучного в сучасних ЗМІ та медіатехнологіях. *Український інформаційний простір*. 2023. № 2 (12). С. 252–265.
6. Guidelines on the responsible implementation of artificial intelligence systems in journalism URL: <https://rm.coe.int/cdmsi-2023-014-guidelines-on-the-responsible-implementation-of-artific/1680adb4c6>.
7. Авдеева Т., Андрієнко О., Дубно О. Рекомендації щодо відповідального використання ШІ в медіа. 11.12.2024. https://webportal.nrada.gov.ua/wp-content/uploads/2024/12/03_Avdyeveva_Prezentatsiya_11.12_SHI-ta-media.pdf.
8. AI and Media Ethics Existing References Overview. *Reporters without borders*. https://rsf.org/sites/default/files/medias/file/2023/09/AI%20and%20media%20ethics%20existing%20references%20overview_0.pdf?utm_source=chatgpt.com.

9. Sonia Parratt-Fernández, Sonia Parratt-Fernández, Victoria Moreno-Gil Journalistic AI Codes of Ethics: Analyzing Academia's Contributions to their Development and Improvement. *Profesional de la información*. 2024. v. 33, n. 6. P. 1–15. <https://doi.org/10.3145/epi.2024.0602>.
10. Yves Bergquist AI ethics a framework for the media industry. *Entertainment technology center*. 14 p. https://www.etccenter.org/wp-content/uploads/2022/04/AI-Ethics-A-Primer-for-the-Media-Industry-ETC.pdf?utm_source=chatgpt.com.
11. Document 32024R1689 Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act). <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>.
12. Ebba Lundberg, Peter Mozelius The potential effects of deepfakes on news media and entertainment. 2024. *AI & SOCIETY* (2025). 40. P. 2159–2170. <https://doi.org/10.1007/s00146-024-02072-1>.
13. Рекомендації Комісії з журналістської етики щодо використання штучного інтелекту для створення журналістських матеріалів <https://cje.org.ua/statements/rekomendatsii-komisii-z-zhurnalistskoi-etyky-shchodo-vykorystannia-shtuchnoho-intelektu-dlia-stvorennia-zhurnalistskykh-materialiv>.
14. Заярний О.А., Деркаченко Ю.В. Деякі особливості обробки персональних даних при використанні чат-ботів зі штучним інтелектом на прикладі Chat GPT. *Юридичний бюлетень*. 2023. Вип. 29. С. 55–62.
15. Clark Merrefield What U.S. audiences want newsrooms to disclose about their AI use: 4 insights from Trusting news. *The Journalist's Resource. Informing the news*. 2025. https://journalistsresource.org/media/ai-use-news-what-audiences-disclose/?utm_source=chatgpt.com.
16. Isabel O. Gallegos, Chen Shani, Weiyan Shi, Federico Bianchi, Izzy Gainsburg, Dan Jurafsky, and Robb Willer Labeling AI-Generated Content May Not Change Its Persuasiveness. *Policy Brief HAI Policy & Society*. July 2025. <https://hai.stanford.edu/assets/files/hai-policy-brief-labeling-ai-generated-content.pdf>.