

УДК: 316.77:004.8

ЕТИКА ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В МАСОВІЙ КОМУНІКАЦІЇ

В. І. Сабат¹, В. В. Мацюк²

¹ Національний університет «Львівська політехніка»,
вул. С. Бандери, 12, Львів, 79013, Україна

² ВСП «Львівський поліграфічний фаховий коледж УАД»,
площа Митна 1, Львів, 79000, Україна

У статті досліджено ключові етичні виклики, що виникають у зв'язку із впровадженням штучного інтелекту (ШІ) у сферу масової комунікації. Розглянуто сучасні технології генеративного ШІ, які активно використовуються для створення медіаконтенту, зокрема у журналістиці, рекламі та політичній комунікації. Особливу увагу приділено питанням маніпулятивного впливу, розмитості авторства, поширення дезінформації та відсутності прозорості у використанні ШІ-алгоритмів. Проаналізовано реальні кейси, пов'язані з використанням deepfake-технологій у виборчих кампаніях, а також таргетованої політичної реклами. В роботі наголошується нагальна проблема етичного регулювання використання штучного інтелекту, особливо в галузях, пов'язаних з соціальними комунікаціями, поширенням рекламних оголошень через засоби масових комунікацій та мережні технології, а також вказується на необхідність навчання користувачів мережі Інтернет та подальших досліджень цифрової етики, права та медіаграмотності.

Ключові слова: штучний інтелект, інформаційні маніпуляції, масова комунікація, дезінформація, медіаграмотність.

Постановка проблеми. Упродовж останнього десятиліття штучний інтелект (ШІ) перетворився з технологічної інновації на потужний інструмент, що активно проникає у всі сфери суспільного життя, зокрема — у сферу масової комунікації. Завдяки здатності обробляти великі обсяги даних, генерувати тексти, зображення, відео й навіть голос, ШІ відкриває нові можливості для журналістики, маркетингу, реклами, соціальних медіа та цифрових платформ. Його застосування обіцяє підвищення ефективності, персоналізацію контенту та створення інноваційних форматів взаємодії з аудиторією [1].

Водночас стрімке впровадження ШІ викликає низку глибоких етичних питань, що стосуються достовірності інформації, прозорості джерел, маніпулятивного впливу, відповідальності за створений контент і потенційного зловживання технологіями. Особливо гостро ці питання постають у сфері масової комунікації, де ШІ може не лише передавати повідомлення, але й формувати суспільну думку, змінювати поведінкові моделі, підсилювати стереотипи або навпаки — приховано їх руйнувати [2].

У цьому контексті постає необхідність переосмислення етичних стандартів, які регулюють інформаційне середовище з використанням ШІ. Класичні засади журналістської та рекламної етики виявляються не завжди готовими до нових викликів, що постають перед суспільством у добу алгоритмічного впливу. Розмитість авторства, непрозорість технологій генерації контенту, а також недостатня інформованість користувачів — усе це створює нову віртуальну реальність, яка потребує критичного осмислення.

Аналіз останніх досліджень та публікацій. Відповідно до мети дослідження проведено аналіз літературних джерел, в яких розглянуто проблеми достовірності інформації, прозорості джерел походження інформації, маніпулятивного впливу та відповідальності за створення недостовірної інформації та зловживання технологіями.

Основні аспекти масової комунікації розглянуті у роботі [2]. Етичні аспекти використання штучного інтелекту досліджуються у роботі [1].

Принципи медіаграмотності, цифрової етики в медіа ресурсах, а також маніпуляцій в засобах масової інформації аналізуються у роботах [4] і [5].

Роль соціальних медіа та штучного інтелекту в інформаційних та кібервійнах досліджується у роботах [3] і [6].

Проте проблеми достовірності інформації, походження інформації, маніпуляцій, які використовуються за допомогою штучного інтелекту залишаються актуальними. За допомогою них відбуваються інформаційні атаки як на конкретних індивідів, так і на цілі прошарки суспільства, метою яких є різного роду маніпуляції людьми, розхитування соціальних та політичних настроїв суспільства та інші злочинні дії. Тому важливо вміти розрізняти та протидіяти такого роду атакам.

Мета дослідження: аналіз ключових етичних викликів, пов'язаних із використанням штучного інтелекту в масовій комунікації, а також окреслення можливих шляхів їх подолання через етичні, правові та освітні підходи.

Виклад основного матеріалу дослідження. Штучний інтелект — це сукупність цифрових технологій, здатних імітувати розумові процеси людини: аналіз, прийняття рішень, навчання, прогнозування, творчість. У контексті масової комунікації під ШІ найчастіше розуміють автоматизовані системи, які:

- генерують або модифікують контент (тексти, зображення, відео, голос);
- персоналізують інформацію відповідно до поведінки користувача;
- наповнюють вміст платформ;
- прогнозують реакцію аудиторії;
- приймають редакційні чи стратегічні рішення.

Ці функції виконуються завдяки алгоритмам машинного навчання, нейронним мережам, генеративним моделям (наприклад, GPT, DALL·E, Midjourney) та іншим технологіям, що постійно вдосконалюються [3].

Етика масової комунікації, як наукова дисципліна, оперує поняттями правдивості, чесності, об'єктивності, поваги до приватності та відповідальності за публічне слово. Ці принципи, сформовані ще в добу традиційної журналістики, слугували нормативною основою діяльності медіа. У XXI столітті, особливо з появою

цифрових платформ, відбувається трансформація самого поняття комунікації: вона стає глибшою та фрагментарнішою.

Застосування ШІ викликає перелік нових етичних питань, на які необхідно надати відповіді, а це зокрема:

1. Чи можна вважати контент, створений ШІ, «справжнім»?
2. Хто є автором або відповідальним за помилкові або шкідливі повідомлення створені штучним інтелектом?
3. Чи має аудиторія знати, що контент створено не людиною?
4. Які ризики виникають внаслідок прихованого алгоритмічного впливу на суспільну думку?

Цифрова етика як напрям філософських технологій фокусується на питаннях справедливості, прозорості, недискримінації та автономії у цифровому середовищі. Вона підкреслює, що розробка і впровадження ШІ повинні ґрунтуватися на етичних принципах, що враховують як права людини, так і соціальні наслідки рішень, прийнятих штучними алгоритмами [5].

Алгоритмічна відповідальність передбачає, що і розробники, і платформи, і користувачі повинні нести моральну (а часто й юридичну) відповідальність за результат дій ШІ. Це особливо важливо у масовій комунікації, де навіть невелика помилка або спотворення інформації можуть мати масштабні суспільні наслідки.

Одним із найбільш помітних наслідків використання ШІ в масовій комунікації є створення персоналізованого інформаційного простору, де контент добирається на основі вподобань, поведінкових моделей і попередніх дій користувача.

Персоналізований інформаційний простір — це цифрове середовище, в якому інформація, сервіси та взаємодія з користувачем адаптовані до його інтересів, поведінки, місцезнаходження, уподобань, потреб у конкретний момент часу. ШІ може обробляти величезні обсяги даних про користувача, а саме: історію переглядів, пошукові запити, місце знаходження, взаємодії з контентом. На основі цього формуються профілі користувачів з урахуванням поведінкових і контекстних моделей.

В табл. 1 наведені ключові алгоритми персоналізації користувачів за допомогою зібраної ШІ про них інформації.

Хоча така адаптація підвищує зручність і залученість користувача у пошуку необхідного йому контенту, вона водночас створює загрозу його інформаційної ізоляції — так званих «інформаційних бульбашок». У таких умовах користувачеві постійно підтверджують його погляди, що обмежує критичне мислення й сприяє радикалізації. Також до ризиків персоналізації ШІ користувачів можна віднести: надмірне збирання персональних даних (порушення конфіденційності); можливість маніпуляцій (наприклад, у політичній рекламі); відсутність прозорості щодо роботи алгоритмів.

Особливу небезпеку становить можливість прихованого маніпулювання поведінкою користувачів, через політичну рекламу або комерційні меседжі, які спрямовані не на раціональну оцінку, а на емоційний вплив.

Таблиця 1

Ключові алгоритми персоналізації

Алгоритм	Принцип роботи	Переваги	Застосування
Колаборативна фільтрація	Рекомендації на основі поведінки схожих користувачів	Навчається без знання контенту, динамічний адаптивний підхід	Рекомендації фільмів (Netflix), товарів (Amazon)
Контентна фільтрація	Враховує властивості контенту, який сподобався користувачу	Може працювати з новими користувачами, не потребує багатьох даних	Рекомендації новин, книг
Гібридна модель	Поєднує колаборативну та контентну фільтрацію	Покращена точність, компенсує недоліки окремих моделей	Spotify, YouTube, TikTok
Rule-based системи	Визначення правил вручну (наприклад, «якщо-умова-то-результат»)	Простота реалізації, контрольована логіка	Онлайн-опитування, банківські послуги
Фільтрація за демографією	Враховує стать, вік, місце проживання, мову тощо	Добре працює при обмежених даних про поведінку користувача	Реклама, регіональні новини
Context-aware персоналізація	Використовує контекст (час, локацію, пристрій, настрій) для рекомендацій	Підвищена точність у реальному часі	Мобільні додатки, навігатори, smart-гаджети
NLP-моделі (GPT, BERT тощо)	Обробка природної мови для аналізу інтересів і генерації контенту	Глибоке розуміння тексту, можливість діалогу	Чат-боти, email-маркетинг, рекомендації статей
Моделі глибокого навчання (Deep Learning)	Нейронні мережі аналізують великі масиви даних і зв'язки між подіями	Висока точність при великих даних	YouTube, Google Discover, Instagram Explore

III сьогодні здатен створювати високореалістичні тексти, зображення, відео та аудіо, які важко відрізнити від оригінальних. Ця технологічна можливість стає особливо небезпечною у випадках створення дезінформації, маніпулятивного або провокаційного контенту.

Застосування deepfake-технологій у політичній боротьбі, дискредитаційних кампаніях або навіть у гумористичних форматах часто виходить за межі етичного, порушуючи право на репутацію, приватність і довіру до інформаційних джерел. В умовах кризи фактів і недовіри до ЗМІ, такі технології лише поглиблюють розрив між медіа і суспільством.

Ще один виклик — недостатня прозорість джерел походження інформації. У більшості випадків користувач не має змоги відрізнити, чи створено контент людиною, чи алгоритмом ШІ. Відсутність маркування матеріалів, згенерованих ШІ, ставить під сумнів достовірність інформації та ускладнює її критичне сприйняття.

Це особливо важливо в контексті журналістики, соціальних медіа та реклами, де етична комунікація базується на принципі поінформованого вибору.

У класичних моделях масової комунікації існує чітке уявлення про автора, редактора, видавця — тих, хто несе відповідальність за зміст інформації. У випадку з ШІ ці межі стираються. Виникає питання: хто відповідає за створений контент — розробник алгоритму, замовник, інформаційна платформа, або сам штучний інтелект?

Сьогодні юридичне поле не встигає адаптуватися до нових викликів, а етична відповідальність стає розмитою. Це створює ризики уникнення відповідальності за наслідки дезінформації — як непередбачувані, так і свідомо спрямовані на дестабілізацію.

Одним із найрезонансніших випадків використання штучного інтелекту у політичному контексті стала поява відео-маніпуляцій із використанням технології **deepfake** — алгоритмічно згенерованих відеозаписів, де обличчя і голоси реальних політиків використовуються для створення фальшивих повідомлень. Наприклад, у період виборів до Європейського парламенту 2024 року в інтернеті ширилися відео, на яких нібито німецькі та французькі політики висловлюють проросійські погляди. Лише завдяки ретельному журналістському розслідуванню вдалося виявити, що ці відеоролики були змонтовані із застосуванням генеративних нейромереж [4].

Такі маніпуляції мають декілька ключових ризиків:

- **дестабілізація довіри до реальних джерел інформації;**
- провокація **емоційних реакцій** в аудиторії, особливо напередодні голосування;
- використання дезінформації як **інструменту в умовах інформаційної війни.**

Ще однією формою впливу ШІ є використання **алгоритмічного аналізу поведінки виборців** з метою мікро-таргетування політичної реклами. Відомий кейс із кампанією **Cambridge Analytica** у США (2016) та Великобританії (Brexit) продемонстрував, як за допомогою алгоритмів машинного навчання можна виявляти психотипи виборців і створювати для них індивідуалізовані політичні меседжі. Хоча формально у тому кейсі використовувалися не генеративні моделі, його логіка була збережена в новіших кампаніях — уже з використанням генеративного ШІ для створення текстів, банерів, гасел і навіть синтетичних відеозвернень.

Ці інструменти:

- підвищують ефективність політичних кампаній;
- **порушують принципи прозорості та чесності** у виборчому процесі;
- можуть використовуватися для **підсилення поляризації суспільства** або впровадження прихованих маніпулятивних наративів.

З етичної точки зору обидва випадки демонструють проблему **відсутності належного регулювання** діяльності ШІ у сфері політичної комунікації. Відповідальність за створення, розповсюдження і споживання подібного контенту розподілена

між багатьма сторонами — розробниками ШІ, інформаційними платформами, партіями, технологічними компаніями та самими виборцями.

Постає запитання: **чи може демократичний процес залишатися відкритим і прозорим**, якщо ключові його елементи опосередковані непрозорими алгоритмами? Як гарантувати рівність доступу до правдивої інформації в умовах інформаційної гібридної війни?

Висновки. Використання штучного інтелекту у сфері масової комунікації відкриває значний потенціал для розвитку персоналізованого, інтерактивного та ефективного контенту. Водночас ці технології породжують низку складних етичних викликів, що потребують глибокого суспільного осмислення та нормативного врегулювання. Як зазначено в дослідженні, головними серед них є: ризик маніпуляцій, використання дезінформації, розмитість відповідальності, зниження прозорості комунікації, а також порушення базових принципів демократії.

Розгляд кейсів політичного використання ШІ — зокрема deepfake-технологій та алгоритмічного мікро-таргетування — доводить, що ці виклики мають не лише етичний, але й безпосередній вплив на якість громадянського вибору, довіру до інституцій та стабільність суспільних процесів, що може призвести до значних соціальних збурень та дестабілізацій у певних груп людей, які піддаються впливу сучасним комунікаційним технологіям. В цьому випадку в умовах кібервійни ШІ може стати інструментом для різноманітних цілеспрямованих атак для деструктивного впливу на свідомість користувачів мережі Інтернет певних регіонів, де проходять військові конфлікти і тих, які формують соціальну думку про ці конфлікти [6].

У цьому контексті важливо не лише формувати етичні стандарти використання ШІ, але й створювати механізми їх контролю та впровадження — як на рівні медіа та розробників технологій, так і в межах державної політики, освіти, правового регулювання.

Перспективи подальших досліджень включають:

- розробку етичних протоколів взаємодії між журналістами, платформами та розробниками ШІ;
- порівняльний аналіз національних та міжнародних підходів до регулювання ШІ у сфері медіа;
- дослідження впливу генеративного ШІ на формування громадської думки в умовах гібридної інформаційної війни;
- емпіричні дослідження рівня критичного мислення аудиторії щодо ШІ-контенту;
- створення правових норм, які б регулювали відповідальність за використання ШІ з метою дестабілізації політичної системи, порушень демократичних норм та принципів масової комунікації, поширень дезінформації тощо.

Загалом, ефективна етична взаємодія з інструментами штучного інтелекту — це не лише технологічне чи правове питання, але й ключовий виклик для збереження демократичного інформаційного простору у XXI столітті.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Гринько А. В. Етичні аспекти використання штучного інтелекту в цифровій журналістиці. Наукові записки Інституту журналістики. 2023. № 1(44). С. 25–32.
2. Павлюк М. І. Масова комунікація: соціокультурний вимір. Львів: ПАІС, 2021. 312 с.
3. Щербина О. В. Інформаційні війни в умовах гібридної загрози: роль соціальних медіа та III. Соціальні комунікації: наукові пошуки. 2022. № 2(18). С. 41–50.
4. Козлова М. Ю. Маніпуляція у ЗМІ та етика нових медіа. Медіакритика. 2020. № 3. С. 12–19.
5. Воронова Л. І. Медіаграмотність і цифрова етика у XXI столітті. Вісник Київського національного університету культури і мистецтв. Серія: Соціальні комунікації. 2021. № 4. С. 74–82.
6. Воеводін І. С. Кібервійна як сучасний метод ведення збройних конфліктів. Протидія кіберзагрозам та торгівлі людьми. Харків, 2019. С. 46–48.

REFERENCES

1. Hrynko, A. V. (2023). Ethical aspects of artificial intelligence use in digital journalism. Scientific Notes of the Institute of Journalism, (1)44, 25–32.
2. Pavliuk, M. I. (2021). Mass communication: A socio-cultural dimension. Lviv: PAIS. [In Ukrainian].
3. Shcherbyna, O. V. (2022). Information wars under hybrid threats: The role of social media and AI. Social Communications: Scientific Researches, (2)18, 41–50.
4. Kozlova, M. Yu. (2020). Manipulation in the media and ethics of new media. MediaKrytyka, (3), 12–19.
5. Voronova, L. I. (2021). Media literacy and digital ethics in the 21st century. Herald of Kyiv National University of Culture and Arts. Series: Social Communications, (4), 74–82.
6. Voievodin, I. S. (2019). Cyberwar as a modern method of armed conflict. In Countering Cyber Threats and Human Trafficking (pp. 46–48). Kharkiv. [In Ukrainian].

doi: 10.32403/1998-6912-2025-1-70-108-115

**ETHICS OF USING ARTIFICIAL INTELLIGENCE
IN MASS COMMUNICATION**

V. I. Sabat¹, V. V. Matsiuk²

¹National University “Lviv Polytechnic”, Lviv, Ukraine
12 Stepana Bandery Street, Lviv, 79000, Ukraine
volodymyr.i.sabat@lpnu.ua,

²Lviv Polygraphic College of the UAD,
Mytna Square 1, Lviv 79000 Ukraine
altentop17@yahoo.com

The article explores the ethical issues surrounding the use of artificial intelligence (AI) in the field of mass communication. AI offers significant potential for the development of personalized, interactive, and efficient content. At the same time, these technologies generate a range of complex ethical challenges that require profound societal reflection and regulatory oversight. As outlined in the paper, key concerns include the risk of manipulation, the use of disinformation, blurred responsibility, reduced communication transparency, and the erosion of core democratic principles.

Particular attention is given to the use of AI in creating a personalized informational environment for users. AI enables a digital space in which content, services, and user interactions are adapted to individual interests, behavior, location, preferences, and needs at any given time. To achieve this, AI systems process vast amounts of user data, including browsing history, search queries, location data, and content interactions. Based on this data, user profiles are built using behavioral and contextual models. While such adaptation enhances user convenience and engagement with relevant content, it also poses a significant risk of informational isolation — the so-called “filter bubbles.” In such conditions, users are constantly exposed to information that reinforces their existing views, thereby limiting critical thinking and contributing to radicalization. Other risks associated with AI-driven personalization include excessive collection of personal data (infringing on privacy), potential for manipulation (e.g., in political advertising), and a lack of transparency in how algorithms function.

A particularly concerning threat is the covert manipulation of user behavior through political or commercial messages aimed not at rational judgment, but at emotional influence.

Effective ethical interaction with AI tools is not only a technological or legal matter but a critical challenge for safeguarding the democratic information space in the 21st century. Therefore, the ethics of AI use in mass communication requires new approaches, thoughtful reconsideration, and accountability on the part of those who integrate AI into their professional practices.

Keywords: *artificial intelligence, information manipulation, mass communication, disinformation, media literacy.*

Стаття надійшла до редакції 07.05.2025.

Received 07.05.2025.