

УДК 004.8:32:316.77:351.86

[https://doi.org/10.52058/3041-1793-2025-11\(16\)-1540-1551](https://doi.org/10.52058/3041-1793-2025-11(16)-1540-1551)

Коваль Ольга Миколаївна кандидат юридичних наук, доцент, доцент кафедри приватного та публічного права, Київський національний університет технологій та дизайну, м. Київ, <https://orcid.org/0000-0003-1509-7258>

Шаповалова Алла Миколаївна кандидат політичних наук, доцент, доцент кафедри приватного та публічного права, Київський національний університет технологій та дизайну, м. Київ, <https://orcid.org/0000-0002-4920-5539>

Фастовець Наталія Валеріївна кандидат наук з державного, доцент, доцент кафедри приватного та публічного права, Київський національний університет технологій та дизайну, м. Київ, <https://orcid.org/0000-0001-8619-8975>

Букорос Тетяна Олександрівна кандидат юридичних наук, доцент, доцент кафедри приватного та публічного права, Київський національний університет технологій та дизайну, м. Київ, <https://orcid.org/0000-0002-4059-2632>

СИНТЕТИЧНА РЕАЛЬНІСТЬ: ЯК DEERFAKE-ТЕХНОЛОГІЇ ТРАНСФОРМУЮТЬ ПОЛІТИЧНУ КОМУНІКАЦІЮ ТА АРХІТЕКТУРУ НАЦІОНАЛЬНОЇ БЕЗПЕКИ

Анотація. Сучасний інформаційний простір стоїть на порозі фундаментальної трансформації: технології deepfake (поєднання «deep learning» і «fake») сформувалися в окрему, стратегічно небезпечну парадигму дезінформації. Ці синтетичні медіа, створені або модифіковані алгоритмами ШІ, не просто фальсифікують окремі факти, а фундаментально підривають довіру до аудіовізуальних «доказів» як до об'єктивного відображення реальності. На відміну від класичних текстових фейків, deepfake атакують саму епістемологічну основу публічної сфери. Вони створюють і посилюють «дивіденд брехуна» (liar's dividend) — руйнівну можливість для зловмисників миттєво оголосити неправдою будь-який незручний, але автентичний аудіо- чи відеозапис.

Стрімка еволюція технології є показовою. Лише за кілька років вона пройшла шлях від перших гучних випадків несанкціонованих порно-відео зі знаменитостями (як-от Галь Гадот у 2017 р.) до складних шахрайських схем. Це включає фінансові махінації через клонування голосу (Vishing), маніпуляції громадською думкою (т.зв. «October surprise» на виборах) та прямі загрози державній безпеці. Водночас deepfake ніколи не працюють ізольовано. Вони є «корисним навантаженням» (payload) у складних, багатоканальних



операціях впливу. Ці операції комбінують бот-мережі для початкового «посіву», псевдо-медіа для «легалізації» контенту та державні дезінформаційні мережі (на кшталт «Doppelganger») для кроссплатформного посилення.

Аналіз deepfake поза цим гібридним, мультимодальним контекстом є стратегічною помилкою. Мета дослідження — показати, що ключова небезпека технології полягає не стільки в окремих актах обману, скільки в довгостроковому, корозійному знищенні інституційної довіри до політичної інформації, медіа, судочинства та легітимності управління в цілому.

Дослідження поєднує технічну деконструкцію (аналіз моделей генерації, включно з GAN та дифузійними моделями, демократизацію інструментів через зниження вартості обчислень, та асиметрію «часу поширення» проти «часу перевірки») з порівняльно-правовим аналізом і кейс-стаді, що охоплюють підходи ключових юрисдикцій (ЄС, США). Основний висновок полягає в тому, що наявні режими регулювання страждають від «проблеми темпу» (racing problem), залишаючись фрагментованими та переважно реактивними. Самі по собі пост-фактум заборони і криміналізація «шкідливого використання» дають обмежений ефект. Вони не здатні зупинити поширення шкоди в режимі реального часу без створення паралельної технічної інфраструктури довіри (наприклад, стандартів простежуваності контенту).

Швидкі процедурні засоби захисту: Створення ефективних механізмів для жертв (оскільки шкода від deepfake є миттєвою та віральною), включно з прискореними судовими приписами на видалення/спростування та «гарячими лініями» під час виборчих періодів. Прозорість і підзвітність платформ: Покладення на онлайн-платформи (як ключових векторів поширення) чітких обов'язків належної обачності (due diligence), аудиту моделей детекції та прозорого журналювання модераторських рішень. Медіаграмотність та інституційна стійкість: Посилення не лише загальної медіаграмотності, але й професійних стандартів верифікації для ЗМІ, правоохоронних органів та публічних інституцій, формуючи «людський брендмауер» як останню лінію захисту.

Запропонована модель розроблена з урахуванням адаптації для правопорядку України. Вона враховує не лише необхідність зближення з європейськими підходами (зокрема, EU AI Act та Digital Services Act (DSA)), але й нагальні потреби національної безпеки в умовах повномасштабної гібридної та конвенційної війни, де дезінформація є прямим інструментом агресії.

Ключові слова: масова комунікація; мас-медіа та політика; методи та технології в політиці; професійна етика й протоколи в політиці, сучасні політичні ідеології.

Koval Olga PhD, Associate Professor, Associate Professor Department of Private and Public Law, Kyiv National University of Technology and Design, Kyiv, <https://orcid.org/0000-0003-1509-7258>

Fastovets Natalia PhD in Public Administration, Associate Professor, Associate Professor of the Department of Private and Public Law, Kyiv National University of Technologies and Design, Kyiv, <https://orcid.org/0000-0001-8619-8975>

Bukoros Tetiana PhD, Associate Professor, Associate Professor of the Department of Private and Public Law, Kyiv National University of Technologies and Design, Kyiv, <https://orcid.org/0000-0002-4059>

Shapovalova Alla Mykolaivna PhD, Associate Professor, Associate Professor, Associate Professor at the Department of Private and Public Law, Kyiv National University of Technologies and Design, Kyiv, <https://orcid.org/0000-0002-4920-5539>

SYNTHETIC REALITY: HOW DEEPPFAKE TECHNOLOGIES TRANSFORM POLITICAL COMMUNICATION AND THE ARCHITECTURE OF NATIONAL SECURITY

Abstract. The contemporary information space stands on the verge of a qualitative transformation: deepfake technologies (a blend of “deep learning” and “fake”) have become a distinct paradigm of disinformation. They are synthetic media (images, audio, video) created or modified by AI algorithms that undermine the long-standing trust in audiovisual “evidence” as an objective reflection of reality. Unlike classical “fake news,” which is predominantly textual, deepfakes strike at the epistemic foundations of the public sphere, fueling the “liar’s dividend”—the ability to dismiss any inconvenient recording as fabricated.

The technology’s rapid evolution is illustrative: from the first high-profile cases of non-consensual pornographic videos involving celebrities (including actress Gal Gadot in 2017) to uses in financial fraud (voice/face cloning to misappropriate funds), market manipulation, electoral interference, and threats to state security. At the same time, deepfakes never operate in isolation: they are the “payload” in multi-channel influence operations that combine bot networks, advertising accounts, pseudo-media, and cross-platform infrastructures (including state actors and networks such as “Doppelganger”). Analyzing them outside this hybrid context is a strategic error.

The aim of this study is to show that the key danger of deepfakes lies not only in discrete acts of deception, but in the long-term erosion of institutional trust in political information, evidentiary contestation, and the legitimacy of governance. We combine technical deconstruction (generation models, compute costs, diffusion speed, and the “time-to-verify” versus “time-to-spread” gap) with comparative-law



analysis and case studies focused on four blocks of risk: (1) privacy, reputation, and protection against non-consensual sexualization; (2) image/voice rights and copyright; (3) elections, information security, and public order; (4) evidence law, procedural safeguards, and standards of due verification.

Main findings: current regulatory regimes are fragmented and largely reactive; prohibitions and the criminalization of “harmful uses” have limited effect without a supporting technical trust infrastructure. A multi-layer model of responsibility and prevention is required: (a) prevention and traceability (content provenance, cryptographic signing, watermarking, default labeling of synthetic media); (b) rapid procedural remedies for victims (expedited takedown/rectification orders, election-period hotlines, cross-platform coordination); (c) platform transparency and accountability (duties of due diligence, audits of models and detection systems, logging of moderation decisions); (d) media literacy and professional verification standards for newsrooms and public institutions.

Originality and practical value lie in integrating technical and legal approaches into a single risk-governance framework oriented toward rapid response, proportionality, and cross-jurisdictional interoperability. The proposed model is adapted to Ukraine’s legal order, taking into account convergence with European approaches, national-security needs, and platform-centric regulatory practice.

Keywords: mass media and politics; foundations of professional ethics and protocols in politics; sociology of public communications and public opinion; global and intercultural communication; contemporary political ideologies.

Постановка проблеми. Початкове, здавалося б, “низькорівневе” використання deepfake стало де-факто глобальним, децентралізованим полігоном для досліджень та розробок [1]. Цей масовий попит стимулював вибухове вдосконалення відкритих (open-source) інструментів [2], які стали доступними на платформах на кшталт GitHub.

Аналітичні звіти, зокрема від Міністерства внутрішньої безпеки США (DHS), пропонують корисну рамку для категоризації цих загроз, розрізняючи два основні типи атак:

“Низький вплив, висока ймовірність” (Low-impact, high-probability): Це вже реалізована та повсякденна загроза. Вона включає фінансове шахрайство, булінг, дифамацію та масове створення порнографії без згоди (non-consensual pornography). Дослідження показують, що 99% жертв такого контенту — жінки, що також робить deepfake гендерною зброєю [3].

“Високий вплив, низька ймовірність” (High-impact, low-probability): Це традиційна оцінка загрози для національної безпеки. Сюди входять сценарії, де державний актор може використати deepfake для провокування війни, віддання фальшивого військового наказу або дестабілізації фінансових ринків.

Мета статті показати, що ключова небезпека deepfake полягає не лише в окремих шахрайських діях, а у довгостроковому зниженні інституційної

довіри до політичної інформації, процедур змагальності доказів і легітимності управління.

Виклад основного матеріалу. В основі більшості deepfake-технологій лежать дві основні архітектури машинного навчання:

Автоенкодер (Autoencoders): Ця архітектура є фундаментом для класичних "face swap" (замін обличчя) [4]. Процес працює наступним чином:

Енкодер (Encoder): Це нейронна мережа, яка навчається на тисячах зображень двох осіб — Джерела (А) та Цілі (В). Вона стискає ці зображення до "латентного простору", по суті, вивчаючи їхні фундаментальні, узагальнені риси обличчя [5].

Декодер (Decoder): Це друга мережа, яка навчається реконструювати оригінальні зображення з цього латентного простору.

Процес підробки: Для створення deepfake, енкодер обробляє зображення Джерела (А), вивчаючи його міміку та положення голови. Потім ці стиснені дані передаються декодеру, який був навчений на обличчі Цілі (В). В результаті, декодер реконструює обличчя В, але з мімікою та виразом обличчя А.

Останні досягнення, такі як дифузійні моделі та мультимодальні великі мовні моделі (LLMs), ще більше спрощують цей процес, дозволяючи створювати високоякісні синтетичні медіа з простих текстових запитів [6].

Найбільш тривожним аспектом є не стільки існування цих технологій, скільки їхня радикальна "демократизація".

Колапс вимог до навичок: Як зазначає у своєму звіті Агентство національної безпеки (АНБ) США, для використання цих інструментів більше не потрібен значний технічний досвід. Часто достатньо мати "персональний ноутбук та мінімальний набір технічних навичок" [7].

Всюдисутність даних: Як зазначив експерт з МІТ, кожне фото, відео чи аудіозапис, опубліковані в Інтернеті, є "потенційними даними для тренування" зловмисників.

Ця демократизація призвела до вибухового зростання загрози, що чітко видно з економічних показників та статистики атак у 2024-2025 роках. Прогнозується щорічне зростання обсягу deepfake-контенту на 900%. За оцінками, загальна кількість deepfake-файлів зросте з 500 000 у 2023 році до 8 мільйонів до кінця 2025 року.

Клонування голосу найбільш демократизований вектор атаки. Сучасні інструменти вимагають лише 3-5 секунд вихідного аудіо для створення переконливого клону голосу зі збігом 85% [8].

Сумнозвісний рободзвінок, що імітував голос президента Джо Байдена у 2024 році, коштував приблизно \$1 і був створений менш ніж за 20 хвилин. А втрати від deepfake-шахрайства лише за перший квартал 2025 року перевищили \$200 мільйонів.¹⁸

Стратегічно найважливішою зміною, яку ілюструють ці цифри, є перехід від data-heavy (жадібних до даних) моделей до data-light (легких) атак. Раніше,



щоб створити "face swap" за допомогою автоенкодера, зловмиснику потрібні були тисячі високоякісних зображень цілі 15, що обмежувало атаки до дуже публічних осіб (політиків, акторів).

Сьогоднішні моделі, що вимагають 3-5 секунд аудіо, означають, що будь-яка людина, яка має профіль у соціальних мережах, запис голосу на автовідповідачі або хто хоча б раз виступав на публічній конференції, стає вразливою ціллю. Це кардинально змінює розрахунок атаки: від масової дезінформації до гранулярних, високоточних, цільових атак. Наприклад, підробка аудіо-наказу військового командира для конкретного підрозділу, цільове шахрайство з видачею себе за СЕО компанії, або викрадення людей з використанням клонованих голосів родичів.

У 2024 рік, коли вибори відбулися у понад 50 країнах, що представляють більше половини населення світу [9], став першим глобальним полігоном для масованого застосування deepfake-технологій у політичній боротьбі. Аналіз цих кампаній демонструє не лише технічну можливість, але й тактичну гнучкість цієї зброї.

Під час загальних виборів в Індії deepfake не були маргінальною загрозою; вони стали центральним інструментом агітації, який відкрито використовувався основними політичними партіями. Тактики включали "воскресіння" померлих політичних лідерів для отримання їхньої "підтримки", створення та поширення підроблених відео популярних зірок Боллівуду, які нібито агітували за певних кандидатів, а також поширення цільової дезінформації про наміри опонентів змінити конституцію [10].

В Індонезії вірусного поширення набув deepfake померлого диктатора Сухарто, який закликав громадян голосувати за певну політичну партію. [11] Цей випадок вийшов за межі простої агітації, спровокувавши масштабну етичну та правову кризу щодо маніпулятивного використання історичної пам'яті та символічної легітимізації через цифрову некромантію.

Аналіз цих випадків показує, що вибір модальності (аудіо чи відео) є свідомим стратегічним рішенням. У Словаччині та США використовувалися аудіофейки. В Індії та Індонезії — відеофейки. Причина тактична: аудіо дешевше, його можна створити за лічені хвилини, і воно може масово поширюватися через традиційні телефонні мережі (рободзвінки), що робить його ідеальною зброєю для швидких атак "в останню хвилину" або придушення явки. Відео, у свою чергу, має більший емоційний вплив і використовується для створення "шокового" ефекту або надання символічної ваги, як у випадку з Сухарто.

Технології генеративного ШІ є за своєю суттю подвійного призначення. Навіть у політичному контексті існують приклади їхнього позитивного або нейтрального використання. Будь-яка регуляторна політика має враховувати цей факт, оскільки повна заборона технології неможлива.

У Венесуелі журналісти використовували ШІ-аватари для анонімного висвітлення новин, що критикують уряд, таким чином уникаючи урядових репресій та покращуючи інформаційне середовище [12].

Місцева новинна організація в США (Arizona Agenda) використала deepfake, щоб наочно пояснити глядачам, як легко маніпулювати відео, посилюючи медіаграмотність. У Каліфорнії кандидат, який втратив голос через ларингіт, прозоро використовував ШІ-клонування голосу для зачитування своїх повідомлень на зустрічах з виборцями.

Непряма загроза, відома як "Дивіденд Брехуна" (термін, введений професорами права Боббі Чесні та Даніель Сітрон), є, можливо, найбільш руйнівною та важковирішуваною [13].

Визначення Концепції: "Дивіденд Брехуна" — це ситуація, коли сама обізнаність громадськості про існування технології deepfake починає працювати на користь зловмисників. Коли з'являються справжні, компрометуючі аудіо- чи відеодокази корупції, злочину або аморальної поведінки політика, він отримує можливість правдоподібно заперечувати їх, просто заявивши: "Це deepfake!".

Це приносить пряму користь (дивіденд) брехуну. Експериментальні дослідження в академічних журналах підтверджують, що політики, які безпідставно заявляють про "фейкові новини" у відповідь на реальний скандал, зберігають значно вищий рівень підтримки серед своїх прихильників, ніж ті, хто вибачається або мовчить. Ця стратегія успішно створює "інформаційну невизначеність", що дозволяє уникнути відповідальності.

Наслідком цього феномену є тотальний суспільний цинізм та ерозія довіри. Якщо будь-який справжній доказ можна відкинути як "фейк", а будь-який фейк може бути сприйнятий як правда, громадяни втрачають здатність довіряти будь-якій політичній інформації та інституціям, що робить демократичну підзвітність неможливою. Мета ворога в гібридній війні — не просто змусити вас повірити в брехню; а змусити вас засумніватися в усій правді.

Феномен "Дивіденду Брехуна" та прямої дестабілізації зійшовся в одному з найбільш показових випадків — спробі державного перевороту в Габоні.

Контекст: У 2018 році президент Габону Алі Бонго зник з публічного простору на кілька місяців, перебуваючи на лікуванні за кордоном. У країні поширювалися чутки про його смерть [14].

Для України deepfake не є теоретичною загрозою; це реальний інструмент, який держава-агресор, Російська Федерація, вже неодноразово застосовувала в рамках повномасштабної гібридної війни [15]. Аналіз двох ключових атак демонструє чітку еволюцію російської тактики.

Невдовзі після повномасштабного вторгнення, у момент найвищої невизначеності, російські актори поширили низькоякісний deepfake. На відео нібито Президент Зеленський закликав Збройні Сили України та громадян скласти зброю і "повернутися до своїх родин". Фейк був швидко викритий через очевидні технічні вади: невідповідний розмір голови до тіла, дивний, неприродний акцент та загальна низька якість.



Атака провалилася не лише через низьку якість. Головною причиною стало те, що українські спецслужби, зокрема Головне управління розвідки (ГУР МО) та Служба безпеки України (СБУ), заздалегідь (ще 3 березня 2022 року) попередили суспільство про те, що Росія готує саме такий фейк про "капітуляцію". Коли фейк нарешті з'явився (16 березня), населення вже було "вакциноване". Ворожа зброя не посіяла сумнівів, а, навпаки, підтвердила довіру до власних спецслужб. Це була видатна перемога в інформаційній війні.

Російські мережі поширили аудіо-deepfake, на якому нібито тодішній Головнокомандувач ЗСУ Валерій Залужний у розмові з американськими посадовцями закликав до військового перевороту та непокори політичному керівництву. Управління стратегічних комунікацій (СтратКом) ЗСУ було змушене негайно виступати з офіційним спростуванням та заявою про єдність керівництва країни.

Порівняння цих двох атак демонструє чітку еволюцію російської тактики. Фейк із Зеленським (березень 2022) був грубою, високоризикованою атакою, спрямованою на стратегічну деморалізацію нації. Вона провалилася. Аудіофейк із Залужним (листопад 2023) був значно хитрішою атакою. Він використовував аудіо (яке легше підробити та важче візуально спростувати) і був націлений на вже існуючий (і активно роздмуханий російською пропагандою) наратив про "конфлікт між військовим та політичним керівництвом".

У відповідь на загрозу deepfake виникла ціла індустрія, що займається їх виявленням (детекцією). Однак аналіз найновіших досліджень 2024-2025 років показує тверезу реальність: ця гонка озброєнь вже програна.

Методи виявлення, на які покладаються організації, поділяються на дві категорії. Пасивні. Засновані на Артефактах: Ці методи покладаються на пошук мікроскопічних недоліків, які генеративні моделі залишають у контенті і які невидимі для людського ока. Вони аналізують патерни моргання (ранні фейки "не моргали"), біологічні сигнали (наприклад, непомітну пульсацію крові на обличчі), невідповідності освітлення та тіней, артефакти на краях об'єктів (edge analysis) та аномалії в аудіоспектрограмах [16].

Активні / Засновані на Глибокому Навчанні: Цей підхід використовує ШІ для боротьби з ШІ. Моделі (такі як XceptionNet або Swin Transformer V2) навчаються на мільйонах прикладів фейків, щоб розпізнавати унікальні "відбитки пальців" або артефакти, залишені конкретними моделями-генераторами [17].

Хоча ці детектори демонструють високу точність (99% і вище) в академічних, лабораторних умовах, вони виявляються катастрофічно неефективними, коли стикаються з реальними загрозами "в дикій природі" (in-the-wild).

Нова парадигма безпеки, яку вже впроваджують провідні світові гравці, звучить так: "Не шукай підробку — верифікуй оригінал". Замість того, щоб намагатися довести, що щось не є справжнім, ця стратегія зосереджена на наданні криптографічного доказу того, що щось є справжнім.

Глобальні підходи до регулювання суттєво відрізняються, створюючи складний правовий ландшафт. Європейський Союз: The AI Act.

ЄС обрав проактивний, орієнтований на процес підхід. The AI Act класифікує deepfake як системи "обмеженого ризику" (limited risk). Ключове правило — прозорість. Користувачі повинні бути чітко поінформовані, що вони взаємодіють зі згенерованим контентом (наприклад, бачать deepfake або спілкуються з чат-ботом). Закон не забороняє технологію, але встановлює величезні штрафи за порушення правил прозорості — до €35 мільйонів або 7% світового річного доходу компанії.

Сполучені Штати: "Клаптиковий" Підхід. США обрали реактивний, орієнтований на шкоду підхід.

Федеральний рівень: У травні 2025 року було ухвалено "TAKE IT DOWN Act". Цей закон не стосується політики; він спрямований виключно на криміналізацію порнографії без згоди (NCII). Він зобов'язує платформи видаляти такий контент протягом 48 годин за запитом жертви.

Рівень штатів: Існує "клаптикова ковдра" законів. До кінця 2024 року 20 штатів ухвалили власні закони. Наприклад, закон Техасу криміналізує використання політичних deepfake з метою впливу на вибори за 30 днів до дня голосування [17].

Україна: Правова Прогалина. Станом на сьогодні, в Україні відсутнє будь-яке спеціальне законодавство, що регулює створення або поширення deepfake. Юристи змушені покладатися на загальні норми Цивільного кодексу про захист честі, гідності та права на образ, а також на закони про інтелектуальну власність та авторське право. Цей інструментарій є абсолютно недостатнім для протидії складним гібридним загрозам національній безпеці.

Висновки. Аналіз технологічних, політичних та безпекових аспектів deepfake-технологій дозволяє сформулювати чотири ключові висновки:

Загроза є Демократизованою та Масштабованою. Технологія пройшла шлях від академічної розробки до зброї, доступної будь-кому. Вартість атаки колапсувала (до \$1 за рободзвінок), а вимоги до даних мінімізувалися (до 3-5 секунд аудіо). Це означає, що загроза більше не обмежується публічними особами, а стосується будь-якого громадянина, військового чи посадовця. Основна Зброя — Психологічна. Найбільш руйнівним наслідком є не сам обман, а "Дивіденд Брехуна". Це здатність зловмисників використовувати саму ідею deepfake для заперечення реальних доказів, що руйнує довіру до інституцій, судової системи та медіа, роблячи неможливою політичну відповідальність. Кейс Габону довів, що підозра у фейку може бути такою ж дестабілізуючою, як і сам фейк.

Рішення — Проактивна Автентифікація. Новий глобальний стандарт — не шукати підробки, а верифікувати оригінали. Впровадження технології C2PA (Content Credentials) такими гігантами, як Google, Meta, OpenAI та, що ключове, Міністерством оборони США, сигналізує про фундаментальний доктринальний зсув, який Україна не може ігнорувати.



Надати державну підтримку (включно з інтеграцією на платформі "Дія.Освіта") для масштабування перевірених моделей, таких як проєкт "Ба та Ді проти брехні". Мета — охопити мільйони людей похилого віку, які, згідно з даними, є найбільш вразливими до шахрайства та дезінформації. Розробити аналогічні цільові програми для інших груп ризику, зокрема, для родин військовослужбовців (вразливих до аудіо-шахрайства) та державних службовців.

Література:

1. Медійна реальність у стилі deep fake // Реанімаційний Пакет Реформ. – URL: <https://rpr.org.ua/news/mediyna-real-nist-u-styli-deep-fake> (дата звернення: 05.11.2025).
2. Loux M., Loux B. What Is Deepfake Technology? Understanding Its Broad Impact / M. Loux, B. Loux // *American Public University System (APUS)*. – URL: <https://www.amu.apus.edu/area-of-study/information-technology/resources/what-is-deepfake-technology/> (дата звернення: 05.11.2025).
3. Boucher P.N. What if deepfakes made us doubt everything we see and hear? / P.N. Boucher // *Think Tank. European Parliament*. – 2021. – URL: [https://www.europarl.europa.eu/thinktank/iv/document/EPRS_ATA\(2021\)690046](https://www.europarl.europa.eu/thinktank/iv/document/EPRS_ATA(2021)690046) (дата звернення: 05.11.2025).
4. Somers M. Deepfakes, explained / M. Somers // *MIT Sloan School of Management*. – URL: <https://mitsloan.mit.edu/ideas-made-to-matter/deepfakes-explained> (дата звернення: 05.11.2025).
5. What Is Deepfake and How to Use It // *DiscoverDataScience.org*. – URL: <https://www.discoverdatascience.org/articles/everything-you-need-to-know-about-how-to-use-deepfake/> (дата звернення: 05.11.2025).
6. Contextualizing Deepfake Threats to Organizations // U.S. Department of Defense. – URL: <https://media.defense.gov/2023/Sep/12/2003298925/-1/-1/0/CSI-deepfakethreats.pdf> (дата звернення: 05.11.2025).
7. NSA, U.S. Federal Agencies Advise on Deepfake Threats // National Security Agency (NSA). – URL: <https://www.nsa.gov/Press-Room/Press-Releases-Statements/Press-Release-View/Article/3523329/nsa-us-federal-agencies-advise-on-deepfake-threats/> (дата звернення: 05.11.2025).
8. Deepfake Statistics & Trends 2025: Key Data & Insights // Keepnet Labs. – URL: <https://keepnetlabs.com/blog/deepfake-statistics-and-trends> (дата звернення: 05.11.2025).
9. Mirza R. How AI deepfakes threaten the 2024 elections / R. Mirza // *Journalist's Resource*. – 16 Feb. 2024. – URL: <https://journalistsresource.org/home/how-ai-deepfakes-threaten-the-2024-elections/> (дата звернення: 05.11.2025).
10. AI-generated misinformation in the election year 2024: measures of European Union // *Frontiers in Political Science*. – 2024. – Vol. 6. – Article 1451601. – URL: <https://www.frontiersin.org/journals/political-science/articles/10.3389/fpos.2024.1451601/full> (дата звернення: 05.11.2025).
11. AI Deepfake of Deceased Dictator Used to Influence Indonesian Election // *OECD.AI*. – URL: <https://oecd.ai/en/incidents/2024-02-12-8d5b> (дата звернення: 05.11.2025).
12. Kapoor S., Narayanan A. We Looked at 78 Election Deepfakes. Political Misinformation Is Not an AI Problem / S. Kapoor, A. Narayanan // *Knight First Amendment Institute at Columbia University*. – 13 Dec. 2024. – URL: <https://knightcolumbia.org/blog/we-looked-at-78-election-deepfakes-political-misinformation-is-not-an-ai-problem> (дата звернення: 05.11.2025).
13. Goldstein J.A., Lohn A. Deepfakes, Elections, and Shrinking the Liar's Dividend / J.A. Goldstein, A. Lohn // *Brennan Center for Justice*. – 23 Jan. 2024. – URL: <https://www.brennancenter.org/our-work/research-reports/deepfakes-elections-and-shrinking-liars-dividend> (дата звернення: 05.11.2025).

14. Deepfake in Gabon: Ali Bongo Case // Mother Jones. – URL: <https://www.motherjones.com/politics/2019/03/deepfake-gabon-ali-bongo/> (дата звернення: 05.11.2025).

15. Amerini I., Barni M., Battiato S. та ін. Deepfake Media Forensics: Status and Future Challenges / I. Amerini, M. Barni, S. Battiato [та ін.] ; ред. H. Cherifi // Journal of Imaging. – 2023. – Vol. 9, № 6. – Article 125. – URL: <https://www.mdpi.com/2313-433X/9/6/125> (дата звернення: 05.11.2025).

16. Harding L. Deepfakes v pre-bunking: is Russia losing the infowar? / L. Harding // The Guardian. – 19 Mar. 2022. – URL: <https://www.theguardian.com/world/2022/mar/19/russia-ukraine-infowar-deepfakes> (дата звернення: 05.11.2025).

17. Gray C.H. Political Deepfakes and Elections / C.H. Gray // The First Amendment Encyclopedia. – 6 Dec. 2024. – Updated: 11 Jan. 2025. – URL: <https://firstamendment.mtsu.edu/article/political-deepfakes-and-elections/> (дата звернення: 05.11.2025).

References:

1. Mediina realnist u styli deep fake [Media reality in the deepfake style]. (n.d.). Reanimatsiyni Paket Reform. Retrieved November 5, 2025, from <https://rpr.org.ua/news/mediyna-real-nist-u-styli-deep-fake> [in Ukrainian]

2. Loux, M., & Loux, B. (n.d.). What is deepfake technology? Understanding its broad impact. American Public University System (APUS). Retrieved November 5, 2025, from <https://www.amu.apus.edu/area-of-study/information-technology/resources/what-is-deepfake-technology/>

3. Boucher, P. N. (2021). What if deepfakes made us doubt everything we see and hear? European Parliament Think Tank. Retrieved November 5, 2025, from [https://www.europarl.europa.eu/thinktank/lv/document/EPRS_ATA\(2021\)690046](https://www.europarl.europa.eu/thinktank/lv/document/EPRS_ATA(2021)690046)

4. Somers, M. (n.d.). Deepfakes, explained. MIT Sloan School of Management. Retrieved November 5, 2025, from <https://mitsloan.mit.edu/ideas-made-to-matter/deepfakes-explained>

5. Discover Data Science. (n.d.). What is deepfake and how to use it. Retrieved November 5, 2025, from <https://www.discoverdatascience.org/articles/everything-you-need-to-know-about-how-to-use-deepfake/>

6. U.S. Department of Defense. (2023). Contextualizing deepfake threats to organizations. Retrieved November 5, 2025, from <https://media.defense.gov/2023/Sep/12/2003298925/-1/-1/0/CSI-deepfakethreats.pdf>

7. National Security Agency (NSA). (n.d.). NSA, U.S. federal agencies advise on deepfake threats. Retrieved November 5, 2025, from <https://www.nsa.gov/Press-Room/Press-Releases-Statements/Press-Release-View/Article/3523329/nsa-us-federal-agencies-advise-on-deepfake-threats/>

8. Keepnet Labs. (n.d.). Deepfake statistics & trends 2025: Key data & insights. Retrieved November 5, 2025, from <https://keepnetlabs.com/blog/deepfake-statistics-and-trends>

9. Mirza, R. (2024, February 16). How AI deepfakes threaten the 2024 elections. Journalist's Resource. Retrieved November 5, 2025, from <https://journalistsresource.org/home/how-ai-deepfakes-threaten-the-2024-elections/>

10. AI-generated misinformation in the election year 2024: Measures of European Union. (2024). Frontiers in Political Science, 6, Article 1451601. Retrieved November 5, 2025, from <https://www.frontiersin.org/journals/political-science/articles/10.3389/fpos.2024.1451601/full>

11. OECD.AI. (n.d.). AI deepfake of deceased dictator used to influence Indonesian election. Retrieved November 5, 2025, from <https://oecd.ai/en/incidents/2024-02-12-8d5b>

12. Kapoor, S., & Narayanan, A. (2024, December 13). We looked at 78 election deepfakes. Political misinformation is not an AI problem. Knight First Amendment Institute at Columbia University. Retrieved November 5, 2025, from <https://knightcolumbia.org/blog/we-looked-at-78-election-deepfakes-political-misinformation-is-not-an-ai-problem>

13. Goldstein, J. A., & Lohn, A. (2024, January 23). Deepfakes, elections, and shrinking the liar's dividend. Brennan Center for Justice. Retrieved November 5, 2025, from <https://www.brennancenter.org/our-work/research-reports/deepfakes-elections-and-shrinking-liars-dividend>



14.Mother Jones. (n.d.). Deepfake in Gabon: Ali Bongo case. Retrieved November 5, 2025, from <https://www.motherjones.com/politics/2019/03/deepfake-gabon-ali-bongo/>

15.Amerini, I., Barni, M., Battiato, S., et al. (2023). Deepfake media forensics: Status and future challenges. *Journal of Imaging*, 9(6), 125. Retrieved November 5, 2025, from <https://www.mdpi.com/2313-433X/9/6/125>

16.Harding, L. (2022, March 19). Deepfakes v pre-bunking: Is Russia losing the infowar? *The Guardian*. Retrieved November 5, 2025, from <https://www.theguardian.com/world/2022/mar/19/russia-ukraine-infowar-deepfakes>

17.Gray, C. H. (2024, December 6; updated 2025, January 11). Political deepfakes and elections. *The First Amendment Encyclopedia*. Retrieved November 5, 2025, from <https://firstamendment.mtsu.edu/article/political-deepfakes-and-elections/>