

5. Mollick E., Mollick L. Assigning AI: Seven Approaches for Students, with Prompts. Social Science Research Network. 2023. 40 p. DOI: 10.2139/ssrn.4475995.
6. Морзе Н. В., Буйницька О. П. Підвищення якості освіти — нагальна проблема сучасної освіти. *Інформаційні технології і засоби навчання*. 2013. № 4 (36). С. 1–14. URL: <https://journal.iitta.gov.ua/index.php/itlt/article/view/890> (дата звернення: 15.03.2025).
7. Lewis P., Perez E., Piktus A. et al. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. *Advances in Neural Information Processing Systems*. 2020. Vol. 33. P. 9459–9474. URL: <https://proceedings.neurips.cc/paper/2020/hash/6b493230205f780e1bc26945df7481e5-Abstract.html> (дата звернення: 20.03.2025).
8. Zawacki-Richter O., Marín V. I., Bond M., Gouverneur F. Systematic review of research on artificial intelligence applications in higher education – where are the educators? *International Journal of Educational Technology in Higher Education*. 2019. Vol. 16, № 39. P. 1–27. DOI: 10.1186/s41239-019-0171-0.

Віталій ПЕТРИК, аспірант кафедри інформаційних технологій ННІ права та психології, Національна академія внутрішніх справ, адвокат

КРИМІНАЛЬНА ВІДПОВІДАЛЬНІСТЬ ЗА НЕЗАКОННЕ ПОШИРЕННЯ ДИПФЕЙКІВ

Розвиток технологій штучного інтелекту зумовив масове поширення дипфейків – штучно згенерованих аудіо-, відео- та фотоматеріалів, створених за допомогою алгоритмів глибинного навчання. Автори дипфейків можуть використовувати персональні дані особи (зовнішність, голос, манеру мовлення) без її згоди.

Відповідно до Закону України «Про захист персональних даних» та Загального регламенту про захист даних (GDPR), зображення обличчя і голос належать до чутливих біометричних даних [1; 2], а тому неправомірне їх використання безпосередньо посягає на честь, гідність, ділову репутацію та приватне життя особи.

За даними Sensity AI, кількість дипфейків у відкритому доступі впродовж 2023–2025 років зросла більш ніж утричі [3].

Відсутність спеціального правового регулювання цього явища в Україні формує прогалину в системі захисту прав людини та інформаційної безпеки, особливо в умовах збройної агресії та інформаційно-психологічних операцій проти держави.

Використання персональних даних особи з метою створення дипфейків, яке спричинило негативні наслідки для особи, дані якої було використано, є непоодинокими випадками.

Яскравим прикладом є кейс директора школи у США, який мав місце у 2024 році. Голос директора школи Пайксвілла (Меріленд) Еріка Айсверта було використано для поширення аудіозапису з расистськими висловлюваннями, внаслідок чого директора відсторонили від роботи, а його родина отримувала погрози. Під час розслідування, шляхом проведення експертиз, зокрема фоноскопичної, було встановлено, що запис створив директор з фізичного виховання школи з мотивів помсти [4; 5].

Такий випадок вказує на високу небезпечність дипфейків:

- легкість створення дипфейку: особа без спеціальних знань може створити аудіоконтент, який сторонні особи сприймають як висловлювання конкретної особи;
- миттєвість завдання шкоди: поширення в мережі Інтернет може відбутись миттєво та охопити невизначену кількість осіб, унаслідок чого шкода потерпілій особі завдається миттєво;

– необхідність часу та витрат на спростування інформації: встановлення неоригінальності та спростування поширеного дипфейку вимагають часу і залучення спеціалістів.

Схожі випадки траплялися в межах українського інформаційного простору. У листопаді 2025 року в TikTok з'явилося дипфейк-відео з головою варшавського фонду «Український дім» Мирославою Керик, у якому від її імені поширювались антипольські заяви. Відео-дипфейк зібрав понад 160 тис. переглядів за кілька годин [6].

У 2022 році було поширено дипфейк-відео Президента України В. Зеленського із закликом до складання зброї [7; 8] — один із перших масштабних прикладів застосування ШІ в інформаційній війні.

Зарубіжні країни запроваджують відповідальність за поширення дипфейків, створених із використанням даних інших осіб без їхньої згоди. Зокрема, до кримінальних кодексів Франції та Італії щодо криміналізації створення та поширення дипфейків з використанням зображення або голосу особи без її згоди було внесено зміни у 2024 та 2025 роках відповідно.

Так, стаття 226-8 Кримінального кодексу Франції встановлює відповідальність за поширення будь-якими засобами дипфейку з використанням слів або зображення особи без її згоди у вигляді позбавлення волі на строк один рік та штрафу у розмірі 15 000 євро [9].

Стаття 612-quater Кримінального кодексу Італії встановлює відповідальність за такі ж діяння у вигляді позбавлення волі від 1 до 5 років [10].

Водночас у відкритих джерелах відсутні відомості щодо судових вироків за цими статтями, що може бути зумовлено, зокрема, складністю розслідування таких злочинів.

Законодавство України не містить аналогічного складу злочину, хоча спроби криміналізації незаконного використання персональних даних під час створення та поширення дипфейків уже мали місце.

Так, 16 січня 2026 року була зареєстрована петиція до Кабінету Міністрів України з вимогою запровадити відповідальність за створення та поширення deepfake-контенту. Однак вона не набрала необхідної підтримки громадян і станом на 17 квітня 2026 року отримала лише 50 голосів із необхідних 25 тисяч [11].

Висновки

Алгоритм протиправних дій щодо незаконного використання зображення або голосу особи для створення та поширення дипфейків, без її згоди та що завдало шкоди, є спільним для всіх таких випадків: створення дипфейку зловмисниками за допомогою алгоритмів ШІ; поширення інформації серед інших осіб; сприйняття іншими особами дипфейку як оригіналу.

Вказані діяння є реальною та зростаючою загрозою правам людини та інформаційній безпеці. Вони можуть завдавати шкоди не лише потерпілим, дані яких було використано, а й близьким цих осіб і третім особам, які були пов'язані з поширенням або сприйняттям такої інформації.

З урахуванням наведеного, вважаємо за необхідне включити до Кримінального кодексу України самостійний склад злочину, який передбачатиме відповідальність за незаконне створення та поширення дипфейків з використанням персональних даних без згоди особи, що завдало шкоди особі, суспільству або державі.

Криміналізація дипфейків є не лише правовою необхідністю, а й елементом системи національної безпеки України.

Список використаних джерел

1. Про захист персональних даних : Закон України від 01.06.2010 № 2297-VI. URL: <https://zakon.rada.gov.ua/laws/show/2297-17> (дата звернення: 05.05.2026).
2. Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 (General Data Protection Regulation). URL: <https://eur-lex.europa.eu/eli/reg/2016/679/oj> (дата звернення: 05.05.2026).

3. Sensity AI. The State of Deepfakes: Landscape, Threats and Impact. 2024. URL: <https://sensity.ai/reports> (дата звернення: 05.05.2026).
4. Athletic Director Charged in Pikesville High School AI Case : Baltimore County Police Department. 25.04.2024. URL: <https://www.baltimorecountymd.gov/departments/police/news/2024/04/25/athletic-director-charged-in-pikesville-high-school-ai-case> (дата звернення: 05.05.2026).
5. Pikesville High School's principal was accused of offensive language on a recording. Authorities now say it was a deepfake. CNN. 26.04.2024. URL: <https://www.cnn.com/2024/04/26/us/pikesville-principal-maryland-deepfake-cec/index.html> (дата звернення: 05.05.2026).
6. Кузьменко О. Deepfake-відео з Мирославою Керик. Наш вибір. 13.11.2025. URL: <https://naszwybir.pl/keryk-deepfake/> (дата звернення: 05.05.2026).
7. Deepfake video of Zelenskiy telling Ukrainians to surrender appears on hacked website. Reuters. 16.03.2022. URL: <https://www.reuters.com/world/europe/deepfake-video-zelenskiy-telling-ukrainians-surrender-appears-hacked-website-2022-03-16/> (дата звернення: 05.05.2026).
8. Deepfake of Zelensky could be 'tip of the iceberg'. BBC News. URL: <https://www.bbc.com/news/technology-60780142> (дата звернення: 05.05.2026).
9. Code pénal. Article 226-8. Légifrance. URL: https://www.legifrance.gouv.fr/codes/article_lc/LEGIARTI000049571542 (дата звернення: 05.05.2026).
10. Codice penale. Art. 612-quater. Brocardi.it. URL: <https://www.brocardi.it/codice-penale/libro-secondo/titolo-xii/capo-iii/sezione-iii/art612quater.html> (дата звернення: 05.05.2026).
11. Українці не підтримали ідею відповідальності за deepfake: петиція не набрала голосів. Судово-юридична газета. URL: <https://sud.ua/uk/news/ukraine/358710-ukraintsy-ne-podderzhali-ideyu-otvetstvennosti-za-deepfake-petitsiya-ne-nabrала-golosov> (дата звернення: 05.05.2026).

Володимир ШАПОШНИКОВ, студент 2-го курсу магістратури, спеціальність «Штучний інтелект», н. гр. КІ-41мн, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського»

МОДЕЛІ МАШИННОГО НАВЧАННЯ ДЛЯ ПРОГНОЗУВАННЯ КІБЕРАТАК НА ОСНОВІ АНАЛІЗУ МЕРЕЖЕВОГО ТРАФІКУ

Актуальні напрями використання методів машинного навчання для прогнозування кіберзагроз. Сучасні інформаційно-комунікаційні системи дедалі частіше стають об'єктами складних багаторівневих кібератак, наслідками яких є фінансові втрати, компрометація конфіденційної інформації, порушення функціонування державних органів, підприємств та об'єктів критичної інфраструктури. У таких умовах традиційні засоби захисту інформації, що базуються переважно на сигнатурному аналізі та заздалегідь визначених правилах виявлення загроз, не завжди здатні ефективно протидіяти новим і модифікованим атакам. Це обумовлює необхідність переходу від реактивної до проактивної моделі кіберзахисту, заснованої на прогнозуванні можливих кіберінцидентів та своєчасному реагуванні на потенційні загрози [1; 2].

Одним із найбільш перспективних інструментів реалізації такого підходу є методи машинного навчання та інтелектуального аналізу даних. Завдяки здатності працювати з