

3. Роботу над даним проектом доцільно проводити за такими етапами:

- організаційно-мотиваційним;
- проміжних звітів учасників проєктних груп щодо виконання досліджень;
- висновків проєктних груп щодо реалізації мети проєкту і варіантів оприлюднення його результатів;
- розробка матеріалів на допомогу використання ІІІ під час розв'язування методичних вправ на засадах дотримання академічної доброчесності.

Список використаних джерел:

1. Ivanova, S., Dimitrov, L., Ivanov, V., & Naleva, G. (2021). The Performance of project teams selected based on student personality types: a longitudinal study. *Advances in Science, Technology and Engineering Systems Journal*, vol. 6, no. 1, pp. 1128-1136. <https://doi.org/10.25046/aj0601126>

2. Ivanova, S., Dimitrov, L., Ivanov, V., & Prokopovych, L. (2021, May). Using Role-playing Game for Professional Skills Formation of Prospective Teachers. In *Society, Integration, Education. Proceedings of the International Scientific Conference* (Vol. 1, pp. 195-206). DOI: <https://doi.org/10.17770/sie2021vol1.6180>

3. Іванова С.В., Орленко І.М., Задоріна О.В. Когнітивно-логічний розвиток дітей в інклюзивному середовищі: досвід використання індивідуальних адаптивних систем у логопедичній практиці. *Інноваційна педагогіка*. Вит. 93, том 1, 2026. С. 124-131.

DOI <https://doi.org/10.32782/ip/93.1.23>

https://www.innovpedagogy.od.ua/archives/2026/93/part_1/25.pdf

Киричук Л.О

Здобувач першого (бакалаврського) рівня вищої освіти Вінницького державного педагогічного університету імені Михайла Коцюбинського

Науковий керівник: Пилипенко Т.І.

старший викладач кафедри фундаментальних і приватно-правових дисциплін Вінницького державного педагогічного університету імені Михайла Коцюбинського

**ПРАВОВЕ РЕГУЛЮВАННЯ ДИПФЕЙКІВ: МІЖНАРОДНИЙ ДОСВІД
ТА ПЕРСПЕКТИВИ ДЛЯ УКРАЇНИ**

Стрімкий прогрес у сфері технологій штучного інтелекту зумовив появу інноваційних інструментів обробки великих обсягів даних та генерації медіаконтенту. Серед них особливе місце посідає технологія «дипфейк»

(deepfake), що базується на використанні генеративно-змагальних нейромереж, які дають змогу з високою точністю відтворювати голос, зовнішність та поведінку людини. Такі розробки відкривають широкі перспективи для кіноіндустрії, реклами, медіабізнесу та сучасних форматів навчання. Проте водночас ця технологія створює безпрецедентні виклики й загрози для безпеки особи, суспільства та стабільності державних інститутів. За відсутності дієвих нормативно-правових бар'єрів генеративний ШІ перетворюється на інструмент маніпуляцій, що вимагає термінового втручання з боку правників і законодавчих органів.

Нагальність правового дослідження цієї проблематики зумовлена геометричним зростанням кількості кіберзлочинів, пов'язаних із використанням дипфейків. Сфабриковані аудіо- та відеоматеріали дедалі частіше застосовуються для фінансових махінацій, шантажу, руйнування ділової репутації компаній, а також для брудних політичних технологій та масового обману аудиторії.

Особливу гостроту зазначене питання набуває для України в умовах ведення повномасштабної війни, де технології ШІ активно експлуатуються ворогом як елемент гібридної агресії та проведення інформаційно-психологічних операцій (ІПСО). Метою таких дій є дезорієнтація цивільного населення, дискредитація військово-політичного керівництва країни та підрив національної безпеки загалом.

Оскільки чинне вітчизняне законодавство тривалий час формувалося без урахування ризиків високотехнологічних цифрових підробок, виникає гостра практична й теоретична потреба в аналізі світового досвіду та розробці власної ефективної моделі правового контролю. Тому мета дослідження полягає у здійсненні комплексного теоретико-правового аналізу правових викликів, пов'язаних із поширенням технології «дипфейк», дослідженні провідних міжнародних і європейських практик юридичного контролю за системами штучного інтелекту, а також у визначенні конкретних шляхів модернізації законодавчої бази України відповідно до сучасних викликів комп'ютерної ери задля надійного захисту прав особи та інформаційної стабільності суспільства.

Міжнародний досвід свідчить про те, що світ уже переходить від теоретичного осмислення до конкретних нормативних рішень. Головним гравцем у цьому процесі виступає Європейський Союз. Ключовим і фундаментальним документом у зазначеній сфері став Регламент Європейського Парламенту і Ради (ЄС) 2024/1689 від 13 червня 2024 року, більш відомий у науковій періодиці як Європейський закон про штучний інтелект (EU AI Act).

Цей нормативно-правовий акт побудований на ризико-орієнтованому підході, що передбачає класифікацію систем ШІ за чотирма рівнями небезпеки: від неприйнятної до мінімальної. Відповідно до положень AI Act, розробники та користувачі генераторів дипфейків належать до категорії систем із високими вимогами до прозорості.

Європейський законодавець імперативно зобов'язує маркувати будь-який штучно створений або маніпулятивний аудіо-, фото- та відеоконтент. Кінцевий споживач інформації має бути чітко, зрозуміло та своєчасно поінформований про те, що він взаємодіє із продуктом штучного інтелекту, а не з реальною фізичною особою чи автентичним записом. Порухення цих вимог тягне за собою значні фінансові санкції та штрафи для компаній-розробників.

Паралельно з європейським вектором власні правові моделі активно розбудовують Сполучені Штати Америки. Регулювання тут розвивається за двома основними напрямками: на загальнодержавному (федеральному) рівні та на рівні окремих штатів. Зокрема, федеральні ініціативи фокусуються на захисті виборчих процесів від маніпуляцій з боку політичних опонентів чи іноземних суб'єктів. Водночас на рівні штатів (наприклад, у Каліфорнії, Нью-Йорку та Техасі) криміналізовано створення та поширення порнографічних дипфейків без згоди особи (*non-consensual deepfake pornography*), що визнано грубим порушенням фундаментальних прав людини на приватність, честь та особисту недоторканність.

Для України питання правового регулювання дипфейків має екзистенційне значення. Традиційні інструменти Цивільного кодексу України, покликані захищати право особи на ім'я, зображення та ділову репутацію, орієнтовані на класичні форми правопорушень (наклеп у пресі, несанкціоноване використання фотографій) і виявляються абсолютно неефективними в умовах миттєвого вірусного поширення ШІ-генерацій у соціальних мережах [5, С. 14]. Донедавна потерпіла особа практично не мала процесуальних можливостей оперативно заблокувати шкідливий контент або ідентифікувати анонімного творця дипфейку для притягнення його до відповідальності. Проте, останнім часом український законодавець демонструє суттєву інтенсифікацію нормотворчої діяльності.

На виконання Плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2025–2026 роки, затвердженого Кабінетом Міністрів, у державі офіційно розпочато процес системної адаптації норм до стандартів Європейського Союзу [3, с. 2]. Наша держава обрала шлях поступової, але впевненої гармонізації. Одним із перших реальних кроків стало прийняття

Закону України «Про академічну доброчесність». Він безпосередньо стосується освіти та науки. Тепер автори зобов'язані чесно вказувати, якщо у своїх наукових роботах вони використовували текст чи дані, згенеровані штучним інтелектом [2, с. 5]. Це закладає міцний фундамент для формування культури відповідального поводження з технологіями генеративного ШІ.

Наступним логічним етапом, що активно обговорюється у наукових та урядових колах, є підготовка спеціалізованого вітчизняного закону про штучний інтелект [4]. Експерти наголошують, що національне законодавство має не просто копіювати європейський AI Act, а враховувати специфічні воєнні та безпекові ризики, з якими стикається Україна. Зокрема, потребує деталізації та реформування Цивільний кодекс України в частині запровадження поняття «цифрового образу особи» та визначення меж цивільно-правової відповідальності платформ-хостингів та соціальних мереж, які допускають поширення немаркованого шкідливого ШІ-контенту.

Підбиваючи підсумки, зазначимо, що протидія загрозам використання технології «дипфейк» в Україні потребує комплексного підходу, який має охоплювати такі ключові напрями:

- запровадження обов'язкового маркування: нормативне закріплення обов'язку розробників ідентифікувати контент штучного інтелекту на етапі його створення за допомогою технології цифрових водяних знаків (watermarking).
- оперативне реагування на правопорушення: надання правоохоронним органам та адміністраторам інтернет-ресурсів дієвих юридичних інструментів для оперативного (зокрема, за спеціальною позасудовою процедурою) видалення шкідливих або недостовірних медіаматеріалів.
- розвиток медіаграмотності: систематичне підвищення рівня цифрової гігієни громадян, оскільки розвинене критичне мислення є первинним елементом захисту особи від дезінформації.

Тільки збалансоване поєднання провідних європейських стандартів із національними безпековими реаліями дасть змогу надійно захистити права людини в умовах стрімкої еволюції систем штучного інтелекту.

Список використаних джерел:

1. Регламент Європейського Парламенту і Ради (ЄС) 2024/1689 від 13 червня 2024 року, що встановлює гармонізовані правила щодо штучного інтелекту (Акт про штучний інтелект) та вносить зміни до деяких законодавчих актів Союзу // Офіційний вісник Європейського Союзу. 2024. L 2024/1689. URL: europa.eu (дата звернення: 04.06.2026).
2. Про академічну доброчесність : Закон України від 18.12.2025 № 4742-IX //

Відомості Верховної Ради України. 2026. № 6. Ст. 42.

3. Про затвердження плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2025–2026 роки : розпорядження Кабінету Міністрів України від 09.05.2025 № 457-р // Урядовий кур'єр. 2025. № 92.

4. Берназюк Я. Проєкт Закону України «Про штучний інтелект»: концептуальні та правові підходи. Constitutionalist. 2026. URL: <http://constitutionalist.com.ua> (дата звернення: 04.06.2026).

5. Харитонов Є. О., Харитонova О. І. Цивільно-правове регулювання відносин у сфері використання штучного інтелекту: виклики та перспективи 2025–2026 рр. Часопис цивілістики. 2025. Т. 54. С. 12–19.

Козак Валентина Андріївна,

кандидат філологічних наук, доцент,

Київський інститут Національної гвардії України

ORCID: 0009-0000-0121-4949

ШТУЧНИЙ ІНТЕЛЕКТ ЯК ВИКЛИК АКАДЕМІЧНИЙ ДОБРОЧЕСНОСТІ: РИЗИКИ «ЦИФРОВОГО ПЛАГІАТУ» У НАУКОВИХ РОБОТАХ ЗДОБУВАЧІВ ВИЩОЇ ОСВІТИ

Генеративний штучний інтелект – зокрема великі мовні моделі типу ChatGPT – стрімко увійшов в освітній простір і змінив усталені уявлення про авторство наукового тексту. Інструменти, покликані полегшити дослідницьку роботу, на практиці розмили межу між власним інтелектуальним зусиллям та машинною генерацією. Це ставить під загрозу не лише якість навчальних робіт здобувачів вищої освіти, а й принципи академічної доброчесності.

Університетська спільнота зіткнулася з явищем «цифрового плагіату» – прихованого академічного шахрайства, за якого здобувач вищої освіти передає інтелектуальну функцію алгоритму, не осмислюючи й не перевіряючи результат. Без сформованих навичок наукового письма та критичного мислення здобувачі вищої освіти легко потрапляють у пастку ілюзорної простоти генерованого тексту.

Широкий контекст етичних та безпекових викликів, пов'язаних із ШІ, сьогодні активно порушується у працях багатьох українських та зарубіжних науковців. Зокрема, стаття О. Петасюк «Штучний інтелект: контекст та осмислення» досліджує феномен ШІ в цифрову епоху, наголошуючи на його перевагах у медицині, науці, освіті та військовій сфері. Авторка переконує, що