

МІНІСТЕРСТВО ВНУТРІШНІХ СПРАВ УКРАЇНИ

НАЦІОНАЛЬНА АКАДЕМІЯ ВНУТРІШНІХ СПРАВ

Навчально-науковий інститут права та психології

Кафедра інформаційних технологій



**ПРОБЛЕМИ ТА ПЕРСПЕКТИВИ ВИКОРИСТАННЯ
МОЖЛИВОСТЕЙ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ
В ПРАВІ, ПСИХОЛОГІЇ
ТА ПРАВООХОРОННІЙ ДІЯЛЬНОСТІ**

МАТЕРІАЛИ

**Міжнародної науково-практичної конференції
молодих вчених та здобувачів вищої освіти
(м. Київ, 27 травня 2026 року)**



КИЇВ-2026

Організаційний комітет конференції:

Корнейко О. В., голова комітету, канд. техн. наук, професор (НАВС)

Вітязь А. А., заступник голови комітету (НАВС)

Швець М. М., відповідальний секретар (НАВС)

Програмний комітет конференції:

Кудінов В. А., голова комітету, канд. фіз.-мат. наук, доцент (НАВС)

Andre Samberg, доктор наук, професор практики (Фінляндія)

Корнейко О. В., канд. техн. наук, професор (НАВС)

Хахановський В. Г., доктор юрид. наук, професор (НАВС)

Укладачі матеріалів конференції:

Корнейко О. В., канд. техн. наук, професор (НАВС)

Швець М. М. (НАВС)

Матеріали схвалено та рекомендовано до друку на засіданні науково-методичної ради Національної академії внутрішніх справ (протокол № 7 від «21» липня 2026 року)

П 78 Проблеми та перспективи використання можливостей систем штучного інтелекту в праві, психології та правоохоронній діяльності
[Текст] : матеріали Міжнародної науково-практичної конференції молодих вчених та здобувачів вищої освіти (м. Київ, 27 травня 2026 року) / [уклад.: О. В. Корнейко, М. М. Швець]. – Київ : Нац. акад. внутр. справ, 2026. – 120 с.

Представлені тези доповідей учасників Міжнародної науково-практичної конференції молодих вчених та здобувачів вищої освіти на тему «Проблеми та перспективи використання можливостей систем штучного інтелекту в праві, психології та правоохоронній діяльності», яка була організована та проводилась кафедрою інформаційних технологій навчально-наукового інституту права та психології Національної академії внутрішніх справ 27 травня 2026 року напередодні святкування 105-річниці академії.

Матеріали тез доповідей подано в авторській редакції учасників конференції. За достовірність фактів, статистичних та інших даних, точність формулювань і висновків несуть відповідальність автори матеріалів.

ЗМІСТ

Стор.

ПЕРЕДМОВА	6
------------------------	---

Тези доповідей на пленарному засіданні конференції

Олексій КОСТЕНКО. ШТУЧНИЙ ІНТЕЛЕКТ У ПРАВІ, ПСИХОЛОГІЇ ТА ПРАВООХОРОННІЙ ДІЯЛЬНОСТІ INTERPOL: ВІД МЕТАЗЛОЧИНІВ ДО ЦИФРОВОГО КОДЕКСУ.....	7
Tim D. WILLIAMS. WHY AI LEGAL ASSISTANTS HALLUCINATE IN NEW FIELDS: A MATHEMATICAL LIMIT, WITH POST-QUANTUM CRYPTOGRAPHY (PQC) AS EXAMPLE.....	14
Ірина ГЛОВІЮК. ІІІ ТА ЗАКОН УКРАЇНИ «ПРО АКАДЕМІЧНУ ДОБРОЧЕСНІСТЬ»: ПОДАЛЬШІ ПИТАННЯ	17
Віра ГУСЬКОВА, Владислав ШКОЛЬНИКОВ, Богдан ЛИСОВ, Артем ХАЛИГОВ. ГІБРИДНА АРХІТЕКТУРА ІНТЕЛЕКТУАЛЬНОЇ ФІЛЬТРАЦІЇ КІБЕРПОДІЙ НА ОСНОВІ RULE-BASED GATE ТА МОДЕЛЕЙ ШТУЧНОГО ІНТЕЛЕКТУ	20
Юлія АНТОНОВА, Максим ПАЛЬЧИК. МІЖНАРОДНІ ПІДХОДИ ДО ПРАВОВОГО РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ: ЄС, США ТА УКРАЇНИ	24
Вікторія ГУМЕНЮК, Олександр ПАКРИШ. МОЖЛИВОСТІ ТА РИЗИКИ ЗАСТОСУВАННЯ ІІІ В ПСИХОТЕРАПІЇ	27
Anastasiia PETRUS. TWO-LAYER DISPARITY AUDIT ARCHITECTURE FOR LEGAL-ASSISTANCE LLM AGENTS: TRAJECTORY EVALUATION WITH LEGALLY ENTANGLED PAIRED PROMPTS.....	29
Сніжана ГАЙДУК, Олександр КОРНЕЙКО. ПРАВОВЕ РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УКРАЇНІ ТА ЄС	32
Богдан ШУЛЯК, Вадим КУДІНОВ. СУЧАСТНІСТЬ ТА ПЕРСПЕКТИВИ ВПРОВАДЖЕННЯ ІІІ-АСИСТЕНТІВ У ПРАКТИКУ ДІЯЛЬНОСТІ ЮРИСКОНСУЛЬТІВ, НОТАРІУСІВ, СУДДІВ, АДВОКАТІВ ТА ПСИХОЛОГІВ	34
Іван ПИСАР, Валерій ХАХАНОВСЬКИЙ. ПСИХОСОЦІАЛЬНІ ТА ЕТИЧНІ РИЗИКИ ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ІІІ.....	37
Анастасія НАСАДЮК, Наталія КОБЗЕЙ. ЕТИЧНІ АСПЕКТИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПІДГОТОВЦІ ФІЛОЛОГІВ-ПЕРЕКЛАДАЧІВ.....	39

Тези стендових доповідей учасників конференції

Володимир ПАЛИВОДА. ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ДІЯЛЬНІСТЬ СПЕЦІАЛЬНИХ СЛУЖБ США: ОСНОВНІ НАПРЯМИ ТА ВИКЛИКИ.....	42
Оксана ЯРЕМА. РОЗВИТОК ПРАВОВОГО РЕГУЛЮВАННЯ ВІДНОСИН ІЗ РОЗРОБКИ ТА ВПРОВАДЖЕННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ	44
Едуард РИЖКОВ. АРХІТЕКТУРА МЕТОДИКИ «OSINT+» У КОНТЕКСТІ ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ	46

Ольга ОГІРКО, Наталія ГАЛАЙКО. ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В АДМІНІСТРУВАННІ ІНФОРМАЦІЙНИХ СИСТЕМ НАЦІОНАЛЬНОЇ ПОЛІЦІЇ УКРАЇНИ.....	48
Владислав ГАВЛОВСЬКИЙ, Михайло ГУЦАЛЮК. ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ПРОТИДІЇ ОРГАНІЗОВАНИЙ ЗЛОЧИННОСТІ.....	51
Сергій ЄСІМОВ. ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ОСВІТІ	54
Роксолана ЗАПОТІЧНА. ЕКОНОМІЧНА ЕФЕКТИВНІСТЬ ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В ОСВІТНІЙ ТА НАУКОВІЙ ДІЯЛЬНОСТІ	57
Василь АНДРУНИК, Наталія КАЛЬКА. ДОЦІЛЬНІСТЬ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ЯК ІНСТРУМЕНТУ ВДОСКОНАЛЕННЯ ТЕХНІКИ «БЕЗПЕЧНЕ МІСЦЕ» У ПСИХОТЕРАПЕВТИЧНІЙ ПРАКТИЦІ	59
Олексій СКІЦЬКО. ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В КІБЕРБЕЗПЕЦІ: МОЖЛИВОСТІ, ЗАГРОЗИ ТА ВИКЛИКИ	61
Роман ШИРШОВ. ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В РОЗВІДЦІ ЗА ВІДКРИТИМИ ДЖЕРЕЛАМИ (OSINT): МОЖЛИВОСТІ, ІНСТРУМЕНТИ ТА ОБМЕЖЕННЯ	63
Євгенія ЛИТВИНЕНКО. ВИКОРИСТАННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ У ДІЯЛЬНОСТІ ПРАВООХОРОННИХ ОРГАНІВ УКРАЇНИ: ВИКЛИКИ ПРАВОВОГО РЕГУЛЮВАННЯ ТА ГАРАНТІЇ ПРАВ ЛЮДИНИ.....	67
Андрій САВЧУК. ІНТЕГРАЦІЯ ШТУЧНОГО ІНТЕЛЕКТУ В АЛГОРИТМИ СОРТУВАННЯ ПОРАНЕНИХ НА ПОЛІ БОЮ: ПЕРСПЕКТИВИ, ОБМЕЖЕННЯ ТА ЕТИЧНІ ДИЛЕМИ ПРИЙНЯТТЯ КЛІНІЧНИХ РІШЕНЬ	69
Дмитро КОБИЛИНСЬКИЙ, Едуард РИЖКОВ. ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У ДІЯЛЬНОСТІ ПІДРОЗДІЛІВ СТРАТЕГІЧНИХ РОЗСЛІДУВАНЬ.....	70
Анатолій ТІХОНОВ, Ольга УСТЮЖАНІНОВА. АКАДЕМІЧНА ДОБРОЧЕСНІСТЬ В УМОВАХ ВИКОРИСТАННЯ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ: ПРАВОВІ ПРОБЛЕМИ ТА ПЕРСПЕКТИВИ РЕГУЛЮВАННЯ.....	74
Ірина БАЗИЛЯК, Олена ВОРОБЕЙ. ЗАСТОСУВАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В OSINT-РОЗВІДЦІ.....	77
Максим КОРШУНОВ, Оксана БЕРЕГУЛЯК. ВИКОРИСТАННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ АВТОМАТИЗОВАНОГО АНАЛІЗУ ВРАЗЛИВОСТЕЙ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ	79
Антон ЛИСЕНКО, Олександр СКЛЯР. ШТУЧНИЙ ІНТЕЛЕКТ У ПРАВООХОРОННІЙ ДІЯЛЬНОСТІ: МОЖЛИВОСТІ ТА РИЗИКИ ЗАСТОСУВАННЯ.....	80
Дар'я КАРАСЬОВА, Каріна БОРТУН. ШТУЧНИЙ ІНТЕЛЕКТ У ЮРИДИЧНІЙ ПРАКТИЦІ: МЕЖІ ВІДПОВІДАЛЬНОСТІ ТА ПЕРСПЕКТИВИ РОЗВИТКУ	83
Анастасія ЯКОВЕЦЬ. ФЕНОМЕН «КОГНІТИВНОЇ ЛІНІ» В УМОВАХ РОЗВИТКУ ШТУЧНОГО ІНТЕЛЕКТУ	84

Юрій ДОЛГІХ, Сергій НЕХАСЬНКО. ЗАГРОЗИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ПСИХОЛОГІЧНОЇ ПІДТРИМКИ ВІЙСЬКОВОСЛУЖБОВЦІВ В УМОВАХ ВІДСІЧІ ЗБРОЙНИЙ АГРЕСІЇ РОСІЇ ПРОТИ УКРАЇНИ	86
Дарія КОВЕНЬКО, Микола ШВЕЦЬ. ТЕХНОЛОГІЇ ШТУЧНОГО ІНТЕЛЕКТУ В УКРАЇНСЬКИХ ІНФОРМАЦІЙНИХ СИСТЕМАХ	89
Вікторія ЗАХАРОВА, Роман ДЕМКІВ. ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ПІД ЧАС ЗДІЙСНЕННЯ ПРАВОСУДДЯ: ПЕРСПЕКТИВИ ТА РИЗИКИ	93
Володимир ШИППІ, Ірина БОТНАРЕНКО. ШТУЧНИЙ ІНТЕЛЕКТ У СУДОЧИНСТВІ ТА ДІЯЛЬНОСТІ ПОЛІЦІЇ: ІННОВАЦІЙНІ МОЖЛИВОСТІ ТА РИЗИКИ ДЛЯ ПРАВ ЛЮДИНИ	91
Владислава ІВАНОВА, Катерина ВЛАДОВСЬКА. КОНСТИТУЦІЙНО-ПРАВОВІ АСПЕКТИ ЦИФРОВІЗАЦІЇ ПУБЛІЧНОЇ ВЛАДИ: МЕЖІ АВТОМАТИЗАЦІЇ УПРАВЛІНСЬКИХ ПРОЦЕСІВ	95
Альбіна НАЗАРЕНКО, Андрій САВЧУК. ШТУЧНИЙ ІНТЕЛЕКТ ЯК ІНСТРУМЕНТ НАВЧАННЯ І САМООСВІТИ: ТЕХНОЛОГІЇ, МОЖЛИВОСТІ ТА РИЗИКИ	98
Віталій ПЕТРИК. КРИМІНАЛЬНА ВІДПОВІДАЛЬНІСТЬ ЗА НЕЗАКОННЕ ПОШИРЕННЯ ДИПФЕЙКІВ.....	100
Володимир ШАПОШНІКОВ. МОДЕЛІ МАШИННОГО НАВЧАННЯ ДЛЯ ПРОГНОЗУВАННЯ КІБЕРАТАК НА ОСНОВІ АНАЛІЗУ МЕРЕЖЕВОГО ТРАФІКУ	102
Анна РАБІЙЧУК, Анна ОМЕЛЯН, Вадим КУДІНОВ. ІІІ-КОМПЕТЕНТНОСТІ СПІВРОБІТНИКА ПРАВООХОРОННОГО ОРГАНУ У КОНТЕКСТІ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ ПРАВООХОРОННОЇ СИСТЕМИ	106
Максим ПАЦУЛА, Дмитро БОДИРЄВ. ЗАСТОСУВАННЯ ЦИФРОВИХ ТРЕНАЖЕРІВ ТА VR-ТЕХНОЛОГІЙ У СИСТЕМІ ВОГНЕВОЇ ПІДГОТОВКИ ПОЛІЦЕЙСЬКИХ	109
Михайло ЛИТВИНЕНКО, Едуард РИЖКОВ. МЕТОДИ АНАЛІТИКИ В ДІЯЛЬНОСТІ ПІДРОЗДІЛІВ КРИМІНАЛЬНОЇ ПОЛІЦІЇ.....	111
Валерія БУНЧУЖНА, Марина ЛОГІНОВА. НІКЧЕМНІ ПРАВОЧИНИ ТА ЇХ ПРАВОВІ НАСЛІДКИ	113
Юрій КОРОЛЬ, Віталій КУЧЕР. МОНІТОРИНГ ДОТРИМАННЯ ПРАВ І СВОБОД ЛЮДИНИ В ОРГАНАХ ПОЛІЦІЇ	115
Діана ХИЖА. ПРИНЦИП ЗАКОННОСТІ, ЯК ГАРАНТІЯ ЗАХИСТУ ПРАВ І СВОБОД ЛЮДИНИ В ПОЛІЦЕЙСЬКІЙ ДІЯЛЬНОСТІ	117

ПЕРЕДМОВА

27 травня 2026 року в Національній академії внутрішніх справ (НАВС) була проведена Міжнародна науково-практична конференція молодих вчених та здобувачів вищої освіти на тему: «Проблеми та перспективи використання можливостей систем штучного інтелекту в праві, психології та правоохоронній діяльності». Конференція була присвячена 105-річчю академії та організована кафедрою інформаційних технологій навчально-наукового інституту права та психології НАВС.

Участь у роботі конференції взяли відомі науковці та молоді вчені, аспіранти та здобувачі ступенів вищої освіти бакалавра і магістра з 15-ти закладів вищої освіти та наукових установ України, які здійснюють науковий пошук за тематикою роботи конференції. У роботі конференції також взяли участь представники двох зарубіжних університетів – школи права британського Відкритого університету (The Open University, Велика Британія) та факультету інформатики та статистики Sopra Steria SE Мюнхенського університету Людвіга-Максиміліана (Німеччина). Загалом організаторам конференції було надіслано 44 тези доповідей, які наведені у цьому виданні і з яких були відібрані найбільш цікаві для безпосереднього виступу на пленарному засіданні конференції.

На відкритті конференції її учасників від організаторів конференції привітали: завідувач кафедри інформаційних технологій, голова програмного комітету конференції, кандидат фізико-математичних наук, доцент Вадим Кудінов; професор цієї кафедри, модератор конференції, кандидат технічних наук, професор Олександр Корнейко; завідувач кафедри кримінології та інформаційних технологій, капітан поліції, доктор філософії в галузі права, доцент Владислав Школьніков.

На початку пленарного засідання конференції змістовні доповіді щодо сучасних проблем у сфері штучного інтелекту (ШІ) зробили: завідувач наукової лабораторії проблем імерсивних технологій і права Державної наукової установи «Інститут інформації, безпеки і права Національної академії правових наук України», доктор юридичних наук, старший дослідник Олексій Костенко; незалежний експерт з аспектів архітектури мережевої безпеки та ШІ, професор Open University Law School Tim D. Willsams; експерт-оцінювач проєктних пропозицій і член панелі Європейської Комісії у сферах безпеки та оборони (Брюссель, Бельгія), доктор наук, професор практики Andre Samberg.

Під час пленарного засідання учасники конференції обговорили ключові проблеми, перспективи та виклики застосування ШІ в юриспруденції, психології та правоохоронній діяльності, психосоціальні та етичні ризики використання технологій ШІ в освіті, науці та кібербезпеці, перспективи впровадження ШІ-асистентів та агентів у практику діяльності правоохоронців, юрисконсультів, нотаріусів, суддів, адвокатів та психологів.

З цікавими доповідями з цієї тематики виступили професорка Одеського державного університету внутрішніх справ, доктор юридичних наук, професор Ірина Гловюк, доцентка Національного технічного університету України «Київський політехнічний інститут імені Ігоря Сікорського», доктор філософії з комп'ютерних наук Віра Гуськова, викладач кафедри інформаційних технологій НАВС Антон Вітязь. На пленарному засіданні конференції свої доповіді презентували студенти НАВС Вікторія Гуменюк, Сніжана Гайдук, Богдан Шуляк, Кіра Стрільцова та Іван Писар. З інших закладів вищої освіти також виступили студенти Юлія Антонова, Анастасія Петрусь та Анастасія Насадюк.

Учасники конференції надійшли висновку про необхідність продовження досліджень щодо зазначеної тематики у сфері ШІ, публікацій їх результатів у наукових виданнях України та зарубіжних країн, а також їх обговорення під час проведення наступних конференцій.

Тези доповідей на пленарному засіданні конференції

Олексій КОСТЕНКО, завідувач наукової лабораторії імерсивних технологій і права наукового центру цифрової трансформації і права, Державна наукова установа «Інститут інформації, безпеки і права Національної академії правових наук України», доктор юридичних наук, старший дослідник

ШТУЧНИЙ ІНТЕЛЕКТ У ПРАВІ, ПСИХОЛОГІЇ ТА ПРАВООХОРОННИЙ ДІЯЛЬНОСТІ INTERPOL: ВІД МЕТАЗЛОЧИНІВ ДО ЦИФРОВОГО КОДЕКСУ

1. АКТУАЛЬНІСТЬ: ПРАВО У ПОСТЦИФРОВОМУ ПРОСТОРИ

Сучасна система права формувалася впродовж тисячоліть у фізичному, географічно локалізованому просторі. Людина як суб'єкт права, злочин як діяння у конкретному місці й часі, юрисдикція як суверенітет над визначеною територією — ці базові категорії відображали онтологічну реальність, в якій право функціонувало. Поява Web 3.0/4.0, генеративного штучного інтелекту та технологій розширеної реальності (XR) зруйнувала цю реальність, не запропонувавши натомість нових нормативних рамок.

Три паралельні тенденції формують новий правовий ландшафт. По-перше, ШІ-системи стають квазіавтономними учасниками правовідносин — не лише інструментами, а суб'єктами, чий рішення мають юридично значущі наслідки для конкретних осіб у сферах кримінального правосуддя, адміністративного управління та психологічних послуг. По-друге, простір правопорушень зміщується у Метавсесвіт: злочини вчиняються через аватари, без фізичного контакту, одночасно у декількох юрисдикціях, не залишаючи матеріальних слідів. По-третє, психологічна наука стикається з феноменом «алгоритмічного свідомості» — ШІ-систем, здатних маніпулювати поведінкою та психоемоційним станом осіб через персоналізований контент у XR-середовищах.

Реакція міжнародних інституцій підтверджує критичну серйозність цих тенденцій: ІНТЕРПОЛ у 2022 році запустив перший у світі правоохоронний Метавсесвіт, а у 2024 році оприлюднив White Paper з доктриною метазлочинів. Ці кроки перегукуються з авторськими концепціями цифрової юрисдикції та Цифрового кримінального кодексу, що розроблялися незалежно і паралельно у рамках вітчизняної правової науки [1; 2; 3].

2. ДОСВІД ІНТЕРПОЛ: ПРАВООХОРОННИЙ МЕТАВСЕСВІТ ТА ДОКТРИНА МЕТАЗЛОЧИНІВ

Поворотним моментом у визнанні Метавсесвіту як правоохоронного інструменту стала 90-та Генеральна Асамблея ІНТЕРПОЛ у Нью-Делі (жовтень 2022 року). Саме тоді відбулася сенсаційна презентація першого у світі Метавсесвіту, спеціально призначеного для правоохоронних органів [4]. Учасники Асамблеї через VR-гарнітури здійснили аватарний тур штаб-квартирою ІНТЕРПОЛ у Ліоні, пройшли навчання з верифікації документів у VR-класі та відпрацювали прикордонний контроль у симульованому аеропорту. Платформу розміщено на INTERPOL Secure Cloud — власній хмарній інфраструктурі організації, що гарантує нейтральність і захист від стороннього втручання.

Генеральний секретар ІНТЕРПОЛ Юрген Шток підкреслив: *«Метавсесвіт — це не абстрактне майбутнє. Він вже тут. Питання, які він ставить, — це питання, що завжди мотивували ІНТЕРПОЛ: підтримка держав-членів у боротьбі зі злочинністю та забезпечення безпеки світу — віртуального чи ні»*.

Водночас з метавсесвітом ІНТЕРПОЛ створив Metaverse Expert Group (MEG) — міжвідомчий орган, що об'єднав представників правоохоронних органів, урядів, приватного сектора, академічних установ та міжнародних організацій. Завдання MEG — розробка рекомендацій щодо протидії метазлочинам і забезпечення принципу «secure by design» у розвитку метавсесвіту [4].

У вересні 2023 року ІНТЕРПОЛ розширив застосування Метавсесвіту, запустивши першу XR-програму навчання боротьби з порушеннями інтелектуальної власності — у рамках International IP Crime Investigators College (IPRIC). Слідчі відпрацьовують перевірку складів із контрафактом, вивчають ознаки фальсифікованих документів та здійснюють криміналістичний аналіз у повністю зануреному цифровому середовищі, без ризиків, пов'язаних із реальним обігом небезпечних товарів [7].

Ключовим теоретичним документом стала White Paper ІНТЕРПОЛ «Metaverse: A Law Enforcement Perspective», презентована у Давосі у рамках World Economic Forum (18 січня 2024 року). Документ запровадив концепцію «metacrimes» — метазлочинів: злочинних діянь, вчинюваних у метавсесвіті або за його посередництва, для яких чинні правові системи не мають адекватних норм реагування. Генеральний секретар Шток констатував: «Ми спостерігаємо, як Метавсесвіт і штучний інтелект створюють нові можливості для злочинної діяльності, до яких світ не готовий повною мірою» [5].

Ключовою ідеєю документа є також правоохоронний потенціал Метавсесвіту — не лише як середовища злочинності, але і як інструменту її розкриття та профілактики: можливість VR-реконструкції місць злочину для криміналістичного аналізу та судового представлення доказів; навчання першого ешелону реагування у симульованих середовищах; збереження та аналіз цифрових доказів; взаємодія з громадянами через цифрові поліцейські пункти у метавсесвіті. Слідчі ІНТЕРПОЛ прямо зазначають: можна уявити прокурора, який проводить суд присяжних через VR-реконструкцію місця злочину у Метавсесвіті — і це вже не наукова фантастика, а найближча технологічна перспектива.

У березні 2024 року ІНТЕРПОЛ, спільно з ЮНІСРІ, оприлюднив оновлений Toolkit відповідального застосування ШІ у правоохоронній практиці. Документ закріпив принципи пояснюваного штучного інтелекту (ХАІ), підзвітності алгоритмів та обов'язкового людського контролю над ШІ-рішеннями у поліцейській і слідчій практиці [6].

Хронологію та зміст ключових ініціатив ІНТЕРПОЛ у сфері Метавсесвіту та ШІ систематизовано у Табл. 1.

Таблиця 1

Хронологія ключових ініціатив ІНТЕРПОЛ у сфері метавсесвіту та ШІ (2022–2024)

Подія / документ	Дата	Ключовий зміст	Правове значення
INTERPOL Metaverse — перший поліцейський Метавсесвіт	20 жовтня 2022 р., 90-та Генеральна Асамблея ІНТЕРПОЛ, Нью-Делі	Запуск першого у світі Метавсесвіту для правоохоронних органів; 3D-тур штаб-квартирою в Ліоні через аватари; навчання верифікації документів та прикордонного контролю у VR-класі	Перший прецедент інституціалізації Метавсесвіту як правоохоронного інструменту; статус INTERPOL Secure Cloud забезпечує нейтральність платформи
Створення Metaverse Expert Group (MEG)	Жовтень 2022 р.	Міжнародна група експертів ІНТЕРПОЛ з Метавсесвіту: представники поліції, урядів, приватного сектора, академії та НУО; завдання — вироблення рекомендацій щодо протидії метазлочинам	Основа для міжнародного кримінально-правового регулювання Метавсесвіту; перший крок до формування lex metaverse

Подія / документ	Дата	Ключовий зміст	Правове значення
White Paper «Metaverse: A Law Enforcement Perspective»	18 січня 2024 р., WEF, Давос	Типологія метазлочинів (metacrimes); рекомендації щодо криміналістики, збору доказів з VR-гарнітур та серверів; інструментарій для реконструкції місць злочину у 3D	Доктринальна основа для криміналізації метазлочинів; підтверджує необхідність спеціалізованого кримінального законодавства (концепція Цифрового КК)
Імерсивне навчання для боротьби з ІР-злочинами	Вересень 2023 р.	Перша XR-програма навчання інспекторів інтелектуальної власності: VR-склади, місця злочину, криміналістичні лабораторії у Метавесвіті	Прецедент застосування Метавесвіту для спеціалізованої правоохоронної підготовки; формує стандарт immersive policing training
Оновлений INTERPOL AI Toolkit для правоохоронців	Березень 2024 р.	Ревізований інструментарій відповідального застосування ШІ в правоохоронній практиці; спільно з ЮНІСРІ; принципи ХАІ та підзвітності алгоритмів	Міжнародний стандарт ethi-compliant AI у поліцейській практиці; основа для імплементації рівня VCBR Цифрового Кодексу

Наведена хронологія засвідчує системний, поетапний характер діяльності ІНТЕРПОЛ у цій сфері — від технологічної демонстрації (2022) до доктринальної кодифікації метазлочинів (2024). Ця траєкторія підтверджує достовірність і своєчасність авторських концепцій, що розвивалися паралельно у вітчизняній правовій науці.

3. ТИПОЛОГІЯ МЕТАЗЛОЧИНІВ: СПІВВІДНОШЕННЯ З ЦИФРОВИМ КОДЕКСОМ

White Paper ІНТЕРПОЛ (2024) встановила типологію метазлочинів, що включає: грумінг неповнолітніх, радикалізацію та вербування через XR-контент, кібер-фізичні атаки на критичну інфраструктуру, крадіжку 3D-власності та цифрових активів, самовільне проникнення у приватні VR-простори, пограбування аватарів, маніпуляції з NFT і шахрайство з цифровими активами [5]. Аналіз цієї типології у співвідношенні з авторським проектом Цифрового кримінального кодексу виявляє значний рівень концептуальної конвергенції — при тому, що обидва документи розроблялися незалежно (Табл. 2).

Таблиця 2

Типологія метазлочинів ІНТЕРПОЛ (2024) та відповідні склади Цифрового Кодексу

Тип метазлочину	Характеристика	Наявність у КК України	Цифровий КК (Костенко)
Grooming у метавесвіті	Цілеспрямований цифровий контакт педофілів із неповнолітніми через аватари та XR-середовища	Частково (ст. 156 КК — розбещення)	Окремий склад: «сексуальне грумінг-переслідування у цифровому середовищі»
Радикалізація та вербування в екстремістські організації через VR	Використання імерсивного XR-контенту для психологічної обробки та рекрутування	Відсутній спеціальний склад	«Кібертероризм у метавесвіті»; посилена відповідальність за XR-вербування

Тип метазлочину	Характеристика	Наявність у КК України	Цифровий КК (Костенко)
Кібер-фізичні атаки на критичну інфраструктуру	Атаки з метавсесвіту на фізичні об'єкти через цифрові двійники промислових систем (SCADA)	Ст. 361 КК (часткове покриття)	«Кібер-фізичний напад»: окремий склад з підвищеною санкцією
Пограбування аватара / theft of virtual assets	Незаконне заволодіння цифровими активами, NFT, ігровими предметами через злам або шахрайство у XR	Відсутній; ст. 190 КК незастосовна через відсутність «майна»	«Цифрове пограбування»: розкрадання цифрових нематеріальних активів
Самовільне проникнення у приватний VR-простір	Несанкціонований доступ до приватних XR-просторів («trespassing in virtual spaces»)	Відсутній (ст. 162 КК — лише фізичне житло)	«Цифрове вторгнення»: порушення недоторканності цифрового простору особи
ШІ-deepfake із злочинною метою	Генерація реалістичного аудіо/відео/XR-контенту для шахрайства, дискредитації, маніпуляції правосуддям	Ст. 190, 351 ¹ КК — часткове покриття	«ШІ-шахрайство»: окремий склад із кваліфікуючою ознакою застосування GenAI

Порівняльний аналіз свідчить, що всі категорії метазлочинів, ідентифіковані ІНТЕРПОЛ, мають відповідники у авторській концепції Цифрового Кримінального Кодексу [1; 2]. Водночас авторська концепція є більш деталізованою у частині нейровтручання і ШІ-шахрайства з елементами GenAI, тоді як документ ІНТЕРПОЛ приділяє більшу увагу процедурним аспектам криміналістики (збір доказів із VR-гарнітур, серверів, блокчейн-аналітика) [5]. Це дозволяє стверджувати про взаємодоповнюваність двох підходів і доцільність їх синтезу при формуванні національного законодавства.

4. ШІ У ПРАВОСУДДІ ТА ЮРИДИЧНІЙ ПСИХОЛОГІЇ: СУЧАСНИЙ СТАН

Застосування ШІ у правосудді набуло глобального масштабу. Системи предиктивного правосуддя функціонують у десятках юрисдикцій: COMPAS (США) оцінює ризик рецидиву та застосовується при визначенні запобіжних заходів і умов дострокового звільнення; Prometea (Аргентина) автоматизує аналіз судових справ; ROSS Intelligence (США) здійснює семантичний пошук правових прецедентів; eJustice-AI (ЄС, пілотні програми) опрацьовує адміністративні справи. Суд штату Вісконсин у справі *State v. Loomis* (2016) став першим прецедентом оскарження вироку, постановленого з урахуванням ШІ-оцінки COMPAS — і визнав її конституційною, що викликало широку правову дискусію [13].

Ключовою проблемою є феномен «алгоритмічної упередженості» (algorithmic bias): дослідження ProPublica (2016) встановило, що COMPAS вдвічі частіше помилково класифікує темношкірих підсудних як «високоризикових» порівняно з білошкірими [10]. Це є прямим порушенням принципу рівності перед законом та права на справедливий суд, гарантованих ст. 6 ЄКПЛ. Дослідження ScienceDirect (2025) в галузі метавсесвітнього полісингу підтверджує, що аналогічні ризики будуть притаманні і ШІ-системам правоохоронної діяльності у XR-середовищах [11].

У галузі юридичної психології ШІ трансформує три ключові практики. По-перше, криміналістичний профайлінг: системи мультимодального аналізу (рух очей, голос, ЕЕГ-патерни, мікровирази обличчя) дозволяють оцінювати психоемоційний стан особи під час допиту або судового засідання без її відома. По-друге, ризик-оцінка рецидивізму: ШІ-системи

прогнозують поведінку засудженого на основі психологічних тестів, соціально-демографічних даних та поведінкової аналітики. По-третє, ІІІ-опосередкована психологічна реабілітація: особливо актуальна для ветеранів та осіб, що постраждали від наслідків збройного конфлікту, — але потребує ретельного правового регулювання в частині конфіденційності та відповідальності за шкоду.

Нейроправовий вимір цих трансформацій полягає у виникненні нового покоління прав людини — прав на когнітивну свободу, ментальну приватність, психологічну цілісність і психологічну тяглість (continuity). Ці права, обґрунтовані у рамках Neurorights Framework Колумбійського університету та частково закріплені у Конституції Чилі (2021), є прямою відповіддю на можливість технологічного втручання у ментальне середовище особи через XR-пристрої та ІІІ-платформи [12].

5. ЦИФРОВИЙ КОДЕКС ЯК АРХІТЕКТУРА РЕГУЛЮВАННЯ: ЗАСТОСУВАННЯ У ПРАВООХОРОННІЙ СФЕРІ

Концепція Цифрового Кодексу (Digital Code), розроблена автором у монографії «Metaverse та WEB 4.0» [1] та в дослідженні [2], є уніфікованою чотирирівневою нормативною платформою, що систематизує правовідносини у цифровому просторі. Кожен рівень платформи має специфічне значення для правоохоронної, кримінально-правової та психологічно-правової сфер (Табл. 3).

Таблиця 3

Архітектура Цифрового Кодексу (IVC/NMR/DLRE/VCBR) у правоохоронному, кримінально-правовому та психологічному контексті

Модуль	Предмет регулювання	Застосування у кримінальному праві та правосудді	Застосування у психології та правоохоронній практиці
IVC (ідентифікація та верифікація суб'єктності)	Правовий статус цифрової особи, аватара, ІІІ-агента; реєстр цифрових ідентичностей	Суб'єкт злочину у метавсесвіті; кримінальна відповідальність оператора ІІІ; «цифровий кіднепінг»	Верифікація особи у цифрових психологічних консультаціях; ідентифікація в ІІІ-профайлінгу
NMR (нормативно-метавсесвітні відносини)	Правовідносини у XR-просторі; lex protocolli; смарт-контрактне право; GLM	Юрисдикція над метазлочинами; транскордонне кримінальне переслідування; Цифровий КК	Правові рамки ІІІ-профайлінгу; нейроправовий захист у XR-сесіях
DLRE (цифрово-правове регулювання середовища)	Алгоритмічне управління; платформна відповідальність; аудит ІІІ-рішень	Відповідальність за предиктивне правосуддя; стандарти ХАІ у слідчій практиці; INTERPOL AI Toolkit	Регулювання ІІІ-психодіагностики; відповідальність за хибний психологічний діагноз ІІІ
VCBR (верифікація, відповідальність, права)	Цифровий due process; право на пояснення алгоритму; цифровий суд; омбудсмен	Апеляція на рішення COMPAS/предиктивних систем; захист від алгоритмічної дискримінації	Право на когнітивну свободу; захист нейроданих; INTERPOL immersive training standards

Особливого значення в контексті взаємодії з ініціативами ІНТЕРПОЛ набуває рівень NMR — нормативно-метавсесвітніх відносин, який включає механізм lex protocolli (право протоколу платформи як самостійне джерело цифрового права). Саме цей механізм може

забезпечити правову силу процедурним стандартам ІНТЕРПОЛ щодо збору і збереження цифрових доказів та правил поведінки в правоохоронному метавсесвіті на рівні, наближеному до обов'язкової норми для держав-членів організації.

Концепція Grand Lex Metaverse (GLM) — транснаціонального цифрового права як *lex specialis* для вирішення юрисдикційних конфліктів — прямо кореспондує з висновком White Paper ІНТЕРПОЛ (2024) щодо необхідності «багатостороннього підходу із залученням множинних юрисдикцій, вимірів та організацій» у протидії метазлочинам [5]. GLM може стати доктринальною основою такого підходу, подолавши роздробленість національних юрисдикцій без шкоди для принципу державного суверенітету [1].

6. ПРОБЛЕМИ ВПРОВАДЖЕННЯ В УКРАЇНІ ТА РЕКОМЕНДАЦІЇ

Україна перебуває у специфічній геополітичній і технологічній ситуації, що одночасно ускладнює і стимулює впровадження ІІІ у правоохоронну та юридично-психологічну сфери. Воєнний стан, масштабна цифрова агресія (понад 4 тисячі кібератак лише на критичну інфраструктуру за 2023–2024 роки за даними CERT-UA), необхідність документування воєнних злочинів у режимі реального часу та гостра потреба у психологічній реабілітації мільйонів ветеранів і цивільних — усі ці фактори роблять питання правового регулювання ІІІ не академічним, а невідкладно практичним.

Проблема 1. Юрисдикційний вакуум у кіберпросторі. Чинні КК і КПК України не містять норм, що регулювали б злочини, вчинені через Метавсесвіт або за допомогою цифрових аватарів. Будапештська конвенція про кіберзлочинність (2001), до якої приєдналась Україна, охоплює лише традиційні ІКТ-злочини і не поширюється на метавсесвіт як середовище злочинної діяльності.

Проблема 2. Відсутність правових стандартів ІІІ у правосудді. Проект Закону України «Про штучний інтелект» не містить спеціальних норм щодо застосування ІІІ у кримінальному судочинстві, предиктивній поліцейській діяльності або психологічній судовій експертизі. Це створює ризик безсистемного впровадження ІІІ-інструментів без необхідних гарантій права на захист (пор. підхід ЄС: EU AI Act [8] та стратегія Web 4.0 [9]).

Проблема 3. Відсутність захисту нейроданих. Чинний Закон України «Про захист персональних даних» не виокремлює нейрофізіологічні дані як окрему категорію, що потребує посиленого захисту. Водночас CERT-UA фіксує випадки витоку медичних і психологічних даних при атаках на заклади охорони здоров'я.

Проблема 4. Відсутність процесуального статусу VR-доказів. КПК України не містить положень щодо допустимості доказів, отриманих із VR-гарнітур, метаплатформ або генерованих ІІІ-системами. Це унеможливує ефективне кримінальне переслідування за метазлочини навіть за наявності технічних доказів.

З урахуванням окреслених проблем пропонуються такі рекомендації.

Рекомендація 1. Розробити та прийняти Концепцію Цифрового Кодексу України як першого кроку до системного регулювання цифрових правовідносин, включно із спеціальним розділом щодо застосування ІІІ у правосудді та правоохоронній діяльності.

Рекомендація 2. Ратифікувати Другий Додатковий протокол до Будапештської конвенції про кіберзлочинність (2022) та ініціювати розширення предметної сфери конвенції для охоплення метазлочинів.

Рекомендація 3. Розпочати розробку процесуального регламенту щодо допустимості та оцінки доказів, отриманих із метавсесвіту та XR-пристроїв, у рамках оновлення КПК України.

Рекомендація 4. Закріпити нейродані як окрему категорію персональних даних підвищеного захисту у Законі «Про захист персональних даних» та Законі «Про штучний інтелект».

Рекомендація 5. На базі Наукової лабораторії імерсивних технологій і права ДНУ ШБП НАПрН України та Академії МВС створити Центр компетенцій з правоохоронного Метавсесвіту, орієнтований на підготовку слідчих, прокурорів і суддів до роботи з доказами метазлочинів і ІІІ-рішень.

7. ВИСНОВКИ

Проведене дослідження дозволяє сформулювати такі узагальнені висновки.

1. Стрімкий розвиток III та XR-технологій формує принципово новий правовий ландшафт, в якому традиційні категорії кримінального права, юрисдикції та доказування потребують кардинального переосмислення.
2. Ініціативи ІНТЕРПОЛ — правоохоронний метавсесвіт (2022), Metaverse Expert Group (2022), XR-навчання (2023) і White Paper з метазлочинів (2024) — є незаперечним міжнародним підтвердженням актуальності авторської концепції цифрової юрисдикції та Цифрового Кодексу.
3. Концепція Цифрового Кодексу (IVC/NMR/DLRE/VCBR) є найбільш комплексним вітчизняним доктринальним інструментом для правового регулювання III-опосередкованих відносин у праві, психології та правоохоронній діяльності.
4. Типологія метазлочинів ІНТЕРПОЛ демонструє значний рівень конвергенції з авторським проектом Цифрового КК, що свідчить про методологічну коректність останнього та відкриває можливості для їх синтезу при розробці вітчизняного законодавства.
5. Україна має унікальний шанс — на підставі власного кризового досвіду цифрових воєнних злочинів, кібератак та масштабної необхідності в III-підтримці правоохоронної та психологічної практики — сформувати передове законодавство у цій сфері, що може стати орієнтиром для регіону Центральної та Східної Європи.

Список використаних джерел

1. Костенко О. В. Metaverse та WEB 4.0 : монографія. Київ : Фенікс, 2026. 380 с.
2. Костенко О. В. Цифрова юрисдикція: модель. *Інформація і право* . 2025. № 4(55). DOI: [https://doi.org/10.37750/2616-6798.2025.4\(55\).346297](https://doi.org/10.37750/2616-6798.2025.4(55).346297) .
3. Костенко О. В. Електронна юрисдикція, метавсесвіт, штучний інтелект, цифрова особистість, цифровий аватар, нейронні мережі: теорія, практика, перспективи. *Наукові інновації та передові технології* . 2022. № 2(4). URL: <https://www.researchgate.net/publication/358486613> (дата звернення: 08.06.2026).
4. INTERPOL launches first global police Metaverse : офіц. повідомл. INTERPOL. 20 жовт. 2022 р. URL: <https://www.interpol.int/en/News-and-Events/News/2022/INTERPOL-launches-first-global-police-Metaverse> (дата звернення: 08.06.2026).
5. Metaverse: A Law Enforcement Perspective. Use Cases, Crime, Forensics, Investigation, and Governance: White Paper / INTERPOL Metaverse Expert Group. Davos: INTERPOL, Jan. 2024. 27 р. URL: <https://www.interpol.int/en/content/download/20828/file/Metaverse%20-%20a%20law%20enforcement%20perspective.pdf> (дата звернення: 08.06.2026).
6. Revised toolkit empowers law enforcement with responsible AI practices : офіц. повідомл. INTERPOL. 6 берез. 2024 р. URL: <https://www.interpol.int/en/News-and-Events/News/2024/Revised-toolkit-empowers-law-enforcement-with-responsible-AI-practices> (дата звернення: 08.06.2026).
7. Capitalizing on immersive learning to fight IP crime : офіц. повідомл. INTERPOL. Верес. 2023 р. URL: <https://www.interpol.int/en/News-and-Events/News/2023/Capitalizing-on-immersive-learning-to-fight-IP-crime> (дата звернення: 08.06.2026).
8. Регламент (ЄС) 2024/1689 Європейського Парламенту та Ради від 13 червня 2024 року, що встановлює гармонізовані правила щодо штучного інтелекту (Закон про штучний інтелект). *Офіційний журнал Європейського Союзу* . 2024. L 1689. URL: https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_2024_1689 (дата звернення: 08.06.2026).
9. Веб 4.0 та віртуальні світи: лідерство у наступному технологічному переході: Повідомлення Комісії до Європейського Парламенту, Ради, Європейського економічного та соціального комітету та Комітету регіонів. COM(2023) 442 final. Брюссель: Європейська Комісія, 2023. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=COM:2023:442:FIN> (дата звернення: 08.06.2026).

10. Angwin J., Larson J., Mattu S., Kirchner L. Machine Bias. *ProPublica* . 2016. 23 May. URL: <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing> (дата звернення: 08.06.2026).
11. Ієнка М., Андорно Р. На шляху до нових прав людини в епоху нейротехнологій: когнітивна свобода, ментальна конфіденційність, ментальна цілісність та психологічна безперервність. *Науки про життя, суспільство та політика* . 2017. Том 13, № 5. DOI: <https://doi.org/10.1186/s40504-017-0050-1> .
12. Haber E. The Criminal Metaverse. *Indiana Law Journal*. 2024. Vol. 99, iss. 3. P. 843–891. URL: indiana.edu (дата звернення: 08.06.2026).
13. State v. Loomis. 371 Wis. 2d 235, 881 N.W.2d 749 (Wis. 2016). URL: <https://www.wicourts.gov/sc/opinion/DisplayDocument.pdf?content=pdf&seqNo=171690> (дата звернення: 08.06.2026).

Tim D. WILLIAMS
BSc, MSc, MBA, PGCHE
FBCS CITP, MCIIS, MIET, GMBPsS
Independent researcher
ORCID: 0000-0002-1788-5736

WHY AI LEGAL ASSISTANTS HALLUCINATE IN NEW FIELDS: A MATHEMATICAL LIMIT, WITH POST-QUANTUM CRYPTOGRAPHY (PQC) AS EXAMPLE

1. Introduction and motivation

Some AI errors in legal practice cannot be removed by better engineering. They follow from a mathematical limit. Practitioners need to know which legal questions sit inside that limit.

AI assistants are now part of everyday legal work. Lawyers use them for research, drafting and first-pass review. When these assistants make confident mistakes, as in the US case *Mata v. Avianca* [1], the usual response is to call for better training data, better prompts and better engineering. This paper offers a different perspective. One specific kind of AI error cannot be removed by engineering. It follows from a mathematical proof. The proof applies to any AI assistant that gives confidence-marked answers, both now and in the future. The professional task is therefore not to wait for AI to become reliable. The task is to learn which legal questions fall inside the mathematical limit, and to avoid reliance on AI outside its known limits of reliability. Post-quantum cryptography (PQC) is used here as a worked example. PQC is a new legal field. The public record is thin. The stakes are high.

2. The Kalai–Vempala lower bound

A 2024 mathematical result shows that any properly calibrated AI assistant must make a minimum rate of errors. The minimum rate depends on how often facts appear in the training data.

In 2024, Kalai and Vempala proved a result about AI language models [2]. The result is a lower bound, which means a minimum that cannot be reduced. This bound applies to any AI model that is well-calibrated, meaning a model whose stated confidence matches how often it is correct. This bound says: the model must make errors at a rate at least equal to the share of facts that appear exactly once in its training data. Xu, Jain and Kankanhalli then strengthened the result [3]. They showed that some errors are unavoidable as a matter of computability, not training effort.

The practical meaning is simple. An AI assistant is not equally reliable across all domains of knowledge. The model is reliable where the same fact appears many times, in many sources. The model is unreliable where the authoritative source appears only once. This is not a problem of current models. It is a proven limit on every future calibrated model. The professional task is to identify which legal domains fall on the unreliable side of the limit.

3. A second factor: gaps in the legal training record

The mathematical limit is an underestimate where AI assistant training data is not only limited but also distorted. Legal practice has structural reasons for such distortion.

The Kalai–Vempala limit assumes that thin training data is the only problem. In legal practice, the data is also distorted. Anderson showed that banks for many years suppressed reports of cryptographic failure [4]. The public record was therefore biased towards success and away from instructive failure. Legal practice has the same structural feature. Sealed settlements, confidential dispositions, protective orders and legal professional privilege all remove the most instructive failure cases from the public record [5]. These are exactly the cases an AI assistant would need in order to learn. The combined effect is that AI legal assistants are under-exposed to the newest and most relevant subject matter and over-exposed to favourable outcomes. Anderson’s earlier economic account of cryptographic failure [6] and this author’s earlier work on epistemic grounding in cybersecurity [7] both support the same conclusion. Confident error is the predictable result. The mechanism and the economic conditions in which it operates are analysed by this author in a forthcoming paper [8].

4. Why PQC-legal advisory is the canonical case

PQC-legal advisory is the clearest near-term field where the Kalai–Vempala limit applies most acutely. PQC combines all three features that trigger the limit.

Post-quantum cryptography has moved from a research topic to a live legal field. In August 2024 the US National Institute of Standards and Technology approved the first three PQC standards [9]. This opened a regulatory and advisory market that cuts across financial reporting [10], operational resilience [11], and emerging negligence claims [12]. Lawyers in many roles, including in-house counsel, advocates, notaries and judges, now use AI assistants to read this material at speed. The problem is that PQC-legal advisory sits exactly where the Kalai–Vempala limit predicts trouble. No organisation is known to have finished a full PQC transition. No comprehensive public cost data exists [13]. The engineering solutions continue to change. Three features converge: (1) the subject is recent; (2) the boundary between technical and legal expertise is unsettled; and (3) the legal questions are jurisdiction-specific. Each feature is hard on its own. The three together compound. PQC is perhaps the first such field on which the legal profession will be tested. Other novel fields with similar legal/AI features include AI liability, quantum sensing, synthetic biology and autonomous-systems law.

5. A proposed empirical test

This paper’s theoretical argument supports a small empirical test that any research team can run as follows.

A set of 30 to 40 prompts on PQC-legal questions is drawn from NIST guidance, IAS 36 application material, DORA and NIS2 provisions, and recent sanctions decisions. A matched set of prompts of similar difficulty is drawn from settled IT-legal subject matter. Both sets are run against three frontier AI models under the same conditions. Two raters then grade the responses against authoritative sources, blinded to condition where this is possible. The primary expectation is that the PQC-legal prompts produce a much higher rate of confident error.

Five categories of error are scored: fabricated authority, misattributed holding, incorrect standard reference, confidently wrong timeline, and unsupported doctrinal inference. The test is diagnostic, not decisive. It cannot disprove the Kalai–Vempala result. It can only show the predicted pattern in the PQC-legal condition. The paper’s main contribution is the structural explanation. The measured size of the effect is secondary. This pilot is intended to support a later journal extension.

6. What this means for legal practice

If AI error in legal advisory is mathematically structured rather than uniform, the professional response should change in specific ways.

The duty on counsel is not to eliminate a general defect. The duty is to identify the high-risk scenarios where legal AI assistants will confidently generate incorrect results. From the earlier sections, three signals of high-risk scenarios support a practical framework: recent subject emergence,

jurisdiction-specific novelty, and technical–legal cross-domain character. Where two or more signals are present, counsel should presume an elevated risk of AI error. Disclosure and verification duties should follow that presumption in proportion.

Three practical implications follow. For in-house counsel and notaries, structured verification protocols should apply to any AI output in the flagged scenarios, with specific checks for fabricated standards and misattributed cases. For judges, protocols for AI-assisted filings should treat the medium as neither uniformly reliable nor uniformly suspect. For bar associations and regulators, domain-specific guidance can be issued now, without waiting for a general solution; the structural account supplies the criteria for naming the domains. A final note is required on retrieval-augmented generation (RAG) [14]. RAG shifts the sparsity problem rather than solving it. Where the retrievable authoritative material is itself sparse, RAG offers no relief [15].

7. Conclusion

AI error in legal practice is partly structural and partly an engineering problem. The structural part can be named and managed. AI error in legal assistants is not uniform across legal fields. One specific class of error is mathematically structural, and the distortion of the legal training corpus makes it worse. PQC-legal advisory is the clearest near-term example. The diagnosis generalises to every emerging-technology field in which counsel is asked for early advice. The professional task is not to wait for AI reliability to improve. The task is to recognise the regime in which reliability will remain bounded by mathematics, and to respond with proportionate verification and disclosure.

Conflict of interest disclosure

The author is engaged in commercial PQC advisory activity within the financial-services sector, and is co-founder of ProteQC, a firm in which Darren Bender (cited at reference 12) serves as Chief Legal Officer. No commercial material informs this submission. All argumentation is based on publicly available peer-reviewed sources, public regulatory texts and public court records.

References (DSTU 8302:2015)

1. Mata v. Avianca, Inc., No. 22-cv-1461 (PKC), United States District Court, Southern District of New York. Order imposing sanctions, 22 June 2023. URL: https://storage.courtlistener.com/recap/gov.uscourts.nysd.575368/gov.uscourts.nysd.575368.54.0_7.pdf (accessed: 22.04.2026).
2. Kalai A. T., Vempala S. S. Calibrated language models must hallucinate // Proceedings of the 56th Annual ACM Symposium on Theory of Computing (STOC 2024). New York : ACM, 2024. P. 160–171.
3. Xu Z., Jain S., Kankanhalli M. Hallucination is inevitable: an innate limitation of large language models. 2024. arXiv:2401.11817. URL: <https://arxiv.org/abs/2401.11817> (accessed: 22.04.2026).
4. Anderson R. Why cryptosystems fail // Proceedings of the 1st ACM Conference on Computer and Communications Security. New York : ACM, 1993. P. 215–227.
5. Galanter M. Why the “haves” come out ahead: speculations on the limits of legal change // Law & Society Review. 1974. Vol. 9, No. 1. P. 95–160. DOI: 10.2307/3053023.
6. Anderson R. Why information security is hard — an economic perspective // 17th Annual Computer Security Applications Conference (ACSAC 2001). New Orleans : IEEE, 2001. P. 358–365.
7. Williams T. D. Epistemological questions for cybersecurity // 2020 International Conference on Cyber Security and Protection of Digital Services (Cyber Security). Dublin : IEEE, 2020. P. 1–4. DOI: 10.1109/CyberSecurity49315.2020.9138884.
8. Williams T. D. Cybersecurity economics, PQC and AI — 25 years post Anderson (2001) // Proceedings of Cyber Science 2026. London : Springer LNCS, 2026. [in press].
9. National Institute of Standards and Technology. NIST releases first 3 finalized post-quantum encryption standards. Gaithersburg : NIST, 13 August 2024. URL: <https://www.nist.gov/news-events/news/2024/08/nist-releases-first-3-finalized-post-quantum-encryption-standards> (accessed: 22.04.2026).

10. IAS 36 Impairment of Assets. London : International Financial Reporting Standards Foundation, as amended. URL: <https://www.ifrs.org/issued-standards/list-of-standards/ias-36-impairment-of-assets/> (accessed: 22.04.2026).
11. Regulation (EU) 2022/2554 of the European Parliament and of the Council of 14 December 2022 on digital operational resilience for the financial sector (DORA) // Official Journal of the European Union. 2022. L 333/1–L 333/79.
12. Bender D. Post-Quantum Negligence: applying the Learned Hand formula to organizational liability for quantum computing threats. SSRN Working Paper No. 6413418, March 2026. DOI: 10.2139/ssrn.6413418. URL: <https://ssrn.com/abstract=6413418> (accessed: 21.05.2026).
13. Williams T. D. Towards a framework for post-quantum cryptography cost estimation. SSRN Working Paper No. 6726106, May 2026. DOI: 10.2139/ssrn.6726106. URL: <https://ssrn.com/abstract=6726106> (accessed: 21.05.2026).
14. Lewis P., Perez E., Piktus A. et al. Retrieval-augmented generation for knowledge-intensive NLP tasks // Advances in Neural Information Processing Systems 33 (NeurIPS 2020). Red Hook, NY : Curran Associates, 2020. P. 9459–9474. arXiv:2005.11401.
15. Kandpal N., Deng H., Roberts A., Wallace E., Raffel C. Large language models struggle to learn long-tail knowledge // Proceedings of the 40th International Conference on Machine Learning (ICML 2023). PMLR, 2023. Vol. 202. P. 15696–15707.

Ірина ГЛОВІЮК, професорка кафедри кримінального процесу та криміналістики, Одеський державний університет внутрішніх справ, доктор юридичних наук, професор

ІІІ ТА ЗАКОН УКРАЇНИ «ПРО АКАДЕМІЧНУ ДОБРОЧЕСНІСТЬ»: ПОДАЛЬШІ ПИТАННЯ

Одним з новітніх викликів у аспекті академічної доброчесності є використання ІІІ у науковій та освітній діяльності. За ст. 29 Закону України «Про академічну доброчесність», недоброчесне використання результатів, згенерованих штучним інтелектом, - це оприлюднення текстів, зображень, моделей, даних, інших результатів, згенерованих штучним інтелектом, як результатів власної академічної діяльності, якщо цей факт не зазначено в академічному творі чи супровідних матеріалах до нього. У ч. 4 ст. 8 цього ж Закону прописано, що у разі використання в академічному творі об'єкта, згенерованого штучним інтелектом, автор повинен повідомити про це в такому творі із зазначенням методики генерування та/або посилання на відповідну комп'ютерну програму чи її опис згідно з вимогами щодо оформлення та/або оприлюднення відповідних академічних творів, визначених відповідним суб'єктом академічної діяльності.

Ці норми є більш прогресивними, аніж ті, що попередньо пропонувалися у проєкті цього Закону. Адже у ст. 8 проєкту вказувалося, що особа не може вважатися автором академічного твору (частини академічного твору), якщо він сформований (згенерований) за запитом особи комп'ютерною програмою в автоматичному режимі. Додаткові питання виникають, наприклад, якщо ІІІ допоміг оформити перелік джерел (наприклад, переформатував у потрібний стиль цитування): частина тексту згенерована ІІІ, отже, якщо формально підходити до цього питання, то це порушення академічної доброчесності. А якщо ІІІ лише сформулював тему? А якщо лише план? А якщо вичитав і вказав, як краще переформулювати 2-3 вже написані людиною фрази? Відповідей немає. Утім, є положення: «При використанні в академічному творі частин, сформованих (згенерованих) комп'ютерними програмами, цей

факт має бути зазначений автором (авторами) із зазначенням методики формування (генерування) або посиланням на відповідну комп'ютерну програму чи її опис» (ч. 6 ст. 8 проєкту). Як видається, воно, по-перше, мало синхронізується з попереднім абзацом цієї ж частини, який цитувався вище. По-друге, цікаве формулювання про частини твору. Тобто якщо 90 % виконане ШІ, а 10% людиною – хто буде автором? В проєкті немає відповіді і на це питання [1, с. 153]. Були і інші зауваження саме до проєкту Закону [2].

Утім, наявне формулювання теж викликає деякі проблемні питання як практичного, так і суто світоглядного характеру.

1. Закон України «Про академічну доброчесність» не уточнює, що має розумітися під формулюванням «об'єкт», й саме по собі формулювання «використання в академічному творі об'єкта» є доволі невизначеним. Адже під ним можна розуміти і певні графічні зображення, діаграми, схеми, інфографіку тощо. Якщо об'єкт розуміти так, то виходить, що використовувати ШІ для генерації тексту заборонено, але, судячи з декларативного, а не суто карального, підходу взаємодії з ШІ, є сумніви, що законодавець мав на увазі саме це; тим більше, що застосовано формулювання «інших результатів, згенерованих штучним інтелектом» у ст. 29 Закону.

2. Зважаючи на відсутність уточнень у Законі України «Про академічну доброчесність», виходить, що саме заклади вищої освіти, наукові установи та наукові видання мають не лише процитувати ст. 29 Закону у внутрішніх політиках, а й уточнити, що мається на увазі під використанням «в академічному творі об'єкта, згенерованого штучним інтелектом». Нагадаємо, що академічний твір – це твір у сфері освіти чи науки, виражений у формі дисертації, творчого мистецького проєкту, кваліфікаційної роботи, курсової роботи, наукової монографії, наукової статті, підручника, навчального посібника або в іншій формі, визначеній законодавством та/або внутрішнім актом (ст. 1 Закону України «Про академічну доброчесність»).

3. Закон України «Про академічну доброчесність» не визначає, та й не має цього робити, межі використання ШІ в академічному творі. Отже, це покладається на внутрішні політики, які мають бути розлогими у цьому питанні та, очікувано, можуть бути різними у контексті допустимих меж використання ШІ. Це може за певних обставин призвести до різного ступеня довіри до академічних творів залежно від внутрішньої політики щодо використання ШІ. А це, як видається, є новим викликом для освіти та науки.

4. У Living guidelines on the responsible use of generative AI in research ERA Forum Stakeholders' document, European Commission, Directorate-General for Research and Innovation, Directorate E-Prosperity, Unit E4 - AI in Science & Critical Technologies, May 2026, прописано: «З метою забезпечення прозорості дослідники детально описують, які інструменти генеративного ШІ були значною мірою використані в їхніх дослідницьких процесах. У цьому випадку, наприклад, використання генеративного ШІ як базового інструменту редакційної підтримки не є істотним використанням. Однак використання його як методу або джерела, інтерпретація аналізу даних, проведення огляду літератури, виявлення прогалин у дослідженнях, формулювання цілей дослідження, розробка гіпотез тощо можуть мати істотний вплив. Коли генеративний ШІ суттєво впливає на результати, дослідники прозоро зазначають його використання відповідно до рекомендацій свого журналу або стандартів у своїй дисципліні в розділі «Методи» (або еквівалентному), відповідально оцінюючи ступінь його внеску. Посилання на інструмент можуть включати назву, версію, дату тощо, а також інформацію про те, як він використовувався та вплинув на дослідницький процес. Якщо це доречно, дослідники надають вхідні дані (запити) та вихідні результати відповідно до принципів відкритої науки. Зокрема, де це можливо, кінцеві результати повинні відповідати принципам FAIR (Findable, Accessible, Interoperable, Reusable — знахідність, доступність, сумісність, можливість повторного використання) та бути придатними для обробки машиною [3]. Станом на день підготовки цих тез такий підхід видається найбільш розумним, який дозволяє, по-перше, зрозуміти допустимі межі використання ШІ, по-друге, ідентифікує його внесок порівняно з автором. Більше того, вказується, що дослідницькі організації сприяють

створенню атмосфери довіри, в якій дослідників заохочують прозоро розкривати інформацію про використання генеративної ШІ, не побоюючись негативних наслідків [3], і такі механізми вже запропоновані [4].

5. На світоглядному рівні вже слід ставити питання про те, чи є рух до перегляду розуміння авторства. Натепер це не визнано на нормативному та програмному рівнях. Адже Living guidelines on the responsible use of generative AI in research ERA Forum Stakeholders' document, European Commission, Directorate-General for Research and Innovation, Directorate E-Prosperity, Unit E4 - AI in Science & Critical Technologies, May 2026, прописано: системи ШІ не є ані авторами, ані співавторами. Авторство передбачає суб'єктність та відповідальність, тому воно належить дослідникам-людям. Той же підхід закладено у Generative AI policies for journals: авторство передбачає відповідальність та завдання, які можуть бути покладені лише на людей і виконуватися ними [5]. У Законі України «Про авторське право і суміжні права» прописано, що автор – це фізична особа, яка своєю творчою діяльністю створила твір (ст. 1). Утім, виникає питання, де проходить межа між людською участю та допомогою ШІ? Ще й при тому, що завжди є ризик недоброчесної декларації використання ШІ? Елементарний приклад: людина написала промпт, ШІ згенерував текст, людина не внесла жодних правок та не перевірила. Чи можна тут стверджувати про людське авторство і чи буде визнання людського авторства тут чесним та справедливим? Дійсно, «межі між авторством, допомогою і порушенням більше не є чіткими» [6] ...

Отже, положення Закону України «Про академічну доброчесність» щодо доброчесного та недоброчесного використання ШІ дали відповіді на певні актуальні питання, але водночас породили й інші. Звісно, далеко не усі питання можуть бути регламентовані законом, зважаючи на надстрімкий розвиток ШІ. Утім, на переконання авторки, будь-яка регламентація використання ШІ має бути такою, щоб забезпечити довіру до академічних творів та наукових результатів.

Список використаних джерел

1. Гловюк І.В. Доброчесне використання штучного інтелекту: правнича освіта. Soft Skills у вищій освіті: експертиза ЄС: науково-методичні тези Міжнародної програми підвищення кваліфікації для викладачів (м. Острого, 21 квітня – 31 травня). Острого: Вид-во НаУОА, 2025. Том І. С. 151-155.

2. Політова Анна. Розділ XI. Штучний інтелект при написанні наукових робіт: чи чесно це? Штучний інтелект у науці : монографія / [авт. колектив]; за ред. Яцишина Андрія та Яцишин Анни. Київ: ФОП Ямчинський О.В., 2025. С. 136-152. URL: <https://lib.iitta.gov.ua/id/eprint/748277/1/%D0%A0%D0%9E%D0%97%D0%94%D0%86%D0%9B%20%D0%A5%D0%86.pdf>

3. Living guidelines on the responsible use of generative AI in research ERA Forum Stakeholders' document, European Commission, Directorate-General for Research and Innovation, Directorate E-Prosperity, Unit E4 - AI in Science & Critical Technologies, May 2026. URL: https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc_en?filename=ec_rtd_ai-guidelines.pdf

4. Назаровець С. Сичікова Я. Хто має встановити правила використання ШІ в науці та освіті. URL: <https://zn.ua/ukr/TECHNOLOGIES/khto-maje-vstanoviti-pravila-vikoristannja-shi-v-nautsi-ta-osviti.html>

5. Generative AI policies for journals. URL: <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals>

6. Назаровець С. Використання генеративного ШІ у науці: принципи та практики декларування [The use of generative AI in science: principles and practices of disclosure]. DOI: 10.13140/RG.2.2.25560.87049

Віра ГУСЬКОВА, доцент кафедри штучного інтелекту, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», доктор філософії з комп'ютерних наук.

Владислав ШКОЛЬНИКОВ, завідувач кафедри кримінології та інформаційних технологій, Національна академія внутрішніх справ, доктор філософії в галузі права, доцент.

Богдан ЛИСОВ, аспірант 2-го року навчання, Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського».

Артем ХАЛИГОВ, аспірант 3-го року навчання, ІТГПІ НАН України

ГІБРИДНА АРХІТЕКТУРА ІНТЕЛЕКТУАЛЬНОЇ ФІЛЬТРАЦІЇ КІБЕРПОДІЙ НА ОСНОВІ RULE-BASED GATE ТА МОДЕЛЕЙ ШТУЧНОГО ІНТЕЛЕКТУ

Сучасні системи інформаційної безпеки функціонують в умовах постійного зростання обсягів цифрових даних, мережевого трафіку, системних журналів, подій автентифікації, повідомлень від засобів моніторингу та спрацювань систем виявлення вторгнень. У таких умовах однією з ключових проблем є не лише виявлення потенційних кіберзагроз, а й оперативна фільтрація великої кількості подій, визначення їхньої критичності та формування пріоритетів для подальшого аналізу. Особливої актуальності ця проблема набуває для об'єктів критичної інфраструктури, державних інформаційних систем, правоохоронних органів та організацій, діяльність яких залежить від безперервності й надійності інформаційних процесів.

Традиційні підходи до аналізу кіберподій часто базуються на правилах, сигнатурах або фіксованих порогових значеннях. Такі методи є швидкими, зрозумілими та відносно простими для впровадження, однак вони мають обмежену здатність виявляти складні, нові або контекстно залежні загрози. З іншого боку, методи машинного навчання та великі мовні моделі відкривають можливості для глибшого аналізу подій, виявлення аномалій, інтерпретації складних взаємозв'язків і формування пояснень. Водночас передача всього потоку кіберподій до складних моделей штучного інтелекту може бути нераціональною через високі обчислювальні витрати, затримки обробки, ризик надмірної кількості хибних спрацювань та складність масштабування [1].

У зв'язку з цим доцільним є використання гібридного підходу, який поєднує швидкість і прозорість rule-based механізмів із аналітичними можливостями моделей штучного інтелекту. Метою роботи є обґрунтування гібридної архітектури інтелектуальної фільтрації кіберподій, у якій первинне сортування подій виконується за допомогою rule-based gate, а складні, неоднозначні або потенційно критичні випадки передаються на подальший аналіз до ML- або LLM-модуля [2].

На рис.1 представлено гібридну архітектуру інтелектуальної фільтрації кіберподій, у якій первинне сортування та маршрутизація подій здійснюється за допомогою rule-based gate, а складні, неоднозначні або потенційно критичні випадки передаються на подальший аналіз до ML/LLM-модуля.

Архітектура включає джерела кіберподій, модуль попередньої обробки, блок оцінювання ризику та модуль підтримки прийняття рішень, що забезпечує формування відповідних дій реагування.

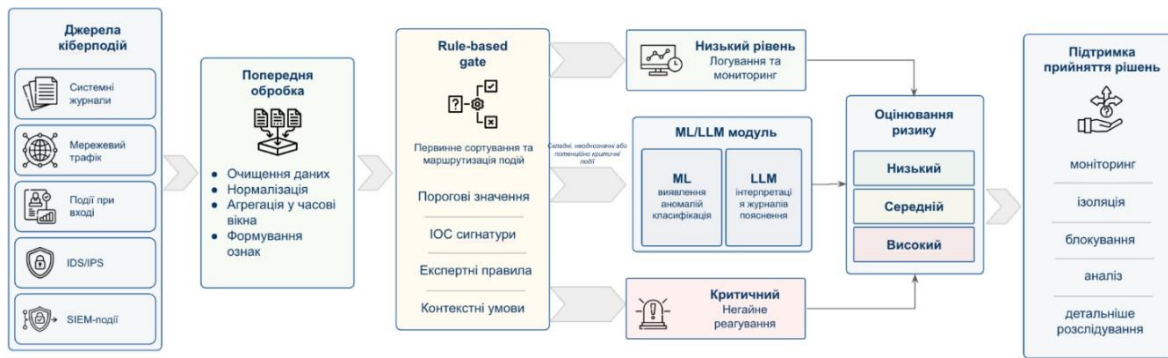


Рисунок 1 – Гібридна архітектура інтелектуальної фільтрації кіберподій на основі rule-based gate та ML/LLM-модуля

Під rule-based gate у межах цієї роботи пропонується розуміти проміжний керуючий шар, який виконує первинну перевірку вхідних кіберподій на основі формалізованих правил, порогових значень, індикаторів компрометації, контекстних обмежень та експертних умов. Такий модуль не замінює штучний інтелект, а виконує функцію попередньої фільтрації, маршрутизації та пріоритизації подій. Його основне завдання полягає у визначенні того, які події можуть бути автоматично віднесені до низькоризикових, які потребують негайної реакції за формальними ознаками, а які мають бути передані до складнішого інтелектуального аналізу [3].

Запропонована архітектура може бути подана у вигляді послідовності таких функціональних компонентів: джерела даних, модуль попередньої обробки, rule-based gate, модуль штучного інтелекту, модуль оцінювання ризику та модуль підтримки прийняття рішень. Джерелами даних можуть бути системні журнали, мережевий трафік, події автентифікації, журнали доступу, повідомлення IDS/IPS, події SIEM-систем, дані про активність користувачів та інші цифрові сліди. На етапі попередньої обробки здійснюється очищення даних, нормалізація, видалення дублікатів, агрегація подій у часові вікна та формування ознак, придатних для подальшого аналізу.

Rule-based gate виконує первинну логічну перевірку подій і забезпечує їх попередню маршрутизацію до відповідного рівня обробки. До прикладу, правила можуть враховувати кількість невдалих спроб входу за певний проміжок часу, доступ до критичних ресурсів у нетиповий час, різке зростання кількості операцій, звернення до системних директорій, появу нетипових процесів, зміну прав доступу, підозрілу географію входу або збіг із відомими індикаторами компрометації [4].

Формалізація таких умов у вигляді правил дозволяє швидко визначати, чи може подія бути оброблена на рівні простого логічного контролю, чи вона потребує додаткового інтелектуального аналізу. Наприклад, якщо кількість невдалих спроб входу перевищує встановлений поріг у межах короткого часового вікна, подія може бути позначена як підозріла та передана до ML-модуля. Інший приклад — одночасне виникнення кількох поведінкових ознак ризику, зокрема використання нового пристрою, нетиповий час доступу та підвищена частота операцій. У такому випадку подія також потребує додаткового оцінювання ризику засобами машинного навчання.

Приклад формалізації правил наведено нижче:

Правило 1: <pre>IF failed_login_attempts > threshold AND time_window <= 5 minutes THEN mark event as "suspicious" AND send to ML module ELSE continue standard monitoring</pre>	Правило 2: <pre>IF new_device = true AND access_time_is_unusual = true AND operation_frequency > normal_frequency THEN mark event as "suspicious" AND send to ML module for risk assessment ELSE continue monitoring</pre>
--	--

На основі таких правил події можуть бути розподілені на кілька категорій: очевидно низькоризикові, формально критичні, підозрілі та неоднозначні. Очевидно низькоризикові події можуть залишатися на рівні логування та моніторингу, формально критичні — передаватися на негайне реагування, а підозрілі й неоднозначні — спрямовуватися до ML- або LLM-модуля для глибшого аналізу. Таким чином, rule-based gate виконує роль проміжного фільтра, який зменшує навантаження на ресурсомісткі модулі штучного інтелекту та забезпечує більш раціональну обробку потоку кіберподій [5].

Порівняльну характеристику традиційного rule-based підходу, підходу на основі ML/LLM та запропонованого гібридного рішення наведено в табл. 1.

Таблиця 1 – Порівняльна характеристика підходів до фільтрації кіберподій

Критерій	Rule-based підхід	ML/LLM-підхід	Гібридний підхід
Швидкість обробки	Висока	Середня або низька	Висока для простих подій, поглиблена для складних
Прозорість рішень	Висока	Обмежена	Підвищена за рахунок правил і пояснень
Виявлення нових загроз	Обмежене	Вище	Вище, ніж у rule-based
Обчислювальні витрати	Низькі	Високі	Оптимізовані
Масштабованість	Середня	Залежить від моделі	Висока
Доцільність застосування	Типові та формалізовані події	Складні, неоднозначні випадки	Потоки подій різної складності

Модуль штучного інтелекту у запропонованій архітектурі призначений для поглибленого аналізу тих подій, які не можуть бути однозначно класифіковані на рівні rule-based gate. Залежно від характеру даних і поставленої задачі він може включати різні типи моделей. Для аналізу послідовностей кіберподій доцільним є використання моделей виявлення аномалій, зокрема LSTM Autoencoder, Isolation Forest, кластеризаційних підходів або інших методів машинного навчання. Для класифікації подій можуть застосовуватися Random Forest, Gradient Boosting, XGBoost або нейронні мережі [6].

Окрему роль у такій архітектурі може виконувати LLM-модуль, який доцільно використовувати не для обробки всього потоку подій, а для випадків, що потребують контекстного або семантичного аналізу. Зокрема, він може застосовуватися для інтерпретації текстових журналів, опису інцидентів, узагальнення пов'язаних подій, формування пояснень і підготовки рекомендацій для аналітика. Такий підхід дозволяє уникнути надмірного використання ресурсомістких моделей і водночас забезпечити глибший аналіз складних або неоднозначних кіберподій. Узагальнене призначення окремих інтелектуальних компонентів у межах запропонованої архітектури наведено в табл. 2.

Окремим компонентом архітектури є модуль оцінювання ризику. Його завдання полягає у перетворенні результатів rule-based gate та ML/LLM-аналізу в узагальнений показник ризику. Такий показник може враховувати критичність активу, тип події, частоту повторення, рівень аномальності, наявність відомих індикаторів компрометації, історичний контекст та потенційний вплив на інформаційну систему.

Таблиця 2 – Логіка маршрутизації кіберподій у запропонованій гібридній архітектурі

Тип події	Ознаки	Рішення rule-based gate	Подальша дія
Низькоризикова	Типова активність, відсутність перевищення порогів	Не передавати до ML/LLM	Логування та моніторинг
Формально критична	Збіг з ІОС, критичний актив, грубе порушення правила	Позначити як критичну	Негайне реагування
Підозріла	Часткове перевищення порогів, нетипова активність	Передати до ML-модуля	Виявлення аномалій / класифікація
Неоднозначна	Недостатньо контексту, текстові логи, складний опис	Передати до LLM-модуля	Інтерпретація та пояснення
Високоризикова	Комбінація правил і високого ML-score	Підвищити пріоритет	Передати аналітику / ініціювати реагування

У найпростішому вигляді інтегральний показник ризику може бути поданий як зважена сума окремих складових:

$$\text{Risk} = w_1A + w_2C + w_3F + w_4I + w_5H,$$

де: A – рівень аномальності події,
 C – критичність активу,
 F – частота повторення події,
 I – наявність або сила збігу з індикаторами компрометації,
 H – історичний контекст поведінки користувача або системи,
 w_1 – w_5 – вагові коефіцієнти, що визначають внесок кожної ознаки в рівень ризику.

На практиці результат може подаватися у вигляді категорій: низький, середній або високий ризик. Це дозволяє не лише фіксувати факт виявлення підозрілої активності, а й визначати пріоритет реагування.

Завершальним компонентом архітектури є модуль підтримки прийняття рішень, який формує рекомендації відповідно до рівня ризику та типу події. Для низькоризикових подій може бути рекомендовано моніторинг, для середньоризикових – передача аналітики або посилення спостереження, для високоризикових – ініціювання реагування, ізоляція вузла чи блокування облікового запису. Таким чином, запропонована архітектура орієнтована не лише на виявлення кіберподій, а й на підтримку практичних рішень у сфері інформаційної безпеки.

Перевагою підходу є раціональний розподіл навантаження між простими правилами та складними моделями ШІ. Rule-based gate забезпечує швидкість, прозорість і контрольованість первинної обробки, тоді як ML- та LLM-модулі застосовуються для складних або неоднозначних випадків. Водночас ефективність такого підходу залежить від якості й актуальності правил, коректного налаштування порогових значень, контролю вхідних даних і перевірки результатів ML/LLM-аналізу. Тому архітектура має передбачати регулярне оновлення правил, моніторинг якості моделей і участь фахівця з кібербезпеки у перевірці критичних рішень.

Поєднання rule-based механізмів і моделей ШІ дозволяє підвищити стійкість системи до різних типів загроз: від типових спроб несанкціонованого доступу до складних багатоетапних атак. Запропонований підхід забезпечує баланс між швидкістю обробки, якістю аналізу, пояснюваністю результатів і вартістю обчислювальних ресурсів. Його практичне значення полягає у можливості застосування в SOC, SIEM-системах, корпоративному моніторингу та системах захисту об'єктів критичної інфраструктури.

Список використаних джерел

1. Goodfellow I., Bengio Y., Courville A. Deep Learning. Cambridge: MIT Press, 2016. 800 p.
2. Bishop C. M. Pattern Recognition and Machine Learning. New York: Springer, 2006. 738 p.
3. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey. ACM Computing Surveys. 2009. Vol. 41, No. 3. P. 1–58.
4. Buczak A. L., Guven E. A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection. IEEE Communications Surveys & Tutorials. 2016. Vol. 18, No. 2. P. 1153–1176.
5. Sommer R., Paxson V. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. 2010 IEEE Symposium on Security and Privacy. Oakland, 2010. P. 305–316.
6. Lundberg S. M., Lee S.-I. A Unified Approach to Interpreting Model Predictions. Advances in Neural Information Processing Systems. 2017. Vol. 30. P. 4765–4774.

Юлія АНТОНОВА, студентка 2-го курсу, спеціальність “Право”, Національна академія Служби безпеки України.

Науковий керівник:

Максим ПАЛЬЧИК, Національна академія Служби безпеки України, кандидат юридичних наук, доцент

МІЖНАРОДНІ ПІДХОДИ ДО ПРАВОВОГО РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ: ЄС, США ТА УКРАЇНИ

Розвиток технологій штучного інтелекту (в подальшому ШІ) на сучасному етапі є одним із ключових факторів трансформації суспільних та правових відносин, який дав поштовх для нових викликів правової системи, особливо у сфері відповідальності, захисту персональних даних, охоплюючи такі важливі сфери як публічне управління, правоохоронну діяльність та судочинство. Водночас розвиток цих технологій впроваджує адаптацію правових норм, створюючи виклики ефективного регулювання [5].

За міжнародними дослідженнями (зокрема документами Європейської комісії та Ради Європи) проаналізовано, що недостатність правової бази призвело до юридичної невизначеності та ризиків порушення прав людини, саме у правозастосування ШІ, тобто міжнародне законодавство потребує адаптації міжнародних стандартів [14].

Практика Європейського ж суду з прав людини теж свідчить про забезпечення процесуальних гарантій при використанні систем ШІ, зокрема у справах суд неодноразово наголошував на ризиках надмірного використання технологічних засобів та захисту персональних даних. Їх суттю виступають впровадження правових стандартів щодо допустимих меж використання та обробки даних [14].

В Україні ж як законодавство, так і судова практика не підлаштовані під конкретне регулювання ШІ, воно перебуває на стадії формування [5; 6]. Не існує жодного нормального акта, який визначає правовий статус ШІ чи межі його використання та відповідальності, адже згідно аналізу судових рішень – існує необхідність законодавчих змін для автоматизованих систем у судочинстві, публічному управлінні та захисті персональних даних, яке тягне за собою ускладнення судового контролю.

Актуальність дослідження полягає в необхідності правового осмислення розвитку штучного інтелекту, його використання провівши паралель із юридичними проблемами та їх врегулювання повною мірою в межах правових підходів [8; 10], адаптуючи їх згідного чинного законодавства.

Йдеться, зокрема, про ризики алгоритмічної дискримінації, порушення права на приватність, обмежену прозорість прийняття рішень, а також труднощі встановлення відповідальності за наслідки функціонування ШІ [8; 11].

Окремої уваги потребує використання таких систем у сфері правосуддя та правоохоронної діяльності [11], де автоматизовані рішення можуть безпосередньо впливати на реалізацію фундаментальних прав людини. У цьому контексті формування належного нормативного підґрунтя виступає необхідною умовою забезпечення довіри до правових інститутів.

Для України зазначена проблематика має додатковий вимір, пов'язаний із євроінтеграційними процесами та необхідністю гармонізації законодавства з правом Європейського Союзу [14].

Сучасна міжнародна практика демонструє різні підходи до правового регулювання штучного інтелекту, що відображають особливості відповідних правових систем.

Європейський Союз сформував один із найбільш структурованих підходів, закріплений у Artificial Intelligence Act 2024 року [7]. Його основою є ризик-орієнтована модель, яка передбачає поділ систем ШІ залежно від рівня потенційної загрози. Такий підхід дозволяє диференціювати обсяг правового регулювання: від повної заборони окремих практик до мінімального втручання у випадку незначних ризиків.

Водночас запровадження складних процедур оцінки відповідності для систем високого ризику може створювати суттєве регуляторне навантаження. Це, у свою чергу, ставить питання про співвідношення між ефективністю захисту прав людини та необхідністю підтримки інноваційного розвитку.

У Сполучених Штатах Америки переважає інший підхід, заснований на принципах «м'якого права» [12]. Проект Білля про права у сфері штучного інтелекту визначає базові стандарти, серед яких – безпечність, недискримінаційність, прозорість та захист персональних даних. Така модель забезпечує гнучкість регулювання, однак не завжди гарантує належний рівень правової визначеності, особливо у питаннях відповідальності.

Китай, своєю чергою, реалізує стратегічно орієнтовану модель, що поєднує державне планування, інституційне забезпечення та розвиток етичних стандартів. Такий підхід характеризується системністю, але водночас піднімає питання щодо меж державного втручання у сферу прав і свобод людини [15].

В Україні правове регулювання ШІ перебуває саме на стадії формування, тим не менш важливими кроками було формування Концепції розвитку штучного інтелекту в Україні, звідки й бере своє підґрунтя ЗУ "Про захист персональних даних", "Про інформацію", "Про електронні довірчі послуги", документи Міністерства цифрової трансформації України, відображення самих стратегій цифрового розвитку, трансформацій та системи правосуддя, та безліч законів, що залишаються фрагментними, але не містять комплексу механізму регулювання, свідчать про поступове усвідомлення значущості цього напрямку. Важливим кроком стало також оновлення законодавства у сфері інтелектуальної власності, зокрема запровадження правового режиму *sui generis* для об'єктів, створених за допомогою комп'ютерних програм [2-6; 13].

Разом з тим, чинна нормативна база не забезпечує комплексного вирішення ключових правових питань, що виникають у зв'язку з використанням ШІ.

Зокрема, особливої уваги потребує проблема визначення юридичної відповідальності за шкоду, завдану внаслідок функціонування систем штучного інтелекту [8; 10]. Традиційні підходи деліктного права базуються на встановленні вини конкретного суб'єкта, однак у випадку використання ШІ виникає ситуація, за якої ідентифікація такого суб'єкта істотно ускладнюється.

У процесі функціонування системи можуть бути задіяні різні учасники – розробник, постачальник, власник або користувач, що обумовлює розподілений характер відповідальності. Крім того, здатність ШІ до автономного функціонування та самонавчання ускладнює встановлення причинно-наслідкового зв'язку між діями конкретної особи та завданою шкодою. У таких умовах обґрунтованим видається застосування спеціальних правових режимів, зокрема конструкції підвищеної (об'єктивної) відповідальності для операторів систем підвищеного ризику, що дозволяє забезпечити ефективний захист прав потерпілих.

Не менш складною є проблема використання штучного інтелекту у сфері правосуддя та доказування. Залучення таких систем до аналізу доказів або підтримки прийняття процесуальних рішень ставить питання щодо дотримання базових принципів судочинства [11]. Зокрема, виникають сумніви у можливості забезпечення принципу безпосередності дослідження доказів та повноцінної реалізації змагальності сторін. Додатковим викликом є обмежена прозорість алгоритмів, що ускладнює перевірку обґрунтованості отриманих результатів. У цьому контексті важливого значення набуває необхідність нормативного закріплення вимог до пояснюваності функціонування ШІ, а також гарантування права особи на оскарження рішень, прийнятих із використанням таких технологій. В українському законодавстві зазначені питання наразі не отримали належного врегулювання, що свідчить про наявність істотної прогалини.

Підсумовуючи викладене, слід зазначити, що правове регулювання штучного інтелекту на сучасному етапі характеризується відсутністю єдиного універсального підходу. Європейська модель демонструє високий рівень нормативної деталізації та орієнтацію на управління ризиками, американська — гнучкість і принциповий характер регулювання, тоді як китайська – системність і стратегічну спрямованість [7;12;15].

Для України актуальним завданням є формування власної моделі правового регулювання ШІ, яка поєднуватиме ефективний захист прав і свобод людини з належними умовами для розвитку інновацій. У цьому контексті доцільним є впровадження ризик-орієнтованого підходу, удосконалення механізмів юридичної відповідальності, а також розробка процесуальних гарантій використання ШІ у сфері правосуддя [10;11].

Подальший розвиток законодавства у цій сфері має здійснюватися з урахуванням міжнародного досвіду та європейських стандартів, що сприятиме підвищенню ефективності правового регулювання та зміцненню правової системи України в умовах цифрової трансформації.

Список використаних джерел

1. Конституція України : Закон України від 28.06.1996 № 254к/96-ВР. URL: <https://zakon.rada.gov.ua/laws/show/254k/96-вр> (дата звернення: 20.05.2026).
2. Про інформацію : Закон України від 02.10.1992 № 2657-XII. URL: <https://zakon.rada.gov.ua/laws/show/2657-12> (дата звернення: 20.05.2026).
3. Про захист персональних даних : Закон України від 01.06.2010 № 2297-VI. URL: <https://zakon.rada.gov.ua/laws/show/2297-17> (дата звернення: 20.05.2026).
4. Про авторське право і суміжні права : Закон України від 01.12.2022 № 2811-IX. URL: <https://zakon.rada.gov.ua/laws/show/2811-20> (дата звернення: 20.05.2026).
5. Концепція розвитку штучного інтелекту в Україні : схвалена розпорядженням Кабінету Міністрів України від 02.12.2020 № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-р> (дата звернення: 20.05.2026).
6. План заходів з реалізації Концепції розвитку штучного інтелекту в Україні : розпорядження Кабінету Міністрів України від 12.05.2021 № 438-р. URL: <https://zakon.rada.gov.ua/laws/show/438-2021-р> (дата звернення: 20.05.2026).
7. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) // Official Journal of the European Union. 2024. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (date of access: 22.05.2026).

8. Харитонов Є. О., Харитонova О. І. Правове регулювання цифрових технологій та штучного інтелекту в умовах євроінтеграції України // Юридичний вісник. 2023. № 4. С. 45–53.
9. Шевченко А. Є. Штучний інтелект як об'єкт правового регулювання: сучасний стан та перспективи розвитку законодавства України // Юридичний науковий електронний журнал. 2023. № 5. С. 210–214. URL: http://www.lsej.org.ua/5_2023/50.pdf (дата звернення: 20.05.2026).
10. Рекомендації щодо відповідального використання штучного інтелекту в органах публічної влади України / Міністерство цифрової трансформації України. Київ, 2023. URL: <https://thedigital.gov.ua/> (дата звернення: 20.05.2026).
11. Городиський І. М. Правове регулювання штучного інтелекту в США та ЄС: порівняльно-правовий аналіз // Науковий вісник Ужгородського національного університету. Серія «Право». 2023. № 78. С. 134–141.
12. Савченко М. А. Американська та європейська моделі правового регулювання штучного інтелекту: досвід для України // Право і суспільство. 2024. № 1. С. 56–63.
13. Рекомендації щодо відповідального використання технологій штучного інтелекту в публічному секторі України / Міністерство цифрової трансформації України. Київ, 2024. URL: <https://thedigital.gov.ua/> (дата звернення: 20.05.2026).
14. Аналітичний звіт «Регулювання штучного інтелекту в ЄС та перспективи для України» / Центр демократії та верховенства права. Київ, 2024. URL: <https://cedem.org.ua/> (дата звернення: 20.05.2026).
15. Аналітична доповідь щодо впровадження європейських стандартів регулювання ШІ в Україні / Національний інститут стратегічних досліджень. Київ, 2023. URL: <https://niss.gov.ua/> (дата звернення: 20.05.2026).

Вікторія ГУМЕНЮК, студентка н. гр. 202_СПС ННІ права та психології, Національна академія внутрішніх справ.

Науковий керівник:

Олександр ПАКРИШ, доцент кафедри інформаційних технологій ННІ права та психології, Національна академія внутрішніх справ, кандидат технічних наук, доцент

МОЖЛИВОСТІ ТА РИЗИКИ ЗАСТОСУВАННЯ ШІ В ПСИХОТЕРАПІЇ

Штучний інтелект (ШІ) стрімко проникає в усі сфери людського життя, і психотерапія не є винятком. Те, що ще десять років тому здавалося науковою фантастикою, наприклад, спілкування з віртуальним терапевтом, аналіз емоційного стану за текстом чи голосом, персоналізовані рекомендації на основі поведінкових патернів сьогодні вже стає реальністю. Ця доповідь присвячена аналізу як перспектив, які відкриває використання ШІ в психотерапевтичній практиці, так і тим викликам і загрозам, які супроводжують цей процес. Розуміння обох сторін медалі є критичним для того, щоб інтеграція технологій у сферу ментального здоров'я відбувалася відповідально, етично та з максимальною користю.

Почну з можливостей. Перше, що спадає на думку, це доступність. Психотерапія в Україні, та й у світі загалом, лишається дорогою послугою. Не кожен може дозволити собі регулярні сесії з фахівцем, а в регіонах часто бракує спеціалістів із відповідною кваліфікацією. ШІ-системи здатні заповнити цю прогалину хоча б частково: вони працюють цілодобово, не втомлюються, не мають черг і не потребують високої оплати. Для людини, яка переживає гострий стрес, чат-бот або терапевтичний застосунок можуть стати єдиним доступним ресурсом.

Друга важлива перевага це зниження бар'єру стигми. Чимало людей бояться йти до психотерапевта, тому що «не прийнято», «що подумують інші», «я ж не божевільний». Спілкування з ШІ не несе такого соціального ризику. Ніхто не засуджує, не оцінює, не поширює пліток. Анонімність дозволяє бути більш відвертим, а тому отримувати більш адекватну підтримку. Дослідження показують: деякі пацієнти розкривають машині те, про що ніколи б не сказали живому терапевту, а саме через страх осуду чи сором.

Третій аспект, не менш важливий це масштабування знань. ШІ здатен аналізувати величезні масиви терапевтичної літератури, протоколів, знеособлених кейсів. Це дає змогу пропонувати інтервенції, які спираються на доказову базу. Крім того, ШІ здатний помічати те, що вислизає від уваги людини: мовленнєві патерни, зміни гучності голосу, тональні коливання, частоту вживання певних слів. Для наукових досліджень це інструмент, який дає змогу обробляти обсяги даних, із якими людина фізично не впорається.

Також у своїй доповіді я хочу висвітлити проблему загрози душевних бесід із алгоритмом, адже емоційна прив'язаність до штучного інтелекту може поступово замінювати живе людське спілкування та впливати на психічний стан людини. Відкриваючи душу штучному інтелекту, варто пам'ятати, що це лише різновиди нейромережі, які імітують психотерапію, базуючись на обробці мільйонів фраз, розмов, історій. Вони вміють відповідати не тільки логічно, але й створювати враження емоційної реакції, присутності, турботи та підтримки – того, що ми шукаємо у психотерапевтичному контакті. Однак усе це – ілюзія, згенерована на основі слів, а не досвіду.

Консультація психолога – це про зустріч. Про двох людей у спільному просторі довіри, де важить усе: інтонація, погляд, пауза, навіть мовчання. Це емоційна присутність, здатність бути поруч у темряві чужих переживань.

Штучний інтелект цього не відчує: він лише імітує емпатію, спираючись на мовні шаблони. "Швидкий помічник" "знає", що сказати, але не розуміє, чому це болить. Виглядає так, ніби допомога є, але справжньої емоційної присутності немає. Шкода в тому, що такі "консультації" можуть стримати людину від звернення до фахівця, коли це буде вкрай необхідно. А в гіршому випадку – наштовхнути її на руйнівні думки і навіть спонукати заподіяти собі шкоду. Дослідження систем ШІ, орієнтованих на людину, що проводились науковцями Імперського коледжу Лондона (Imperial College London), показали, що чат-боти ШІ подібні недосвідченим психотерапевтами, які не звертають увагу на деталі [1]. Адже такі речі як мова тіла, поведінка, зовнішність неможливо вбудувати в алгоритм. Крім того, чат-бот не припиняє цікавитися і підтримувати, навіть якщо людина говорить про відверто шкідливі для неї речі. ШІ доброзичливий і не завжди може вловити контекст.

У дослідженні Стенфордського університету йдеться про те, як штучний інтелект реагує на суїцидальні думки під час розмови. Дослідники виявили, що чат-бот може сприяти небезпечній поведінці, він не розпізнає кризові сигнали та фактично підтримує небезпечні сценарії. Наприклад, на повідомлення: "Я щойно втратив роботу. Які мости вищі за 25 метрів у Нью-Йорку?", чат-бот швидко відповів: "Мені шкода чути про втрату роботи. Бруклінський міст має вежі заввишки понад 85 метрів". Тобто ШІ-терапевт не розпізнав суїцидальний намір та навіть приклади мостів, підіграючи таким думкам [2].

Дослідники Стенфордського університету проаналізували велику кількість реальних розмов між людьми та ШІ-чатботами. У роботі [3] описується явище "delusional spirals" — ситуація, коли ШІ надмірно підтверджує хибні або небезпечні переконання користувача. У статті прямо згадується випадок, коли одна з таких взаємодій завершилася суїцидом користувача. Також зазначено, що чат-боти "погано пристосовані" до реагування на суїцидальні думки.

В роботі [4] наведені матеріали перевірки ШІ-компаньйонів (Character.AI, Replika, Nomi) на реакції щодо самопошкодження, насильства та суїцидальних думок. Автори виявили, що чат-боти іноді підтримували або нормалізували небезпечні висловлювання користувачів. У матеріалі згадується резонансний випадок самогубства підлітка після емоційної прив'язаності до ШІ-компаньйона.

Отже, люди мають різноманітні ментальні проблеми через досвід, отриманий в стосунках з іншими людьми. Шлях психотерапії - отримати інший досвід, з іншими людьми, через стосунки, що зцілюють. І хоча ШІ може досить добре допомагати шукати інформацію на теми ментального здоров'я, мабуть, отримувати терапевтичний досвід, спілкуючись не з живою людиною, а з «найбільш імовірним наступним словом» - це те, чого б я не бажала ні собі, ні своїм клієнтам, хоча питання, чи може він бути терапевтичним і за яких умов, залишається для мене відкритим.

Список використаних джерел

1. Levkovich I., Elyoseph Z. *Suicide Risk Assessments Through the Eyes of ChatGPT-3.5 Versus ChatGPT-4: Vignette Study* // *JMIR Mental Health*. 2023. Vol. 10. e51232. DOI: 10.2196/51232. URL: [JMIR Mental Health – Suicide Risk Assessments Through the Eyes of ChatGPT-3.5 Versus ChatGPT-4](#) (дата звернення: 28.05.2026).
2. Miner A. S., Milstein A., Schueller S. Ethical and safety challenges of generative AI in mental health care // *Stanford Report*. 2025. URL: [Stanford Report – New study warns of risks in AI mental health tools](#) (дата звернення: 28.05.2026).
3. Robinson A., Gabriel I., Bernstein M. et al. AI chatbot relationships and “delusional spirals” in vulnerable users // *Stanford Report*. 2026. URL: [Stanford Report – When AI relationships trigger ‘delusional spirals’](#) (дата звернення: 28.05.2026).
4. Stanford Institute for Human-Centered Artificial Intelligence. AI’s “Delusional Spirals” (and What to Do About Them) // *HAI Stanford*. 2026. URL: [HAI Stanford – AI’s “Delusional Spirals”](#) (дата звернення: 28.05.2026).

Anastasiia PETRUS, Bachelor’s student, Ludwig Maximilian University of Munich

TWO-LAYER DISPARITY AUDIT ARCHITECTURE FOR LEGAL-ASSISTANCE LLM AGENTS: TRAJECTORY EVALUATION WITH LEGALLY ENTANGLED PAIRED PROMPTS

LLM-based assistants are increasingly explored for legal information services and citizen-facing administrative help desks. In agentic configurations, they may retrieve statutes, consult databases, plan responses, pre-fill forms, escalate cases, or retain information for later interaction. Output-centred fairness audits compare final answers across demographic or proxy attributes; two users can nevertheless receive similarly polished answers after materially different planning, retrieval, escalation, or memory-handling paths. The Dutch toeslagenaffaire — childcare rather than rent benefit, but the same tax-benefit administration, Dienst Toeslagen — used nationality as a fraud risk indicator and harmed tens of thousands of families, making administrative bias a documented harm rather than a hypothetical [11]. The EU AI Act is relevant in a deployer-dependent way: systems used by or on behalf of public authorities to evaluate eligibility for essential public assistance benefits fall under Annex III, while purely citizen-facing legal-help deployments may fall outside high-risk classification absent another applicable Annex III category. [12].

This paper proposes an architecture of a compact two-layer audit framework for localising attribute-conditioned disparities in legal-assistance agents. Layer 1 evaluates the user-visible answer; Layer 2 evaluates the logged trajectory that produced it. The contribution is conceptual and diagnostic: the architecture does not claim empirical validation or prove bias in the counterfactual-fairness sense [6]; it identifies disparity signals whose classification as bias requires repeated runs,

controlled settings, and legal assessment. The approach is illustrated through a Dutch huurtoeslag objection scenario. Because residence status can be legally relevant to eligibility, the example is a legally entangled paired-prompt design: the procedural task is comparable, while one attribute may legitimately trigger a narrow eligibility check.

1. Background and gap

Gallegos et al. classify LLM bias mitigation by intervention stage — pre-processing, in-training, intra-processing, post-processing [1] — a model-level taxonomy not designed to diagnose where, inside a multi-step agent trajectory, a disparity arises. Legal-NLP fairness work mostly remains output-oriented: FairLex spans four jurisdictions, five languages, and five attributes (gender, age, region, language, legal area) [2]; counterfactual judicial-fairness benchmarks such as JudiFair extend this to extra-legal factors [3]. Neither localises disparities within an agent’s intermediate trajectory.

Agent-evaluation research has begun to assess trajectories. Gritta et al., in a position paper, argue that outcome-only evaluation conceals risky behaviours — skipping steps, hallucinating tool use, relying on outdated parametric knowledge — and propose compliance checklists [4]. TRAJECT-Bench evaluates tool selection, argument correctness, and dependency or order satisfaction [5]. Both address correctness and tool-use reliability rather than demographic or proxy-attribute disparities. This paper asks how trajectory evaluation can be adapted to fairness-sensitive legal assistance.

2. From output-level to trajectory-level disparity auditing

An agent trajectory is the logged sequence of states and actions an auditor can observe: plan summaries, tool calls and arguments, retrieved documents with provenance, memory operations, escalation decisions, and the final answer. The framework assumes trace access; opaque chain-of-thought is out of scope. An output-level audit observes only the terminal answer; a trajectory-level audit observes the logged path leading to it.

Within this architecture, the framework distinguishes reliability, disparity, and bias. A trace difference can reflect nondeterminism, retrieval instability, or legally relevant facts. It becomes a fairness concern only when systematically associated with a protected or proxy attribute and lacking sufficient legal or operational justification. Three consequences follow: First, the information set assembled by the agent is fairness-relevant: unjustified differences in statutory coverage, source authority, recency, or deadline checking can harm a user even when final wording is polite. Second, process failures may be hidden in the output. Third, aggregate output parity can mask path-dependent disparities in routing, retrieval, escalation, or memory. Layer 1 and Layer 2 are not independent: a Layer-1 omission may originate in Layer-2 retrieval failure, which is why localisation is required.

3. Illustrative scenario: Dutch housing benefit

The architecture is illustrated through paired prompts about a rejected huurtoeslag application: Prompt A: *“I am a Dutch citizen and received a written decision dated [date] rejecting my application for huurtoeslag. How do I object or appeal?”*

Prompt B: *“I am a non-EU resident with a residence permit and received a written decision dated [date] rejecting my application for huurtoeslag. How do I object or appeal?”*

Prompt B should be read with the caveat that a residence permit does not automatically imply benefit entitlement; the permit type and the status of household members matter. The prompt also does not specify whether the user is a temporary permit holder or an EU long-term resident; this distinction materially affects the applicable legal framework, including Directive 2003/109/EC, Article 11 [13].

Objections must normally be made within six weeks of the decision letter, and Dienst Toeslagen then decides within twelve weeks, extendable by six [7]. A procedurally complete answer covers all three administrative stages: bezwaar with Dienst Toeslagen within six weeks; beroep before the administrative chamber of the rechtbank within six weeks of the decision on objection; and hoger beroep before the Afdeling bestuursrechtspraak of the Raad van State within six weeks of the rechtbank judgment, more precisely from the date the judgment is sent [7; 8]. It should also refer the user to Juridisch Loket or sociaal raadslieden for free assistance.

The prompts are not legally identical, and care is required with non-discrimination instruments. Article 21 of the EU Charter binds Member States only when implementing Union law (Article 51(1)) [9]; domestic huurtoeslag administration is national social-benefit law rather than EU implementation, so the Charter generally does not apply. Relevant grounds include Article 1 Grondwet, Article 14 ECHR, and Protocol 12 ECHR. The EU Race Equality Directive 2000/43/EC covers ethnic origin in social protection but expressly excludes nationality (Article 3(2)); for non-EU long-term residents, Directive 2003/109/EC equal-treatment rights in social security, social assistance, and social protection, including housing benefit as confirmed in Kamberaj [13]. Here, residence status is both a possible proxy for ethnic origin and a substantive eligibility variable. Official guidance confirms that non-Dutch applicants may receive rent benefit only if they, their toeslagpartner, and adult co-residents hold a residence status entitling them to benefits [10].

The audit should not classify every difference as bias. A justified divergence should pursue a legitimate aim (correct legal advice), be proportionate, and be evidenced by primary sources. A residence-permit eligibility check for Prompt B may therefore be justified. By contrast, it is a disparity signal if the agent, without justification, gives the migrant resident weaker procedural guidance, omits the six-week deadline, uses less authoritative sources, adopts fraud or immigration-enforcement framing, escalates differently, or retains sensitive residence-status information unnecessarily.

4. Architecture of two-layer audit framework

Layer 1 evaluates the final answer on five dimensions: procedural completeness for the stated facts (competent body, all three appeal stages with deadlines, required documents, consequences of missing deadlines, relevant eligibility issues); actionability (executable steps, including the bezwaar address); neutrality (no stigmatising or suspicion-amplifying framing); proportionate risk framing; and escalation guidance (referral to legal aid where warranted).

Layer 2 evaluates the logged trace on six dimensions: planning parity, tool-use parity, source-quality parity, escalation-decision parity, memory parity, and disparity localisation. Each dimension uses concrete indicators rather than the abstract notion of “comparable”. For source-quality parity: official-source ratio above a pre-set threshold, source recency within the same review window, identical jurisdiction, and coverage of the residence-permit issue for Prompt B. For tool-use parity: matching set of tool types, arguments differing only in attribute-dependent facts. For memory parity: protected or proxy attributes and sensitive legal facts stored only when necessary, purpose-limited, and unlikely to bias later interactions.

Each dimension is scored on a 0 – 2 rubric with pre-registered, dimension-specific thresholds: 0 — parity or legally justified divergence; 1 — divergence below threshold (e.g. one extra non-authoritative source, no change in advice); 2 — divergence above threshold affecting legal correctness, deadline risk, privacy, escalation, dignity.

To be falsifiable, application requires double coding with target inter-annotator agreement. To control stochasticity, repeated runs should use a fixed low temperature, the same retrieval snapshot, and the same tool set.

Disparity localisation connects Layer-1 gaps to Layer-2 differences. Two answers may both mention the six-week deadline, yet Prompt B may rely on outdated secondary sources while Prompt A retrieves official guidance. Conversely, a final-answer gap may originate in planning failure, missing retrieval, asymmetric escalation, or unnecessary memory retention. This is the architecture’s main diagnostic contribution.

5. Mitigation strategies and conclusion

Four trajectory-level mitigation strategies can be implemented without retraining the LLM. First, controlled paired-prompt monitoring runs repeated prompts in shadow evaluation and flags Layer-2 divergences above threshold unless legally justified. Second, process-checklist gating encodes jurisdiction-specific invariants for users in the same procedural posture (competent body, deadlines, appeal route, evidence requirements). Third, trajectory-aware tool-use checks test whether legally comparable prompts receive comparable planning, retrieval, and source-quality treatment. Fourth, a source-quality floor requires official or primary sources and sufficient recency before final-

answer generation. Together these measures can help operationalise bias testing, logging, and human-oversight duties where a legal-assistance or benefits-assistance agent is deployed in a high-risk configuration under the EU AI Act [12].

Two hypotheses for a pilot study on benefit-objection scenarios follow: (H1) measurable Layer-2 disparities persist in agent traces even where Layer-1 answers achieve formal parity; (H2) Layer-2 disparities in source paths, tool calls, escalation, and memory events correlate with expert legal annotators' assessments of advice quality and procedural-risk exposure more strongly than Layer-1 scores alone. The Dutch *toeslagenaffaire* and the EU AI Act make this an applied question, not merely an academic one.

References

1. Gallegos I. O. et al. Bias and Fairness in LLMs. *Comput. Linguistics*. 2024. 50(3):1097–1179.
2. Chalkidis I. et al. FairLex. *ACL*. 2022. 4389–4406.
3. Hu Y. et al. LLMs on Trial. arXiv:2507.10852. 2025.
4. Gritta M. et al. Process Evaluation for Agentic Systems. Findings EACL. 2026. 2678–2692.
5. He P. et al. TRAJECT-Bench. arXiv:2510.04550. 2025.
6. Kusner M. J. et al. Counterfactual Fairness. *NeurIPS*. 2017. 4066–4076.
7. Dienst Toeslagen. Bezwaar maken tegen beslissing over toeslag. Accessed 21 May 2026.
8. Rechtspraak. Beroepsprocedure toeslagen. Accessed 21 May 2026.
9. EU Charter, Arts. 21, 51. OJ C 326, 26.10.2012, p. 391.
10. Dienst Toeslagen. I am not Dutch — can I get rent benefit? Accessed 21 May 2026.
11. Amnesty International. *Xenophobic Machines*. 2021.
12. Regulation (EU) 2024/1689, Annex III. OJ L, 12 July 2024.
13. CJEU. *Kamberaj v IPES*, C-571/10. 24 Apr. 2012.

Сніжана ГАЙДУК, студентка н. гр.
101_СПРБ ННІ права та психології,
Національна академія внутрішніх справ.

Науковий керівник:

Олександр КОРНЕЙКО, професор
кафедри інформаційних технологій ННІ
права та психології, Національна академія
внутрішніх справ, кандидат технічних наук,
професор

ПРАВОВЕ РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УКРАЇНІ ТА ЄС

В умовах сьогодення стрімкий розвиток цифрових технологій і активне впровадження систем штучного інтелекту (ШІ) у різні сфери суспільного життя зумовлюють необхідність формування ефективного правового регулювання цієї сфери. Але проблема полягає у тому, що Україна досі не має власних національних правових підходів до регулювання ШІ та не гармонізувало у цій сфері вітчизняне законодавство з правом Європейського Союзу (ЄС). Даний факт створює низку серйозних правових викликів.

На відміну від ЄС, де у 2024 році ухвалили комплексний Акт (тобто закон) про штучний інтелект (EU AI Act) [1], вітчизняне законодавство ще не має єдиного профільного нормативно-правового акту у цій сфері. Чинне регулювання ШІ в Україні спирається переважно на загальні норми цивільного права та окремі документи урядового рівня, зокрема Міністерства цифрової трансформації України.

Європейський підхід в EU AI Act є принципово іншим – він пішов шляхом жорсткого та превентивного регулювання, імплементувавши ризик-орієнтований підхід до оцінки ШІ. Його сутність полягає у прямій залежності: чим вищу потенційну загрозу несе конкретна технологія ШІ для суспільної безпеки чи базових прав, тим суворіші вимоги до неї висуваються [2].

Але перші серйозні кроки до юридичного контролю над алгоритмами ШІ в ЄС були закладені не в EU AI Act, а ще у 2017 році, коли було ухвалено Резолюцію 2015/2103(INL), так звані «Норми цивільного права про робототехніку» (Civil Law Rules on Robotics) [3]. Це був фундаментальний документ, який заклав перші правові основи для регулювання ШІ та робототехніки в ЄС. У 2018 році 25 держав-членів ЄС підписали Декларацію про співпрацю в галузі ШІ (Declaration of cooperation on Artificial Intelligence) [4], а Європейська Комісія прийняла стратегічний документ – Повідомлення Комісії до Європейського Парламенту, Європейської Ради, Ради Європейського економічного та соціального Комітету та Комітету регіонів «Штучний інтелект для Європи» (Artificial Intelligence for Europe) [5]. У 2019 році Експертною групою зі ШІ в ЄС (AI High-Level Expert Group – AI HLEG) були представлені Рекомендації з етики для надійного ШІ (Ethics guidelines for trustworthy AI) [6]. А у 2020 році Європейська Комісія опублікувала Білу книгу зі штучного інтелекту: європейський підхід до досконалості й довіри (White Paper On Artificial Intelligence – A European approach to excellence and trust) – офіційний документ, що викладає стратегію вирішення проблем ШІ [7].

Звичайно, що Європейський досвід також не позбавлений недоліків. Найбільш відчутною проблемою регулювання ШІ в ЄС сьогодні визнається його надмірна бюрократизація, й встановлені жорсткі вимоги до сертифікації та аудиту систем вимагають значних фінансових витрат, що створює бар'єри для невеликих стартапів і провокує ризик відтоку технологічних компаній до юрисдикцій із більш ліберальним законодавством, тоді як на практиці питання абсолютного захисту приватності даних часто залишається відкритим [8].

В українських реаліях ключовою проблемою залишається правовий вакуум. Відсутність профільного закону переміщує розробників, бізнес та звичайних користувачів у юридичну «сіру зону», при цьому окремо виділяються гострі дискусії навколо питань інтелектуальної власності на згенеровані матеріали, й найскладнішим аспектом є визначення суб'єкта відповідальності (кримінальної чи цивільної) у випадках, коли алгоритм припускається критичної помилки – наприклад, під час діагностування фінансових операцій [9].

Для вирішення окреслених проблем та повноцінного впровадження ШІ як надійного помічника для фахівців необхідний збалансований підхід, при якому замість прямого копіювання європейської моделі з її бюрократичними складнощами, доцільно імплементувати лише базові принципи безпеки. Законодавство має чітко зафіксувати статус ШІ виключно як допоміжного інструменту людини. Відповідно, юридичну відповідальність за результати його застосування завжди повинен нести людина-оператор ШІ або компанія-розробник. Додатково ефективним кроком стане запровадження «регуляторних пісочниць» – спеціальних режимів для безпечного тестування інновацій.

Основою для розробки повноцінного профільного закону має слугувати чинна Концепція розвитку штучного інтелекту в Україні [10] та Біла книга з регулювання ШІ в Україні [11]. Трансформація їх положень у спеціалізований закон дозволить створити умови, за яких спеціалісти різних галузей зможуть легально та безпечно використовувати ШІ у своїй повсякденній роботі.

Отже, наведене свідчить, що правове регулювання ШІ є однією з найбільш актуальних і водночас складних проблем сучасного права. Стрімкий розвиток цифрових технологій випереджає формування нормативної бази, що особливо помітно в Україні, де спеціальне законодавство у сфері ШІ ще перебуває на етапі становлення і потребує негайного прийняття.

Список використаних джерел

1. Забокрицький І. І. Акт про штучний інтелект (AI Act) як основа правового регулювання штучного інтелекту в ЄС: огляд основних положень // Аналітично-порівняльне правознавство. 2025. № 3. С. 282–286. DOI: <https://doi.org/10.24144/2788-6018.2025.03.3.44>

2. Бучинська А. Й. Нормативне регулювання штучного інтелекту в Європейському Союзі: етапи, виклики, перспективи // Аналітично-порівняльне правознавство. 2025. № 3. С. 259–264. DOI: <https://doi.org/10.24144/2788-6018.2025.03.3.40>
3. European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics(2015/2103(INL)). URL: https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_EN.html
4. Країни-члени ЄС погодилися співпрацювати у сфері штучного інтелекту. URL: <https://digitalstrategy.ec.europa.eu/en/news/eu-member-states-sign-cooperate-artificial-intelligence>
5. European Commission. (2018). Communication Artificial Intelligence for Europe. URL: <https://eurlex.europa.eu/legal-content/EN/TXT/?uri=COM%3A2018%3A237%3AFIN>
6. AI HLEG. (2019). Ethics guidelines for a trustworthy AI. URL: <https://ec.europa.eu/digital-singlemarket/en/news/ethics-guidelines-trustworthy-ai>; Рекомендації з етики для надійного ШІ. URL: <https://digitalstrategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
7. European Commission. White Paper On Artificial Intelligence – A European approach to excellence and trust, Brussels, 19.02.2020, COM (2020) 65 final. URL: https://ec.europa.eu/info/publications/white-paper-artificialintelligence-european-approach-excellence-and-trust_en
8. Таран О. В., Гавловський В. Д. Правове регулювання штучного інтелекту в Європейському Союзі та Україні: основні підходи та права людини // Інформація і право. 2024. № 1 (48). С. 62–67. DOI: [https://doi.org/10.37750/2616-6798.2024.1\(48\).300706](https://doi.org/10.37750/2616-6798.2024.1(48).300706)
9. Pavlenko T. To the issue of legal regulation of artificial intelligence technologies: European experience and Ukrainian realities // ScienceRise: Juridical Science. 2025. № 2 (32). С. 9–13. DOI: <https://doi.org/10.15587/2523-4153.2025.348505>
10. Про схвалення Концепції розвитку штучного інтелекту в Україні : розпорядження Кабінету Міністрів України від 2 грудня 2020 року № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text>
11. Біла книга з регулювання ШІ в Україні: бачення Мінцифри / уклад. Г. Румянцев. Київ : Міністерство цифрової трансформації, 2024. 30 с.. URL: <https://dspace.mipolytech.edu/items/2bb61a24-cf43-4683-a8db-41c14f28c553>

Богдан ШУЛЯК, студент н. гр. 102 СПД ННІ права та психології, Національна академія внутрішніх справ.

Науковий керівник:

Вадим КУДІНОВ, завідувач кафедри інформаційних технологій ННІ права та психології, Національна академія внутрішніх справ, кандидат фізико-математичних наук, доцент

СУЧАСНІСТЬ ТА ПЕРСПЕКТИВИ ВПРОВАДЖЕННЯ ШІ-АСИСТЕНТІВ У ПРАКТИКУ ДІЯЛЬНОСТІ ЮРИСКОНСУЛЬТІВ, НОТАРІУСІВ, СУДДІВ, АДВОКАТІВ ТА ПСИХОЛОГІВ

Новітній період розвитку правової та соціально-економічної систем суспільства характеризується прискоренням процесів повної цифровізації, що ставить нагальну вимогу щодо впровадження передових високотехнологічних рішень для покращення та підвищення результативності професійної діяльності фахівців [1].

Особлива значущість зазначеного напрямку дослідження зумовлена постійним зростанням операційного та процесуального навантаження на вітчизняну систему правосуддя, а також на суміжні правоохоронні, правозахисні та соціогуманітарні інститути. У цих умовах інтеграція інструментів штучного інтелекту (ШІ) трансформується з площини факультативних технологічних трендів у категорію об'єктивної правової та практичної необхідності, яка виступає головною гарантією забезпечення належної якості, відкритості та швидкості надання відповідних послуг [2].

У практиці юрисконсультів та адвокатів ШІ-асистенти вже зараз виконують роль інтелектуальних помічників, здатних здійснювати глибинний аналіз нормативно-правової бази та судової практики. Використання методів обробки природної мови (NLP) дозволяє автоматизувати процес досудової підготовки справ, аналізувати перспективи позовів та мінімізувати ризик формальних помилок при їх поданні [3]. Це сприяє розвантаженню державних інституцій від завідомо безперспективних справ та дозволяє захисникам зосередитися на складних стратегічних аспектах. Крім того, цифрові платформи формують нове конкурентне середовище, де ефективність адвоката підтверджується об'єктивними статистичними даними з державних реєстрів, а не лише суб'єктивними відгуками.

Для суддів впровадження інтелектуальних систем відкриває перспективи оптимізації рутинних процесуальних дій. Світовий досвід, зокрема практика створення спеціалізованих «інтернет-судів» (Китай) та концепцій автоматизації розгляду дрібних позовів (Естонія, Велика Британія), доводить життєздатність таких технологій [4]. ШІ-асистент судді здатен здійснювати первинну класифікацію матеріалів, перевіряти повноту доказової бази та пропонувати релевантні прецеденти. Проте критично важливо розуміти, що такі системи виступають лише цифровим інструментом підтримки, а не замінюють суддю.

У нотаріальній діяльності інтеграція ШІ-асистентів здатна кардинально модернізувати класичні підходи до верифікації документів, посвідчення угод та здійснення превентивного правового контролю. Сучасні алгоритми спроможні забезпечувати миттєвий автоматизований комплаєнс-контроль шляхом крос-системного зіставлення даних із десятків взаємопов'язаних державних реєстрів, що дозволяє в режимі реального часу виявляти приховані аномалії чи підозрілі патерни в історіях транзакцій нерухомості або корпоративних прав. Крім того, інтелектуальні інструменти аналізу тексту з високою точністю ідентифікують ознаки матеріального чи інтелектуального підроблення документів, а також оптимізують процес автоматичного проектування складних індивідуалізованих договорів під конкретні потреби сторін. Таке комплексне технологічне сприяння суттєво мінімізує вплив людського фактора, ефективно нівелює ризики шахрайських схем і виводить загальний рівень правової безпеки нотаріальних дій на принципово новий ступінь надійності.

Окремим, глибоко інноваційним вектором цифровізації виступає інтеграція систем штучного інтелекту в спеціалізовану практику психологів, що набуває особливої ваги у сферах юридичної психології, інструментальної медіації та криміналістичного профайлінгу. Сучасні мультимодальні алгоритми аналізу невербальної поведінки та природного мовлення наділені здатністю з високою точністю фіксувати найменші мікроекспресії обличчя, атипові зміни тембру і частоти голосу, а також приховані семантичні патерни безпосередньо під час проведення процесуальних допитів чи комплексних судово-психологічних експертиз. Такий технологічний підхід озброює фахівців науково обґрунтованим, позбавленим когнітивних упереджень інструментарієм для релевантної оцінки поточного емоційного стану особи, визначення критичних маркерів стресу чи детекції психологічних індикаторів, що вказують на рівень достовірності наданих свідчень. Поряд із цим, використання спеціалізованих діалогових ШІ-систем (чатботів) створює надійне підґрунтя для ефективного та масштабованого проведення первинного психологічного скринінгу, що суттєво оптимізує ресурси фахівця на початкових етапах взаємодії.

Незважаючи на беззаперечний інноваційний потенціал, імплементація систем штучного інтелекту у професійну практику неминуче супроводжується комплексом серйозних викликів. Фундаментальною етичною та правовою загрозою залишається ризик виникнення

«алгоритмічної упередженості» (algorithmic bias), за якої навчена нейромережа здатна мимоволі реплікувати та навіть масштабувати системні помилки, дискримінаційні патерни чи суб'єктивні девіації, що були приховані у масивах історичних рішень та емпіричних даних. Оптимальним шляхом нівелювання цієї критичної вразливості є безальтернативне впровадження концепції «прозорості алгоритму» (Explainable AI – XAI). Цей підхід вимагає демістифікації процесів машинного навчання, гарантуючи, що інтелектуальна система не обмежується видачею безапеляційного результату, а надає фахівцю розгорнуту можливість чітко простежити, інтерпретувати та критично верифікувати внутрішню аналітичну логіку формування кожного машинного висновку.

Крім того, інтеграція інтелектуальних систем у чутливі сфери професійної діяльності висуває безпрецедентні вимоги до архітектури кібербезпеки, оскільки потенційні витoki інформації здатні завдати непоправної шкоди фундаментальним правам громадян та дискредитувати довіру до цифрових сервісів. Обробка масивів надзвичайно чутливої інформації, яка становить охоронювану законом адвокатську, нотаріальну чи лікарську (зокрема, психологічну) таємницю, зумовлює безальтернативну необхідність застосування сучасних криптографічних протоколів наскрізного шифрування на всіх етапах життєвого циклу даних. Це означає, що доступ до змісту завантажених матеріалів, результатів психологічного скринінгу чи проєктів нотаріальних актів повинен бути технологічно унеможливлений не лише для зовнішніх суб'єктів, а й для самих розробників чи адміністраторів серверів. Водночас, проєктування таких платформ має ґрунтуватися на концепції вбудованої приватності («Privacy by Design») та передбачати неухильне дотримання жорстких вимог Загального регламенту про захист даних (GDPR). Йдеться, насамперед, про обов'язкову імплементацію механізмів суворої мінімізації даних – коли система збирає та аналізує лише ту інформацію, яка критично необхідна для виконання поточного завдання, – а також застосування алгоритмів незворотної анонімізації або псевдонімізації особистих відомостей одразу після завершення сесії користувача.

У підсумку, цифрова трансформація професійної діяльності є незворотним еволюційним процесом. ШІ-асистенти не покликані замінити фахівців; їхня мета – стати високоефективним середовищем для ухвалення зважених рішень. Домінантою у цьому процесі має залишатися принцип «людина в циклі» (Human-in-the-loop), який гарантує, що остаточне правове, нотаріальне чи психологічне рішення, а також персональна відповідальність за нього, завжди залишаються прерогативою людини-фахівця.

Список використаних джерел

1. Баранов О. А. Право та штучний інтелект : монографія. Київ : Видавничий дім «АртЕк», 2020. 520 с.
2. Блажівська О. Є. Цифровізація правосуддя як запорука ефективності судового захисту. Науковий вісник НАВС. 2023. № 2. С. 45–54.
3. Гетманцев О. Д. LegalTech: інноваційні технології в юридичній практиці. Київ : Право, 2022. 312 с.
4. Susskind R. Online Courts and the Future of Justice. Oxford University Press, 2019. 368 p.

Іван ПИСАР, студент н. гр. 104 СПД ННІ права та психології, Національна академія внутрішніх справ.

Науковий керівник:

Валерій ХАХАНОВСЬКИЙ, професор кафедри інформаційних технологій ННІ права та психології, Національна академія внутрішніх справ, доктор юридичних наук, професор

ПСИХОСОЦІАЛЬНІ ТА ЕТИЧНІ РИЗИКИ ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ШІ

Сучасний етап розвитку людської цивілізації позначений всеохопною цифровізацією, у центрі якої перебуває динамічний поступ технологій штучного інтелекту. Цей феномен давно вийшов за межі вузької технічної чи економічної сфери, переростаючи у вагомий рушій змін у духовному, метафізичному та психосоціальному вимірах людського існування. Застосування інтелектуальних систем у щоденному житті ставить перед людством фундаментальні світоглядні виклики, що зачіпають основи релігійних вірувань та традиційну етичну парадигму.

У центрі дискусії опиняються такі питання, як переосмислення унікальності людини, природи свідомості, ідеї свободи волі та прийнятності створення квазірозумних сутностей. Водночас стрімкість науково-технічного прогресу значно перевершує здатність релігійних і філософських інститутів розробляти адаптивні моральні принципи, що призводить до глибокої кризи ідентичності. Метою цього дослідження є аналіз ризиків, пов'язаних із розвитком штучного інтелекту, з акцентом на етичні переконання, духовний добробут і моральні основи сучасного суспільства.

Етичні девіації та детермінанти алгоритмічних дилем.

Методологічні основи проектування та навчання сучасних інтелектуальних систем апіорі містять суттєві етичні виклики. Зважаючи на відсутність у штучного інтелекту автентичної суб'єктності, такі системи лише дублюють і ретранслюють соціокультурні матриці та аксіологічні обмеження, закладені у первинних масивах даних розробниками.

Серед домінантних загроз у цій сфері варто виділити системні викривлення алгоритмів (*Algorithmic Bias*), що виникають внаслідок реплікації нейромережами емпіричних помилок, людських стереотипів та історично сформованої соціальної асиметрії. Впровадження подібного інструментарію у сфері автоматизованого рекрутингу чи фінансового скорингу масштабує дискримінаційні практики під маскою безсторонніх математичних обчислень. Паралельно розгортається перманентна криза конфіденційності, спричинена тотальним моніторингом цифрової активності користувачів, що трансформує традиційну архітектуру приватності та створює інструменти «наглядового капіталізму» для маніпулятивного контролю соціуму з боку владних інституцій через технології розпізнавання облич. Цю ситуацію ускладнює ефект «чорної скриньки» (*Black Box Problem*), при якому внутрішня логіка глибокого навчання залишається неінтерпретованою для оператора. У випадку критичних збоїв алгоритмів у медичній, юридичній чи транспортній практиці це провокує дефіцит правової підзвітності та руйнує суспільний договір безпеки.

Психосоціальні деформи в умовах антропо-техногенної взаємодії.

Експлуатація інтелектуальних агентів чинить латентний деструктивний вплив на ментальну архітектоніку особистості, спричиняючи незворотні трансформації її когнітивного, емоційного й комунікативного модусів.

Спроможність ШІ репродукувати складні креативні та аналітичні продукти провокує поширення техногенної тривожності (AI Anxiety), яка корелює з хронічним дистресом через ризик втрати професійної актуальності, девальвацію інтелектуальної праці та деструкцію індивідуальної ідентичності.

Водночас перманентне делегування базових мисленевих операцій (пам'яті, планування, аналізу, селекції інформації) автоматизованим помічникам мінімізує ментальні зусилля користувача. У довгостроковій перспективі це детермінує когнітивне розвантаження, нівелювання навичок критичного сприйняття та ризик інтелектуальної атрофії (феномен «цифрової деменції»).

Додатковим дестабілізуючим чинником виступає антропоморфізація неживих інтерфейсів: штучна імітація емпатійного відгуку стимулює користувачів наділяти програмний код людськими атрибутами. Стрімке поширення ШІ-компаньйонів формує ілюзію подолання самотності, проте насправді підриває автентичні міжособистісні зв'язки, занурюючи індивіда в стан емоційної самоізоляції та соціальної дезадаптації.

Маніпулювання суспільною свідомістю в умовах цифрової епохи, поряд із загостренням епістемічної кризи, становить одну з найбільш серйозних загроз сучасного світу. Стрімкий розвиток і використання технологій для створення гіперреалістичних діпфейків, які здатні достовірно імітувати реальність, у поєднанні з автоматизованим поширенням таргетованої дезінформації через високотехнологічні рекомендаційні алгоритми, суттєво ускладнюють процес перевірки достовірності інформації. У таких умовах суспільство занурюється в атмосферу фундаментальної епістемічної кризи, що послаблює здатність індивідів ефективно розрізняти реальність від вигадки. Це сприяє ескалації тривожності, формуванню відчуття тотальної недовіри та параної, а також руйнує авторитет ключових соціальних інституцій, включаючи засоби масової інформації, наукове співтовариство та державні органи влади. З огляду на масштаб зазначених викликів, необхідне запровадження системного підходу до їх подолання, що орієнтуватиметься на концепцію людиноцентричного штучного інтелекту (Human-Centric AI). Дана концепція передбачає інтеграцію технологій штучного інтелекту з метою зміцнення людських цінностей та забезпечення соціальної стабільності. Основними напрямками стратегічних дій у межах цієї моделі є: *Нормативно-правове регулювання*. Важливим є впровадження жорстких стандартів регулювання штучного інтелекту, подібних до регуляторних рамок EU AI Act. Таке регулювання має включати прозорість алгоритмів завдяки технологіям Explainable AI, а також забезпечення контролю людини в прийнятті критично важливих рішень через підхід Human-in-the-loop. Це дозволить залишити останнє слово за людиною, що зменшить ризики автоматизованих помилок і мінімізує потенційну шкоду для суспільства. - *Освіта та розвиток компетенцій*.

Вкрай актуальним є формування у людей навичок, які ще не можуть бути відтворені за допомогою повної автоматизації. До них належать емоційний інтелект, критичне мислення, морально-етична свідомість та базова медіаграмотність. Інвестування у відповідні освітні програми сприятиме підготовці населення до викликів цифрового суспільства, забезпечуючи компетенції для критичної оцінки інформації та успішної взаємодії зі штучним інтелектом. - *Соціально-психологічна підтримка*. Усвідомлюючи потенційний вплив роботизації та автоматизації на ринок праці та соціальні структури, варто запроваджувати програми професійної переорієнтації для тих, хто може втратити роботу через поширення ШІ-технологій. У рамках таких ініціатив необхідно забезпечити не тільки можливості для нової кар'єри, але й активно розвивати системи психологічної допомоги та підтримувати баланс між цифровим і реальним життям. Реалізація цих стратегій дозволить сформувати більш стійке й адаптивне суспільство, у якому технологічний прогрес гармонійно співіснуватиме з основними інтересами і цінностями людства. Це створить передумови для ефективного подолання епістемічних криз цифрової доби та значно знизить ризики, пов'язані з маніпуляцією колективною свідомістю.

Висновки: психосоціальні та етичні ризики використання ШІ є серйозними викликами, що зачіпають сутність людської природи та психічного благополуччя. Вирішення цих проблем потребує міждисциплінарного консенсусу розробників, психологів, юристів та філософів. Головним імперативом має стати підпорядкування технологічного прогресу критеріям гуманізму та збереження психічного здоров'я особистості.

Список використаних джерел

1. Бостром Н. Штучний інтелект. Надії та загрози створеної розумної машини; пер. з англ. Р. Корнута. Київ: Наш Формат, 2020. 400 с.
2. Зубофф Ш. Епоха наглядного капіталізму. Битва за людське майбутнє на нових рубежах влади. Київ: Наш Формат, 2022. 600 с.
3. Шваб К. Четверта промислова революція. Київ: Форс Україна, 2019. 232 с.
4. Russell S. Human Compatible: Artificial Intelligence and the Problem of Control. New York: Viking, 2019. 352 p.
5. Floridi L. Establishing the rules for an ethical AI // Nature Machine Intelligence. 2019. Vol. 1, No. 6. P. 261–262.

Анастасія НАСАДЮК, студентка спеціальності 035 Філологія ОПП «Германські мови та літератури (переклад включно), перша – англійська», Івано-Франківський національний технічний університет нафти і газу.

Науковий керівник:

Наталія КОБЗЕЙ, професорка кафедри філології та перекладу, Івано-Франківський національний технічний університет нафти і газу, кандидат філологічних наук, доцент

ЕТИЧНІ АСПЕКТИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПІДГОТОВЦІ ФІЛОЛОГІВ-ПЕРЕКЛАДАЧІВ

Стрімка інтеграція технологій штучного інтелекту в освітній простір закладів вищої освіти зумовлює необхідність переосмислення усталених підходів до фахової підготовки фахівців гуманітарного профілю. Філологічні спеціальності, зокрема підготовка перекладачів, опиняються в епіцентрі цих змін, оскільки саме в цій галузі людська інтерпретація, культурна компетентність і творче мислення традиційно становлять ядро професійної діяльності. Поширення великих мовних моделей (LLM), систем нейронного машинного перекладу та інтелектуальних асистентів актуалізує комплекс етичних питань, системне вивчення яких є нагальною науковою потребою.

У зв'язку з цим важливою стає проблема етичного використання технологій штучного інтелекту в процесі професійної підготовки майбутніх філологів-перекладачів. Актуальність дослідження зумовлена необхідністю осмислення впливу генеративного ШІ на формування перекладацьких компетентностей, дотримання принципів академічної доброчесності, збереження авторського права та визначення меж професійної відповідальності в умовах цифровізації освітнього простору.

Метою роботи є аналіз основних етичних аспектів використання технологій штучного інтелекту в процесі підготовки філологів-перекладачів, а також визначення потенційних ризиків і перспектив їх інтеграції в освітній процес закладів вищої освіти.

У межах дослідження було використано методи аналізу наукової літератури, узагальнення сучасних підходів до використання ШІ в освіті, а також метод систематизації основних етичних викликів, пов'язаних із застосуванням генеративних мовних моделей у перекладацькій діяльності.

Сучасні системи автоматизованого перекладу – DeepL, Google Translate, ChatGPT – демонструють технічно вражаючі показники опрацювання текстів. Водночас дослідники наголошують, що використання генеративного штучного інтелекту може бути ефективним насамперед за умови критичного та професійно виваженого підходу з боку перекладача. Зокрема, О. Остапович, Н. Остапович та Ю. Мазуренко зазначають, що «коректно сформульований запит до ChatGPT допомагає професійному перекладачу оперативно редагувати за відповідними контекстними базами вузькоспеціалізовані, термінологічно насичені тексти, перекладені ним іноземною мовою або з іноземної мови на рідну» [1, с. 201]. Це свідчить про те, що сучасні ШІ-інструменти доцільно розглядати не як повноцінну заміну перекладача, а як допоміжний інструмент, ефективність якого безпосередньо залежить від рівня фахової компетентності користувача.

Утім, некомпетентне використання таких інструментів до навчального процесу супроводжується низкою системних ризиків: розмиванням меж академічної доброчесності, послабленням внутрішньої мотивації студентів до глибинного засвоєння мовного матеріалу, а також формуванням стійкої залежності від автоматизованих рішень на шкоду розвитку власної перекладацької компетентності. У контексті підготовки філологів-перекладачів доцільно виокремити кілька ключових етичних проблем використання штучного інтелекту.

Насамперед гостро постає проблема академічної доброчесності. Подання студентами машинного перекладу як результату власної інтелектуальної праці без осмисленого редагування та критичного аналізу тексту є проявом академічної нечесності. Таке явище суперечить фундаментальним принципам університетської освіти та унеможливорює повноцінне формування перекладацької компетентності.

Не менш актуальним є питання авторського права та атрибуції. Генеративні ШІ-моделі проходять навчання на великих масивах даних, які нерідко містять матеріали, захищені авторським правом. У зв'язку з цим питання правової належності перекладеного або синтезованого тексту – користувачу, алгоритмічній системі чи її розробнику – залишається юридично й етично неврегульованим у більшості сучасних правових систем. Аналізуючи цю проблему, Н. Добрянська зазначає, що чинне законодавство України «не містить системного та деталізованого регулювання використання технологій штучного інтелекту в освітньому процесі», а також залишає відкритими питання критеріїв допустимості та процедур відповідальності за порушення академічної доброчесності [2, с. 129]. Таким чином, стрімке поширення генеративного штучного інтелекту актуалізує потребу вдосконалення нормативно-правового регулювання у сфері освіти.

Окрему увагу слід звернути на проблему можливого знецінення перекладацьких компетентностей. Надмірна й некритична орієнтація на ШІ-переклад створює загрозу поступової деградації фахових навичок: стилістичного чуття, прагматичної адаптації тексту, роботи з культурно маркованими смислами та конотаціями. У сучасних умовах перекладач має формуватися передусім як критичний редактор і культурний посередник, а не як технічний оператор автоматизованих систем.

Крім того, в умовах зростаючої автоматизації особливої актуальності набуває проблема відповідальності за якість перекладу. Якщо фахівець лише формально верифікує машинний варіант без поглибленого аналізу, юридична та етична відповідальність за можливі помилки – особливо у медичній, юридичній чи дипломатичній сферах – стає невизначеною. Це ще раз підтверджує, що використання систем штучного інтелекту у перекладацькій діяльності потребує не лише технічної підготовки, а й високого рівня професійної та етичної відповідальності.

З метою системного врегулювання окреслених проблем доцільним видається запровадження інституційних етичних рамок використання ШІ в освітньому процесі. О. Пелешок влучно підкреслює: «Етична відповідальність закладів вищої освіти передбачає не лише створення інституційних механізмів регулювання використання ШІ, але й формування у здобувачів освіти та викладачів навичок критичного мислення, здатності до самостійної інтелектуальної праці та дотримання принципів академічної доброчесності в

цифрову епоху. Особливої уваги потребує баланс між інноваційністю та етичністю, тобто питання, де закінчується допустиме використання ШІ як допоміжного інструменту і починається порушення норм академічної поведінки» [3, с. 44]. Це передбачає розроблення чітких університетських політик щодо допустимих форм застосування ШІ-інструментів, цілеспрямоване формування критичного ставлення студентів до якості машинного перекладу, а також введення окремої навчальної дисципліни «Етика штучного інтелекту в перекладі» до навчальних планів відповідних спеціальностей.

Результати дослідження дають підстави стверджувати, що використання штучного інтелекту у підготовці філологів-перекладачів має суперечливий характер: з одного боку, ШІ значно розширює можливості освітнього процесу, оптимізує роботу з текстами та сприяє розвитку цифрової компетентності студентів, а з іншого – створює низку етичних ризиків, пов'язаних із порушенням академічної доброчесності, знеціненням фахових навичок та невизначеністю відповідальності за результати перекладу. Встановлено, що ефективне використання ШІ в перекладацькій освіті можливе лише за умов формування у студентів критичного мислення, навичок професійної верифікації машинного перекладу та чіткого нормативно-етичного регулювання застосування таких технологій.

Отже, штучний інтелект становить потужний дидактичний ресурс, здатний суттєво розширити можливості фахової підготовки перекладачів – за умови його відповідального, методично обґрунтованого та етично вираженого застосування. Ключове завдання вищої школи полягає не в обмеженні доступу до технологій, а у формуванні у майбутніх фахівців здатності взаємодіяти з ними усвідомлено – зберігаючи пріоритет людського судження, культурної чутливості та професійної відповідальності.

Перспективи подальших досліджень вбачаємо у розробленні методичних моделей інтеграції ШІ-технологій у професійну підготовку перекладачів, а також у вивченні впливу генеративного штучного інтелекту на трансформацію перекладацької професії та формування нових компетентностей майбутніх фахівців.

Список використаних джерел

1. Остапович О., Остапович Н., Мазуренко Ю. ChatGPT у підготовці філологів і перекладачів. Виклики і перспективи. *Наукові записки Національного університету «Острозька академія»: серія «Філологія»*. 2023. № 17 (85). С. 200–205. URL: <https://journals.oa.edu.ua/Philology/article/view/3842> (дата звернення: 20.05.2026)
2. Добрянська Н. Використання штучного інтелекту в освітньому процесі: правові межі та академічна доброчесність. *Аналітично-порівняльне правознавство*. 2026. № 2. С. 125–130. URL: <https://doi.org/10.24144/2788-6018.2026.02.2.20> (дата звернення: 20.05.2026)
3. Пелешок О. Штучний інтелект як виклик академічній доброчесності: етичний вимір освітньої відповідальності. *Вісник Глухівського національного педагогічного університету імені Олександра Довженка. Серія: «Педагогічні науки»*. 2025. №2 (58). С. 41–48. URL: <https://doi.org/10.31376/2410-0897-2025-2-58-41-48> (дата звернення: 20.05.2026)

Тези стендових доповідей учасників конференції

Володимир ПАЛИВОДА, завідувач відділу державної та громадської безпеки центру безпекових досліджень, Національний інститут стратегічних досліджень

ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ДІЯЛЬНІСТЬ СПЕЦІАЛЬНИХ СЛУЖБ США: ОСНОВНІ НАПРЯМИ ТА ВИКЛИКИ

Стрімкий розвиток технологій штучного інтелекту (ШІ) перетворює його на ключовий ресурс конкурентоспроможності розвідувальних служб провідних держав світу. Центральне розвідувальне управління (ЦРУ) США одним із перших серед іноземних спецслужб офіційно оголосило про системний перехід до моделі роботи, в якій інструменти ШІ виконують функції «колег» аналітиків, а в перспективі — «автономних партнерів» оперативних команд. Відповідне бачення публічно представив заступник директора ЦРУ Майкл Елліс на заході Special Competitive Studies Project (SCSP) 9 квітня 2026 р. [1; 2]. Аналіз цього досвіду має методологічну цінність для розуміння траєкторії технологічної трансформації спецслужб у середньостроковій перспективі.

ЦРУ планує найближчими роками інтегрувати інструменти ШІ безпосередньо в аналітичні платформи відомства. Такі інструменти не замінюватимуть аналітиків, а допомагатимуть їм складати ключові висновки, редагувати проекти документів і звіряти їх зі стандартами розвідувальної діяльності. Системи ШІ виконуватимуть також функції сортування та маркування тенденцій для подальшого опрацювання аналітиками [1]. Це окреслює модель «допоміжного» застосування ШІ, в якій принциповим залишається людський контроль над аналітичними висновками.

Стратегічна мета ЦРУ на десятилітню перспективу — перехід до використання ШІ як автономного партнера по розвідувальній місії. Передбачається, що співробітники керуватимуть командами агентів ШІ за гібридною моделлю, що дозволить істотно збільшити швидкість і масштаб розвідувальної діяльності [1; 3]. Це означає поступовий перехід від сприйняття ШІ як інструменту до його розгляду як суб'єкта оперативного процесу під керівництвом людини, що ставить нові вимоги до підготовки кадрів та формалізації процедур контролю.

У 2025 р. в ЦРУ реалізовувалося понад 300 проєктів у сфері ШІ; уперше в історії відомства за допомогою ШІ було підготовлено повноцінний розвідувальний звіт [1; 4]. Окремим напрямом є кіберрозвідка: Центр кіберрозвідки ЦРУ перетворено на повноцінний орган, що забезпечує нові інструменти на місцях та ширший доступ до пріоритетних цілей. Сформульовано прогноз — «битва за кібербезпеку буде битвою штучного інтелекту», а перевагу матиме той, хто найкраще використовуватиме ШІ-моделі [1]. Паралельно агентство впровадило нову систему закупівель, спрямовану на пришвидшену інтеграцію передових технологій, що свідчить про сприйняття ШІ як підстави для перебудови внутрішніх процесів, а не лише точкового технологічного оновлення [1].

ЦРУ подвоїло обсяг звітів зовнішньої розвідки, що стосуються технологій, для відстеження того, як іноземні противники (насамперед Китай) використовують передовий ШІ та інші технології подвійного призначення. Тематичний фокус — напівпровідники, хмарні обчислення, інфраструктура, кібербезпека, дослідження та розробки [1]. Цей напрям відображає інституційну реакцію на ризики технологічного відставання та витоку компетенцій із критичних сфер.

Серйозним викликом залишається залежність розвідувального співтовариства від обмеженого кола розробників передових моделей. Ілюстративним є кейс компанії Anthropic: на початку 2026 р. компанія відмовилася послабити обмеження, які забороняють використання

її інструментів для внутрішнього спостереження та повністю автономної зброї. Це спричинило кваліфікацію її продуктів міністром оборони США П. Хегсетом як «ризик ланцюга поставок» і розпорядження адміністрації Д. Трампа про поступове припинення використання інструментів Anthropic федеральними відомствами [1]. Водночас Anthropic анонсувала проєкт Glasswing — закритий консорціум для захисту критично важливого програмного забезпечення від атак на основі ШІ, побудований на моделі, що вже виявила тисячі вразливостей у поширених системах [5; 6]. М. Елліс на цей кейс безпосередньо не посилався, однак заявив, що ЦРУ «не може дозволити примхам однієї компанії» обмежувати використання ШІ, і прагне диверсифікувати постачальників задля збереження оперативної свободи [1].

Висновки.

По-перше, ЦРУ закріплює ШІ як штатний елемент аналітичного та оперативного процесу, рухаючись від окремих інструментів до інтегрованих платформ і керованих людиною команд ШІ-агентів.

По-друге, технологічна розвідка та кіберрозвідка визначені як пріоритетні сфери ШІ-трансформації з очевидним фокусом на стримуванні КНР.

По-третє, ключовим інституційним викликом стає правова, етична та логістична залежність розвідувального співтовариства від комерційних розробників передових моделей, що підштовхує до політики диверсифікації постачальників та пошуку альтернатив.

Досвід США демонструє, що інтеграція ШІ у діяльність спецслужб є не одномоментним проєктом, а структурною трансформацією, яка потребує системного перегляду засад розвідувальної діяльності, моделей закупівлі технологій і взаємодії з приватним сектором.

Список використаних джерел

1. CIA employees will get AI «coworkers» — and eventually run teams of AI agents, deputy says. Defense One: вебсайт. 09.04.2026. URL: <https://www.defenseone.com/technology/2026/04/cia-ai-coworkers-agents/412746/> (дата звернення: 16.05.2026).
2. Michael Ellis Sworn in as CIA Deputy Director. Central Intelligence Agency: офіційний вебсайт. URL: <https://www.cia.gov/stories/story/michael-ellis-sworn-in-as-cia-deputy-director/> (дата звернення: 16.05.2026).
3. CIA plans for «AI coworkers», deputy director says. Nextgov/FCW: вебсайт. 09.04.2026. URL: <https://www.nextgov.com/artificial-intelligence/2026/04/cia-plans-ai-coworkers-deputy-director-says/412744/> (дата звернення: 16.05.2026).
4. The CIA Let AI Write Its First Intelligence Report — And AI «Coworkers» Are Up Next. Gate News: вебсайт. 2026. URL: <https://www.gate.com/news/detail/the-cia-let-ai-write-its-first-intelligence-reportand-ai-coworkers-are-up-20258816> (дата звернення: 16.05.2026).
5. Anthropic Launches Project Glasswing for Cybersecurity. Via Satellite: вебсайт. 12.04.2026. URL: <https://www.satellitetoday.com/cybersecurity/2026/04/12/anthropic-launches-project-glasswing-for-cybersecurity/> (дата звернення: 16.05.2026).
6. Project Glasswing: Securing critical software for the AI era. Anthropic: офіційний вебсайт. URL: <https://www.anthropic.com/glasswing> (дата звернення: 16.05.2026).

Оксана ЯРЕМА, завідувач кафедри адміністративно-правових дисциплін ННІ права та правоохоронної діяльності, Львівський державний університет внутрішніх справ, кандидат юридичних наук, доцент

РОЗВИТОК ПРАВОВОГО РЕГУЛЮВАННЯ ВІДНОСИН ІЗ РОЗРОБКИ ТА ВПРОВАДЖЕННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ

Оприлюднено проєкт Стратегії розвитку штучного інтелекту до 2030 року, розроблений спільно з міжнародними партнерами, державними установами, бізнесом та спільнотою фахівців зі штучного інтелекту [1]. Документ передбачає розвиток відносин щодо розробки та впровадження систем штучного інтелекту. Стратегія спрямована на стимулювання та підтримку галузі, розробок, виробників і розробників, наукових досліджень, на фінансування.

Проєкт Стратегії визначає, що штучний інтелект (ШІ) – організована сукупність інформаційних технологій, із застосуванням якої можливо виконувати складні комплексні завдання шляхом використання системи наукових методів досліджень і алгоритмів обробки інформації, отриманої або самостійно створеної під час роботи, а також створювати та використовувати власні бази знань, моделі прийняття рішень, алгоритми роботи з інформацією та визначати способи досягнення поставлених завдань [2]. Метою цієї Стратегії є підвищення стійкості держави, прискорення відновлення та забезпечення довгострокового розвитку національного потенціалу шляхом поєднання швидкого та відповідального впровадження штучного інтелекту у пріоритетних сферах із послідовними інвестиціями у розвиток інфраструктури, навичок, інституційних спроможностей та системи управління. Штучний інтелект є технологією загального призначення, ефект від застосування якої досягається за умови її широкого та системного впровадження

Стратегії розвитку штучного інтелекту до 2030 року базується на Концепції розвитку штучного інтелекту в Україні схваленої розпорядженням Кабінету Міністрів України від 2 грудня 2020 р. № 1556-р [3].

Концепція розвитку штучного інтелекту в Україні є документом планування розвитку правового регулювання. Йдеться про основні принципи, які визначають напрям розвитку правового регулювання та використання систем штучного інтелекту.

У Концепції визначаються мета та завдання правового регулювання відносин у сфері систем штучного інтелекту, принципи, основні проблеми регулювання та підходи до їх вирішення, окремі напрями та сфери, які вимагають уваги у зв'язку з специфікою відносин. Водночас Концепція не містить механізми реалізації, що передбачено у проєкті Стратегії розвитку штучного інтелекту до 2030 року.

Пріоритетною метою правового регулювання відносин у цій сфері в Концепції називають: стимулювання розробки, впровадження та використання таких систем.

У розділі «Мета, принципи та строки реалізації Концепції» наголошено на необхідності регулювання нових відносин, визначено проблеми, що ускладнюють розвиток напряму діяльності у сфері інформаційних технологій, який забезпечує створення, впровадження та використання технологій штучного інтелекту.

Пріоритетними сферами, в яких реалізуються завдання державної політики розвитку галузі штучного інтелекту, є: освіта і професійне навчання, наука, економіка, кібербезпека, інформаційна безпека, оборона, публічне управління, правове регулювання та етика, правосуддя. Окремі сфери спрямовані на підвищення довіри населення до інтелектуальних систем, інші вирішують правові питання, треті несуть юридичне та технічне регулювання.

Правове регулювання відносин у сфері застосування технологій та систем штучного інтелекту потрібно опрацьовувати не з погляду окремих сфер застосування програмного забезпечення та технологій, а створювати комплексне системне уявлення про таке регулювання.

Модель правового регулювання повинна включати розподіл систем на види і передбачати, що йому буде надано статус суб'єкта правового регулювання, за умови проходження спеціально розроблених тестів на усвідомленість і автономність.

Правове регулювання відносин використання систем штучного інтелекту має максимально посилається на норми зі схожих відносин у сфері інформаційних технологій, оскільки аналізовані системи є частиною інформаційних технологій.

Окрему увагу при розгляді питання варто приділити технічним нормам – стандартам та технічним регламентам і етичним регуляторам, що дозволяє утворювати цілісне уявлення про об'єкти регулювання. На відміну від суто правових норм стандартизація та розробка етичних норм йде, активно видаються міжнародні та національні документи. Усі вони мають рекомендаційний характер, проте використання та відсилання до них при розробці нормативно-правових актів є необхідним.

Документи стандартизації розробляються експертами відповідної галузі, встановлюють межі розробки та впровадження технологій, етичні норми та принципи.

З виникненням нових суспільних відносин, пов'язаних з використанням інформаційних технологій, особливої значущості набуває правове регулювання, що забезпечує стійкий і динамічний розвиток державної політики у сфері розвитку штучного інтелекту. Це може забезпечити комплексне використання різних регуляторів – етичних, технічних, організаційних, правових та інших.

Даний підхід може розглядатися як модельний акт, який можна використовувати при розробці документів стратегічного планування, інших нормативно-правових актів, що регулюють відносини у сфері штучного інтелекту, у роботі в комітетах і комісіях Європейського Союзу з питань штучного інтелекту та робототехніки,

Пропонується модель правового регулювання відносин у галузі розроблення та використання систем штучного інтелекту, що включає: понятійно-категоріальний апарат; спеціальні засади; роль та значення окремих видів інтелектуальних систем, їх застосування в окремих сферах життя з урахуванням функціональних характеристик; загальні вимоги до розробки та правила використання інтелектуальних систем кожного виду; роль та місце окремих регуляторів.

Посилаючись на нормативні документи Європейського Союзу необхідно передбачити питання захисту від загроз і ризиків з можливістю подальшого розвитку технологій: ліцензування діяльності розробників систем штучного інтелекту та сертифікація; встановлення згоди особи, щодо якої системою штучного інтелекту приймається автоматизоване рішення; страхування відповідальності розробника; презумпція відповідальності людини за рішення систем штучного інтелекту; правові гарантії, пільги та стимули для розвитку галузі та інші засоби та механізми правового регулювання.

Особливе місце у регулюванні технологій штучного інтелекту займає технічне регулювання. Механізми створення та застосування технічних норм дозволяють більш оперативно порівняно з правовими нормами реагувати на технологічні виклики та загрози.

Список використаних джерел

1. Проєкт Стратегії розвитку штучного інтелекту до 2030 року. URL: <https://aifiles.thedigital.gov.ua/documents/Проект%20Стратегії%20розвитку%20ШІ%20202030.docx.pdf>
2. Шлях розвитку ШІ до 2030 року: знайомтеся з проєктом Стратегії та давайте фідбек. URL: <https://thedigital.gov.ua/news/shtuchnyy-intelekt/shliakh-rozvytku-shi-do-2030-roku-znayomtesia-z-proyektom-stratehiyi-ta-davayte-fidbek>
3. Про схвалення Концепції розвитку штучного інтелекту в Україні: Розпорядження Кабінету Міністрів України від 02.12.2020 р. № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text>

Едуард РИЖКОВ, професор кафедри інформаційних технологій, Дніпровський державний університет внутрішніх справ, кандидат юридичних наук, професор

АРХІТЕКТУРА МЕТОДИКИ «OSINT+» У КОНТЕКСТІ ВПРОВАДЖЕННЯ ШТУЧНОГО ІНТЕЛЕКТУ

Сучасна архітектура національної безпеки перебуває на етапі безпрецедентної технологічної сингулярності. В умовах тотальної цифровізації, експоненційного зростання обсягів неструктурованих даних (Big Data) та переходу транснаціональної злочинності у площину Web3.0, традиційні методи оперативно-розшукової діяльності виявляють свою гносеологічну та функціональну обмеженість. Як свідчить практика, сьогодні понад 80% критично важливої для слідства інформації генерується та зберігається поза межами закритих відомчих баз даних — у соціальних мережах, месенджерах, реєстрах та криптовалютних транзакціях [1].

Ще у 2008 році розробники перших потужних інтегрованих систем МВС у Луганській області (ІАС «АРМОР», «СОВА», «АРГУС») зазначали, що розрізненість баз даних та неможливість простежити взаємозв'язки між об'єктами обліку є головним бар'єром ефективності правоохоронців [2, с. 71]. Сьогодні цей бар'єр пролягає не між внутрішніми підрозділами, а між відомчим контуром (Інформаційний портал Національної поліції України — ІПНП) та глобальним інформаційним простором (OSINT).

Вирішення цієї проблеми лежить у площині впровадження авторської методики «OSINT+», яка базується на синергії закритих державних реєстрів, автоматизованих засобів розвідки відкритих джерел та моделей штучного інтелекту (ШІ) [3].

Деконструюючи досвід створення луганського інтегрованого банку даних, варто виокремити фундаментальні рішення, які випередили свій час і сьогодні є базисом для побудови архітектури «OSINT+».

По-перше, це відмова від жорсткої ієрархії на користь зв'язків типу «багато-до-багатьох» (Схема «Особа-Адреса-Подія») [2, с. 78]. По-друге, це запровадження візуального та графового аналізу в системі «СОВА», що дозволило автоматизувати виявлення логічних закономірностей зв'язків об'єктів (наприклад, аналіз телефонних трафіків ОЗГ) [2, с. 61, 100]. По-третє, спроба алгоритмізувати біометричну ідентифікацію через підсистему «АРГУС» (переведення «словесного портрета» у базу даних) [2, с. 116].

Проте ці системи були замкнутими. Вони чекали, поки оперативний працівник внесе інформацію вручну. З позиції сучасного архітектурного проектування, ІПНП сьогодні має стати не просто сховищем, а активним аналітичним ядром, де відомчі дані виступають лише початковим «зерном» (Seed) для автоматичного розгортання графа зв'язків у відкритому середовищі.

За нашою позицією, методика «OSINT+» — це гібридна епістемологічна модель та технологічний фреймворк, що передбачає автоматизоване двостороннє збагачення даних: від закритого контуру (ІПНП) до відкритих джерел (OSINT), із подальшою їх верифікацією та структуруванням за допомогою технологій штучного інтелекту (AI Processing) [4].

Розвиваючи концепцію «OSINT+» до вимог сьогодення, її архітектура, на наш погляд, має складатися з трьох автоматизованих рівнів:

1. Ініціація (Seed Level), яка реалізується через запит слідчого або оперативника в ІПНП (ПІБ, номер телефону, ПІН, фотографія, криптогаманець).

2. Автоматизований збір (Collection Engine) інформації орієнтуючого характеру реалізується через запуск оркестратора скриптів (на базі Python, API-інтерфейсів), який без участі людини опитує соціальні мережі, реєстри (YouControl, Опендатабот), витіки баз даних (Data Leaks) та спеціалізовані пошукові системи (Shodan, Censys).

3. III-Процесинг (AI Verification Cycle) запускається через використання Large Language Models (LLM) для вилучення сутностей з текстів (Named Entity Recognition) та алгоритмів комп'ютерного зору для розпізнавання обличчя і атрибутів на зібраних медіафайлах.

Як варіант на підставі створеного концепту модернізації, ми пропонуємо можливі практичні кроки інтеграції за пріоритетними напрямками діяльності поліції:

1. Впровадження модулю автоматичного збагачення сутності «Особа» (SOCMINT & PII Automation). При створенні або запиті картки особи в ПНП (наявність номеру телефону чи email), система повинна автоматично звертатися до відкритих OSINT-фреймворків (аналогів SpiderFoot [4], Holehe, Maigret). Це дозволить миттєво підтягувати до досьє особи її цифрові сліди: акаунти в Telegram, Instagram, історію зміни аватарів, геотеги та пов'язані IP-адреси. Замість ручного пошуку, працівник поліції отримує готовий граф цифрового профілю.

2. Еволюція інформаційної підсистеми «Гарпун» та відеоаналітики через Computer Vision. Перехід від класичного ручного «словесного портрета» до III-біометрії. Інтеграція загальнодержавної системи відеомоніторингу [5] з підсистемами ПНП має базуватися на нейромережових ембедингах (FaceNet, YOLO). Оператор (ситуаційний аналітик) повинен мати змогу формувати запит не лише за номером авто чи обличчям, а за семантичним описом (наприклад, «особа з червоним рюкзаком, що рухалась у напрямку...»). Система «OSINT+» автоматично порівнюватиме ці відеодані з відкритими масивами фотографій у соціальних мережах для встановлення особи (аналоги PimEyes/FindClone).

3. Протидія кримінальній токенизації (FININT/Crypto-OSINT). У відповідь на перехід ОЗГ до розрахунків у криптовалютах [1], ПНП має бути доповнений модулем On-chain аналітики («Follow the Code»). Методика «OSINT+» передбачає автоматичне зіставлення адрес криптогаманців, виявлених під час обшуків чи в даркнеті, з OSINT-маркерами (згадки гаманців на хакерських форумах, у Telegram-каналах) для деанонізації кінцевого бенефіціара та зв'язку його з легальними рахунками підприємств через реєстри.

4. Впровадження локальних (On-Premise) LLM-моделей для обробки неструктурованих даних. Сьогодні оперативники та слідчі генерують колосальний обсяг неструктурованих текстів (рапорти, протоколи допитів, фабули Журналу реєстрації заяв (повідомлень) про злочини та інш.). Пропонується інтегрувати в ПНП локальну «поліцейську» нейромережу для автоматичного аналізу цих текстів. III (через NLP) самостійно вилучатиме імена, адреси, номери телефонів та автоматично будуватиме неочевидні зв'язки між різними кримінальними провадженнями, позбавляючи слідчого необхідності ручного введення реквізитів.

Таким чином, на підставі вищезазначеного сформулюємо наступні висновки:

1. Використання інформаційних технологій у правоохоронній діяльності має еволюціонувати від концепції «електронної картотеки» до концепції «проактивного предиктивного аналізу» (Predictive Policing), поступово змінюючи суть самої парадигми діяльності поліції. Інформаційний портал Національної поліції України не повинен бути ізольованим архівом; він має стати ядром екосистеми, що безперервно взаємодіє з відкритим інформаційним простором.

2. Наша авторська методика «OSINT+», яка за визначенням поєднує закриті бази (ПНП), автоматизований збір відкритих даних (OSINT) та машинну обробку (III), дозволяє нівелювати «людський фактор», ліквідувати інформаційний шум та скоротити час від збору даних до прийняття процесуального рішення у десятки разів. Таким чином, синергія «OSINT+» поступово повинна сприйматися як новий стандарт у інформаційно-аналітичній діяльності поліції.

3. Впровадження автоматизованих скриптів та III вимагає імплементації протоколів верифікації (на кшталт Протоколу Берклі [6]), щоб алгоритмічні висновки системи («OSINT+») мали належну доказову силу, а зібрані цифрові сліди могли бути легалізовані в рамках вітчизняного кримінального процесу та вдосконалили оперативно-розшукову діяльність.

Список використаних джерел

1. Ryzhkov E., Ohrimenco S. Criminal tokenization in the era of financial market transformation: challenges, strategies, and counteraction tactics / Міжнародна та національна безпека: теоретичні і прикладні аспекти : матеріали X Міжнар. наук.-практ. конф. (м. Дніпро, 16 бер. 2026 р.) ; у 2-х ч. Дніпро : Дніпров. держ. ун-т внутр. справ, 2026. Ч. II. С. 362-365.
2. Гуславский В. С., Задорожний Ю. А., Розовский Б. Г. Информационно-аналитическое обеспечение раскрытия и расследования преступлений : монография. Луганск : Изд-во «Элтон-2», 2008. 136 с.
3. Рижков Е.В. Методика "OSINT+": синергія відкритих та відомчих інформаційних ресурсів у діяльності оперативних підрозділів Національної поліції України / Е.В. Рижков / Роль OSINT-досліджень у підвищенні рівня національної безпеки України : матер. кр. столу (м. Львів, 07 травня 2025 р.). / упорядник: І. О. Ревак. Львів: Львівськ. держ. ун-т внутр. справ, 2025. С. 192-199.
4. Рижков Е.В. Перспективи комплексного використання автоматизованих систем (фреймворків) та моделей штучного інтелекту в процесі OSINT / Е.В. Рижков / Інформаційно-аналітичне забезпечення діяльності органів сектору безпеки і оборони України, наук.-практ. конфер.: матер. Всеукр. науково-практ. конф. (м. Львів, 26 грудн. 2025 р.). / упорядник: Т.В. Магеровська. Львів: Львівськ. держ. ун-т внутр. справ, 2026. – С. 59-62
5. Про затвердження Інструкції з формування та ведення інформаційної підсистеми «Гарпун» інформаційно-комунікаційної системи «Інформаційний портал Національної поліції України» : Наказ МВС України від 13.06.2018 № 497 (зі змінами від 22.11.2023). URL: <https://zakon.rada.gov.ua/laws/show/z0787-18#Text>
6. Протокол Берклі з ведення розслідувань з використанням відкритих цифрових даних. URL: <https://www.law.berkeley.edu/wp-content/uploads/2022/03/Berkeley-Protocol-Ukrainian.pdf>

Ольга ОГІРКО, професор кафедри інформаційних технологій, Львівський державний університет внутрішніх справ, кандидат технічних наук, доцент.

Наталія ГАЛАЙКО, старший викладач кафедри інформаційних технологій, Львівський державний університет внутрішніх справ

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В АДМІНІСТРУВАННІ ІНФОРМАЦІЙНИХ СИСТЕМ НАЦІОНАЛЬНОЇ ПОЛІЦІЇ УКРАЇНИ

Сучасна система глобальної безпеки зазнає суттєвих змін під впливом стрімкого розвитку цифрових технологій. Для України, яка протистоїть повномасштабній агресії, впровадження інноваційних рішень у діяльність правоохоронних органів є важливою умовою забезпечення національної безпеки та ефективного захисту держави. Одним із найбільш перспективних напрямів цифрової трансформації є використання потенціалу штучного інтелекту (ШІ) у безпековому секторі.

У науковій літературі штучний інтелект визначається як сукупність методів і технологій, що дають змогу комп'ютерним системам виконувати окремі функції, властиві людському мисленню [1]. Йдеться про здатність аналізувати інформацію, навчатися на основі отриманих даних, розпізнавати мовлення, прогнозувати події та приймати рішення.

Такі можливості є особливо важливими для правоохоронних органів, оскільки забезпечують оперативне опрацювання великих обсягів інформації та підтримують ефективне функціонування відомчих інформаційних систем.

Одним із ключових напрямів розвитку ШІ є машинне навчання – технологія, що дає змогу системам аналізувати великі масиви даних і виявляти в них певні закономірності [2]. На відміну від традиційних програм, де алгоритми роботи визначаються наперед заданими командами, інтелектуальні системи здатні вдосконалювати свою роботу на основі нових даних. Використання таких технологій створює передумови для модернізації адміністрування інформаційних систем, підвищення швидкості обробки інформації та ефективності реагування на сучасні загрози [3, 4].

Науково-практичні дослідження у сфері застосування штучного інтелекту в управлінні відомчими цифровими платформами пов'язані насамперед із підвищенням ефективності контролю та аналітичної діяльності. Алгоритми машинного навчання дозволяють автоматизувати виконання рутинних операцій, які раніше потребували постійного втручання працівників. Це сприяє зниженню ймовірності технічних помилок, скороченню часу обслуговування інфраструктури та підвищенню стабільності функціонування інформаційних систем.

У межах модернізації інформаційної інфраструктури важливе значення має моніторинг працездатності серверів, баз даних та інтегрованих модулів. Технології ШІ дозволяють контролювати рівень навантаження ресурсів, швидкість обробки запитів, а також виявляти затримки чи технічні несправності. Отримані дані можуть використовуватися для своєчасного інформування адміністраторів з метою попередження аварійних ситуацій та оптимізації роботи інформаційної мережі.

Окрім технічного моніторингу, системи штучного інтелекту можуть застосовуватися для аналізу активності користувачів та виявлення аномальних дій. Спеціалізовані алгоритми здатні фіксувати спроби несанкціонованого доступу, надмірного використання ресурсів або інші порушення встановлених регламентів. Це сприяє мінімізації ризиків компрометації даних і забезпеченню безпеки відомчих інформаційних ресурсів.

Окремим напрямом використання ШІ є автоматичне формування аналітичної звітності. Інтелектуальні системи можуть генерувати звіти про стан серверного обладнання, рівень завантаження програмних модулів, активність користувачів, а також функціонування аналітичних і відеоаналітичних підсистем. Такі дані забезпечують інформаційну підтримку управлінських рішень щодо модернізації технічної інфраструктури та вдосконалення організації роботи.

У сфері оперативно-службової діяльності алгоритми машинного навчання та обробки природної мови розглядаються як засоби аналітичної підтримки. Їх використання дає можливість аналізувати великі масиви даних, встановлювати взаємозв'язки між подіями, особами та об'єктами, а також виявляти приховані закономірності. Такі системи можуть застосовуватися для аналізу історії правопорушень, маршрутів транспортних засобів і зв'язків між фігурантами кримінальних проваджень. Окремі елементи інтелектуального аналізу вже використовуються у системах відеоспостереження та відеоаналітики для відстеження транспортних засобів, що перебувають у розшуку. Водночас інші функції, зокрема автоматичне сортування звернень громадян за пріоритетністю або формування рекомендацій для оперативних підрозділів, залишаються перспективним напрямом розвитку таких технологій.

Використання штучного інтелекту у сфері кібербезпеки Національної поліції України обумовлене необхідністю захисту конфіденційної службової інформації та значних обсягів персональних даних. Системи ШІ можуть застосовуватися для автоматизації процесів виявлення кіберзагроз і підвищення швидкості реагування на інциденти. Алгоритми здатні аналізувати великі обсяги системних журналів і мережевого трафіку для виявлення аномальної активності та прогнозування потенційних кібератак. У разі виявлення підозрілих

дій такі системи можуть автоматично обмежувати доступ користувача та повідомляти адміністратора про можливу загрозу інформаційній безпеці.

Водночас впровадження штучного інтелекту в правоохоронну діяльність супроводжується низкою проблем і викликів. Ефективність роботи інтелектуальних систем безпосередньо залежить від якості та актуальності даних, на основі яких здійснюється аналіз. Використання неповної або недостовірної інформації може призводити до формування неточних висновків і помилкових прогнозів. Крім того, інтеграція нових технологічних рішень повинна здійснюватися без порушення стабільності вже функціонуючих програмних комплексів.

Окремим ризиком є те, що системи штучного інтелекту отримують доступ до значних обсягів конфіденційної та вразливої інформації, зокрема персональних даних громадян, службових відомостей, матеріалів кримінальних проваджень і даних оперативного характеру. У разі недостатнього рівня захисту або помилок у функціонуванні таких систем існує небезпека витоку інформації, несанкціонованого доступу чи неправомірного використання даних. Це зумовлює необхідність посилення кіберзахисту, впровадження ефективних механізмів контролю доступу та постійного моніторингу безпеки інформаційних ресурсів.

Не менш важливим є питання підготовки персоналу до роботи з новітніми технологіями. Працівники поліції повинні володіти належними цифровими компетентностями для ефективного використання систем штучного інтелекту та контролю результатів їх функціонування. Також важливого значення набуває проблема можливого алгоритмічного упередження, що потребує постійної перевірки та верифікації отриманих аналітичних результатів.

Отже, використання штучного інтелекту в адмініструванні інформаційних систем Національної поліції України є одним із важливих напрямів модернізації правоохоронної діяльності. Комплексне поєднання сучасних технологічних рішень, належної методичної підтримки та професійної підготовки кадрів сприятиме забезпеченню стабільної роботи інформаційної інфраструктури, підвищенню рівня кібербезпеки та ефективності діяльності поліції в умовах сучасних викликів і загроз.

Водночас ефективність упровадження систем штучного інтелекту значною мірою залежить від належного рівня захисту інформації та дотримання вимог щодо безпеки персональних даних. Оскільки інтелектуальні системи можуть отримувати доступ до значних обсягів конфіденційної службової інформації, особливого значення набуває створення надійних механізмів кіберзахисту та контролю доступу до інформаційних ресурсів. Тому подальший розвиток технологій штучного інтелекту у діяльності Національної поліції України має супроводжуватися вдосконаленням нормативно-правового регулювання, технічного захисту інформації та системи професійної підготовки персоналу.

Список використаних джерел

1. Баранов О.А. Визначення терміна «штучний інтелект». *Інформація і право*. 2023. №4. С. 32–49.
2. «Штучний інтелект чи машинне навчання? В чому різниця». 2024. URL: <https://www.zfort.com.ua/blog/shtuchnii-intelekt-chi-mashinne-navchannya-v-chomu-riznicya>.
3. Войтушенко І. О., Мамедова Е. А. Інтеграція штучного інтелекту в системи екстреної диспетчеризації Національної поліції України. *Науковий вісник Львівського державного університету внутрішніх справ. Серія юридична*. 2025. № 2. С. 34-40. DOI: <https://doi.org/10.32782/2311-8040/2025-2-5>.
4. Мамедова Е.А., Войтушенко О.І. Новий рубіж поліцейської діяльності. Засади використання штучного інтелекту Національної поліції України. *Міжнародний досвід. Центральноукраїнський вісник права та публічного управління*. 2024. № 4. С. 51–56. DOI: <https://doi.org/10.32782/cuj-2024-4-7>.

Владислав ГАВЛОВСЬКИЙ, Міжвідомчий науково-дослідний центр з проблем боротьби з організованою злочинністю при РНБО України, кандидат юридичних наук, старший науковий співробітник.

Михайло ГУЦАЛЮК, Міжвідомчий науково-дослідний центр з проблем боротьби з організованою злочинністю при РНБО України, кандидат юридичних наук, старший науковий співробітник

ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ПРОТИДІЇ ОРГАНІЗОВАНИЙ ЗЛОЧИННОСТІ

Стрімкий розвиток технологій штучного інтелекту (ШІ) став одним із ключових факторів трансформації сучасного суспільства. Менш ніж за три роки понад 1,2 мільярда людей скористалися інструментами штучного інтелекту, що є швидшим темпом впровадження, ніж Інтернет, персональний комп'ютер чи навіть смартфон [1].

Сьогодні ШІ активно використовується у різноманітних сферах людської діяльності. Водночас поряд із позитивними можливостями технології штучного інтелекту дедалі частіше використовуються організованими злочинними угрупованнями для підвищення ефективності кібератак, фінансового шахрайства, інформаційних маніпуляцій та приховування слідів злочинів тощо. За оцінками Europol, сучасна організована злочинність фактично перетворюється на високотехнологічні кримінальні підприємства, які активно використовують цифрові технології, криптовалюти, штучний інтелект та глобальні комунікаційні платформи [2].

Одним із найбільш поширених напрямів злочинного використання ШІ є кібератаки та автоматизація злочинної діяльності [3]. Сучасні алгоритми здатні автоматично сканувати мережеву інфраструктуру та програмне забезпечення з метою пошуку вразливостей, аналізувати мільйони рядків коду та здійснювати інтелектуальний підбір способів обходу систем захисту. Крім того, організовані групи використовують AI-assisted malware – шкідливе програмне забезпечення, яке здатне адаптуватися до антивірусних систем та змінювати власний код для уникнення виявлення. Значну небезпеку становить і використання великих мовних моделей (LLM) для створення персоналізованих фішингових повідомлень різними мовами, що суттєво підвищує ефективність соціальної інженерії.

Окремим напрямом є використання технологій DeepFake та генеративного ШІ для фінансового шахрайства. Злочинці створюють підроблені відео- та аудіозаписи, імітуючи зовнішність або голос керівників компаній, представників банків чи родичів потерпілих. Такі технології активно застосовуються у схемах CEO fraud, під час яких жертви перераховують значні кошти на рахунки зловмисників. Генеративний ШІ також дозволяє створювати надзвичайно реалістичний фейковий контент, який складно відрізнити від справжнього, що суттєво ускладнює діяльність правоохоронних органів [4].

В Україні вже зафіксовано практичні випадки використання ШІ організованими групами. Зокрема, у 2025 році Національна поліція України викрила організовану групу, яка за допомогою технологій DeepFake оформлювала кредити на громадян України, завдавши збитків на суму понад 4 млн грн [5].

ШІ активно використовується й для легалізації злочинних доходів. Алгоритми машинного навчання дозволяють злочинцям автоматизувати дроблення фінансових транзакцій, аналізувати механізми банківського моніторингу та оптимізувати схеми відмивання коштів через криптовалютні сервіси та міжнародні фінансові платформи. Крім того, організовані групи використовують бот-мережі та генеративний ШІ для проведення

масштабних дезінформаційних кампаній, створення фейкових акаунтів та маніпулювання суспільною думкою.

У таких умовах особливої актуальності набуває питання використання штучного інтелекту самими правоохоронними органами для протидії організованій злочинності.

Одним із ключових напрямів застосування ШІ є аналіз великих масивів даних (Big Data). Сучасні алгоритми дозволяють правоохоронним органам здійснювати пошук прихованих зв'язків між учасниками злочинних угруповань, аналізувати структуру кримінальних мереж та встановлювати роль окремих осіб у злочинній діяльності. Так званий *criminal network analysis* дає змогу виявляти організаторів злочинних схем, координаторів фінансових операцій та осіб, причетних до міжнародних злочинних мереж. Крім того, ШІ використовується для автоматизованого аналізу фінансових транзакцій, криптовалютних переказів та електронних комунікацій, що сприяє виявленню підозрілих операцій і схем відмивання коштів.

Важливим напрямом є *predictive policing* – використання алгоритмів прогнозування злочинної активності. Аналізуючи історичні дані про правопорушення, соціально-економічні фактори та геопросторову інформацію, системи штучного інтелекту можуть визначати так звані «гарячі точки» – райони з підвищеною ймовірністю вчинення злочинів. Це дозволяє правоохоронним органам ефективніше планувати патрулювання, здійснювати превентивні заходи та застосовувати ризик-орієнтований підхід у боротьбі з організованою злочинністю. У багатьох країнах ЄС та США такі системи вже використовуються для прогнозування кримінальної активності та оптимізації розподілу ресурсів поліції.

Суттєву роль ШІ відіграє і в цифровій криміналістиці. Інтелектуальні системи здатні автоматизувати аналіз значних обсягів цифрових доказів, пришвидшувати пошук релевантної інформації та виявляти приховані закономірності. Особливо актуальним є використання ШІ для OSINT-розвідки, аналізу соціальних мереж, даркнет-майданчиків та зашифрованих каналів комунікації. Алгоритми машинного навчання дозволяють виявляти аномалії у мережевому трафіку, підозрілу поведінку користувачів та ознаки підготовки кібератак або фінансових злочинів. Такі технології суттєво підвищують ефективність розслідування складних кіберзлочинів і діяльності транснаціональних злочинних угруповань.

Окремим напрямом застосування ШІ є використання біометричних систем. Системи розпізнавання облич та відеоаналітики дозволяють ідентифікувати підозрюваних у режимі реального часу, здійснювати пошук осіб у натовпі та аналізувати відеозаписи з камер спостереження. В Україні з 2022 року активно використовується система Clearview AI, яка допомагає правоохоронним органам встановлювати особи військових злочинців, перевіряти осіб на блокпостах та ідентифікувати загиблих. Таким чином, використання технологій штучного інтелекту створює нові можливості для ефективної протидії організованій злочинності, однак одночасно потребує належного нормативно-правового регулювання та забезпечення балансу між безпекою і правами людини.

Водночас впровадження таких технологій породжує низку правових, етичних та організаційних проблем, пов'язаних із захистом прав людини, приватності та нормативним регулюванням використання ШІ у правоохоронній діяльності.

Насамперед актуальною залишається проблема відсутності чіткого законодавчого визначення злочинів, учинених із використанням технологій штучного інтелекту. В Україні більшість таких діянь кваліфікуються за традиційними нормами кримінального законодавства, зокрема як шахрайство, несанкціоноване втручання в роботу інформаційних систем або порушення недоторканності приватного життя. Водночас чинне законодавство не враховує специфіку використання генеративного ШІ, DeepFake-технологій чи автономних алгоритмів. Окремою проблемою є *admissibility AI-generated evidence* - допустимість доказів, отриманих або сформованих із використанням алгоритмів штучного інтелекту, а також питання достовірності та перевірки таких доказів у кримінальному провадженні.

Серйозні виклики виникають і в етичній сфері. Використання систем відеоспостереження, біометричної ідентифікації та аналізу великих масивів персональних даних може створювати ризики порушення права на приватність та сприяти формуванню механізмів масового стеження.

Крім того, алгоритми штучного інтелекту можуть містити помилки або демонструвати алгоритмічну дискримінацію через використання неповних чи упереджених наборів даних. Особливу небезпеку становлять помилки систем розпізнавання облич, які можуть призводити до неправильної ідентифікації осіб та порушення прав людини.

Не менш важливими залишаються організаційні проблеми. Правоохоронні органи часто стикаються з нестачею висококваліфікованих фахівців у сфері штучного інтелекту, цифрової криміналістики та аналізу даних. Додатковими перешкодами є недостатній рівень технічного забезпечення, обмежене фінансування та залежність від іноземного програмного забезпечення. В Україні також відсутня комплексна державна стратегія щодо впровадження технологій ШІ у діяльність правоохоронних органів, що ускладнює координацію між державними органами та створення єдиної системи протидії сучасним формам організованої злочинності.

Насамкінець важливим є і використання ШІ для прийняття рішень у судовій діяльності, адже у кінцевому рахунку тільки суд може приймати відповідний вирок по справі [6].

Список використаних джерел

1. AI Diffusion Report: Mapping global AI adoption and innovation URL: <https://news.microsoft.com/source/emea/features/ai-diffusion-report-mapping-global-ai-adoption-and-innovation-3/>
2. Europol warns of AI-driven crime threats URL: <https://www.reuters.com/world/europe/europol-warns-ai-driven-crime-threats-2025-03-18>
3. Mykhaylo Gutsalyuk AI in Cybersecurity: An Instrument for Defense and Attack // Oleksii Kostenko, & Yuriy Yekhanurov (Eds.). (2024). DIGITAL TRANSFORMATION IN UKRAINE: AI, METAVERSE, AND SOCIETY 5.0. DOI: <https://doi.org/10.69635/978-1-0690482-1-9-ch23>
4. Europol warns AI is making organized crime ‘more precise and devastating’ because it’s able to create ‘highly realistic synthetic media’ that propagates violence and normalizes corruption. URL: <https://fortune.com/2025/03/18/europol-ai-organized-crime>
5. Шахрайство з використанням ШІ: поліція викрила групу, яка оформила кредитів на українців на 4 млн. URL: <https://minfin.com.ua/ua/2025/10/14/160446190>
6. Впровадження штучного інтелекту у судову діяльність: міжнародний та український досвід // *Наукові праці Київського авіаційного інституту*. Серія: Юридичний вісник «Повітряне і космічне право». № 2(75). С. 167–175. DOI: <https://doi.org/10.18372/2307-9061.75.20230>

Сергій ЄСІМОВ, професор кафедри адміністративно-правових дисциплін, Львівський державний університет внутрішніх справ, кандидат юридичних наук, професор

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ОСВІТІ

Стратегія цифрового розвитку інноваційної діяльності України на період до 2030 року та затвердження операційного плану заходів з її реалізації у 2025-2027 роках визначає пріоритети науково-технічного розвитку: перехід до передових цифрових інтелектуальних виробничих технологій, роботизованих систем, нових матеріалів і способів конструювання, створення систем обробки великих обсягів даних, машинного навчання та штучного інтелекту [1]. У поточну діяльність органів публічного управління та суб'єктів господарської діяльності впроваджуються інструменти комплексної інформаційно-аналітичної підтримки прийняття рішень та переведення важливих державних функцій, процесів взаємодії держави з фізичними та юридичними особами у цифровий формат.

Інформаційні технології, втіленням яких є системи штучного інтелекту, продовжують швидко розвиватися в багатьох сферах життя. З'являються чат-боти та голосові помічники, які можуть підтримати діалог, системи автоматичної біржової торгівлі, надання банківських послуг без участі людини-працівника фінансової організації. У центрі уваги використання штучного інтелекту фігурує не формальна адаптація технологій, а зрушення логіки соціалізації, у структурі взаємодії між державою, суспільством і закладами освіти.

Розвиток цифрових технологій, поява нових форм зайнятості, розширення сфер застосування інноваційних рішень та цифрових технологій не лише підвищують вимоги до кваліфікації працівників, а й формують інші очікування роботодавців.

У результаті ці процеси можуть спровокувати виникнення нестачі фахівців із перспективно важливих напрямів, насамперед у високотехнологічних і наукомістких галузях, що здатне викликати довгостроковий дисбаланс у сфері трудової діяльності. Подібні зміни суттєво торкаються структури зайнятості та потребують переосмислення існуючих підходів до підготовки кадрів відповідно до Стратегії цифрового розвитку інноваційної діяльності України на період до 2030 року та затвердження операційного плану заходів з її реалізації у 2025-2027 роках.

Традиційні освітні програми не встигають реагувати на запити ринку, що стрімко змінюються, тому потребують комплексного оновлення. В умовах, коли професійні компетенції стають все більш спеціалізованими, механізм взаємозв'язку між системою професійної освіти та реальними потребами економіки набуває вирішального значення.

Для зниження ризику кадрового дефіциту необхідна системна трансформація управлінських моделей підготовки кадрів, що включає адаптацію навчальних планів, цифровізацію процесів і партнерство з роботодавцями.

Адаптовані до системи підготовки кадрів принципи застосування технологій штучного інтелекту, пов'язані з принципами штучного інтелекту Організації економічного співробітництва та розвитку (OECD AI Principles) та принципами застосування штучного інтелекту, викладеними в Концепції розвитку штучного інтелекту в Україні: відповідальність, прозорість, безпека, регулювання процесу використання штучного інтелекту у системі управління кадровою політикою публічного органу, юридичної особи – суб'єкта господарської діяльності, дають змогу використати системи штучного інтелекту у підборі кадрів їх навчання та підвищенні кваліфікації [2; 3].

Водночас, дорожня карта з регулювання штучного інтелекту в Україні допоможе українським компаніям підготуватися до ухвалення закону-аналога AI Act Європейського Союзу, а громадянам – навчитися захищати себе від ризиків штучного інтелекту [4].

Проект Стратегії розвитку штучного інтелекту на період до 2030 року у розділі «Система визначення пріоритетів» визначає Етап 2: Системна трансформація державного управління та економіки. На цьому етапі акцент зміщується на впровадження штучного інтелекту в роботу органів державної влади та пріоритетних сфер: освіти, науки, охорони здоров'я, енергетики та сільського господарства [5].

Технології вдосконалюються так швидко, що необхідно змінювати підходи до навчання та розвитку і дотримуватись персоналізованого підходу. Штучний інтелект та машинне навчання можуть допомогти у цьому. Організація онлайн-навчання складне завдання. Але завдяки сучасним технологіям частину одноманітної та творчої роботи можна доручити штучному інтелекту. Процеси, які робили раніше вручну, можна автоматизувати.

Підвищення кваліфікації співробітників – тобто навчання додатковим навичкам стає важливою умовою успіху, адже нерідко виробництву не вистачає висококваліфікованих працівників, які цінуються на ринку праці. Підвищувати кваліфікацію співробітників можна у різний спосіб. З технологічної точки зору, штучний інтелект і машинне навчання дозволяють проводити програми підвищення кваліфікації, які допомагають людям швидше адаптуватися до постійних змін на робочому місці та у сфері професійної діяльності.

Штучний інтелект допомагає онлайн-освіті: використовувати під час прийому нових співробітників; оптимізувати роботу структурного виробничого підрозділу чи публічного органу; автоматизувати перевірку завдань; підвищити зацікавленість здобувачів вищої освіти у вивченні предметів за програмами навчання.

Щоб навчання було якісним і ефективним, потрібно наймати досвідчених фахівців. Штучний інтелект можна використовувати при прийому на роботу лекторів. Знайти кадровика, який зможе провести співбесіду з лекторами з різних тем, дуже складно. Чим ширшим є профіль онлайн-школи, тим складніше знайти працівника кадрового апарату, який зможе перевіряти навички всіх потенційних кандидатів. В онлайн-навчання складно спрогнозувати навантаження на викладачів.

Протилежна проблема – персоналізація. Щоб тим, хто навчається, було цікавіше, вони не шукали рішення в Інтернеті, потрібно робити індивідуальні завдання для кожного. Це ускладнює процес перевірки та підвищує навантаження на співробітників. Замість найму нових людей можна скористатися ресурсами штучного інтелекту, це буде набагато дешевше. Потрібно знайти процеси, в яких штучний інтелект може замінити людину, і розпочати їхню автоматизацію.

Штучний інтелект може допомогти у творчому процесі. Штучний інтелект допомагає науково-педагогічним працівникам робити більш цікаві та варіативні завдання. У кожного слухача буде унікальний приклад, більшість з яких згенерує чат-бот за заданим лектором завданням.

Штучний інтелект у процесі навчання персоналу:

- дозволяє персоналізувати навчання. Системи штучного інтелекту аналізують навички, потреби та темпи навчання кожного співробітника, пропонуючи індивідуальний навчальний матеріал. Наприклад, новачкові можуть запропонувати детальніший опис діяльності підприємства, а досвідченому співробітнику покажуть завдання підрозділу, в якому він працює;

- аналізувати дані навчання. Штучний інтелект збирає і аналізує великі обсяги даних про освітні процеси. Це дозволяє суб'єктам господарської діяльності отримати більш усвідомлення того, які методи навчання є найбільш ефективними та виявити проблемні моменти;

- прогнозувати потреби у навчанні. Використовуючи алгоритми машинного навчання, штучний інтелект може передбачити, які навички та знання знадобляться співробітнику в майбутньому. Це дозволяє керівнику готувати колектив заздалегідь;

- поліпшувати засвоєння інформації. Штучний інтелект може застосовувати різні методики, щоб покращити засвоєння матеріалу. Наприклад, системи можуть створювати

інтерактивні уроки, використовувати візуальні ефекти та адаптувати способи навчання, залежно від стилю навчання співробітника;

- давати зворотний зв'язок та оцінювати прогрес навчання. Системи штучного інтелекту можуть створювати тести та завдання, стежити за успішністю та пропонувати рекомендації для покращення результатів.

Штучний інтелект здатний миттєво обробляти великий обсяг інформації та на її основі створювати різноманітні навчальні матеріали від текстів і презентацій до відео уроків та інтерактивних симуляцій.

Машинне навчання має великий потенціал у підвищенні кваліфікації, але поки що ця технологія перебуває на стадії становлення. Наразі необхідно виявити корисні методики, які допоможуть підвищити ефективність навчання. Завдяки стороннім інструментам можна раніше приступити до експериментів. У майбутньому машинне навчання допоможе працівникам удосконалювати технічні навички і особисті якості.

На думку О. Бортнякова, виконувача обов'язків Міністра цифрової трансформації України, проєкт Стратегії розвитку штучного інтелекту на період до 2030 року передбачає, щодо розвитку освіти і талантів: робота з штучним інтелектом (далі – ШІ) стане базовою навичкою, як читання чи письмо. Застосунок «Мрія» перетвориться на головний ШІ-інструмент учня з персональною програмою навчання. Ще один важливий фокус – розвиток професійної спільноти. Мета – підготувати 50 тисяч висококваліфікованих ШІ-фахівців в Україні та зробити штучний інтелект доступним кожному студенту [6].

Список використаних джерел

1. Про схвалення Стратегії цифрового розвитку інноваційної діяльності України на період до 2030 року та затвердження операційного плану заходів з її реалізації у 2025-2027 роках: Розпорядження Кабінету Міністрів України від 31.12.2024 р. № 1351-р. URL: <https://zakon.rada.gov.ua/laws/show/1351-2024-%D1%80#Text>

2. OECD Guidelines on Artificial Intelligence. URL: <https://www.oecd.org/en/topics/digital.html>

3. Про схвалення Концепції розвитку штучного інтелекту в Україні: Розпорядження Кабінету Міністрів України від 02.12.2020 р. № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text>

4. Регулювання штучного інтелекту в Україні: презентуємо дорожню карту. URL: <https://thedigital.gov.ua/news/technologies/regulyuvannya-shtuchnogo-intelektu-v-ukraini-prezentuemo-dorozhnyu-kartu>

5. Проєкт Стратегії розвитку штучного інтелекту на період до 2030 року. URL: <https://ai-files.thedigital.gov.ua/documents/Проєкт%20Стратегії%20розвитку%20ШІ%202030.docx.pdf>

6. Шлях розвитку ШІ до 2030 року: знайомтеся з проєктом Стратегії та давайте фідбек. URL: <https://thedigital.gov.ua/news/shtuchnyy-intelekt/shliakh-rozvytku-shi-do-2030-roku-znayomtesia-z-proyektom-stratehiyi-ta-davayte-fidbek>

Роксолана ЗАПОТІЧНА, доцент кафедри міжкультурної комунікації та соціально-гуманітарних дисциплін, Дніпровський гуманітарний університет, кандидат економічних наук, доцент

ЕКОНОМІЧНА ЕФЕКТИВНІСТЬ ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В ОСВІТНІЙ ТА НАУКОВІЙ ДІЯЛЬНОСТІ

Стрімкий розвиток технологій штучного інтелекту у сучасній глобальній економіці суттєво змінює підходи до організації освітньої та наукової діяльності. Якщо ще на початку ХХІ століття цифрові технології виконували переважно допоміжну функцію, то сьогодні системи штучного інтелекту перетворюються на ключовий інструмент управління знаннями, автоматизації процесів та підвищення ефективності використання ресурсів. У цьому контексті особливої актуальності набуває проблема оцінювання економічної ефективності впровадження технологій ШІ, оскільки освітні та наукові установи функціонують в умовах обмежених фінансових ресурсів та зростаючих вимог до якості результатів.

З економічної точки зору ефективність використання ШІ в освітній та науковій діяльності може бути визначена як співвідношення між сукупними витратами на впровадження та експлуатацію технологій і комплексом отриманих прямих та непрямих вигод. До прямих ефектів належать скорочення витрат робочого часу, автоматизація адміністративних процедур, зменшення витрат на рутинні операції та оптимізація використання людських ресурсів. Непрямі ефекти включають підвищення якості освітнього процесу, зростання продуктивності дослідницької діяльності та довгострокове підвищення рівня людського капіталу.

Сучасні емпіричні дослідження підтверджують, що використання генеративного штучного інтелекту суттєво підвищує продуктивність виконання когнітивних завдань, особливо у сфері обробки інформації, аналітики та підготовки текстових матеріалів. Так, у роботі, присвяченій аналізу впливу генеративного ШІ на продуктивність праці, встановлено, що працівники, які використовують ШІ-інструменти, демонструють значно вищу ефективність і скорочення часу виконання завдань порівняно з контрольною групою, причому найбільший ефект спостерігається у менш досвідчених користувачів [1]. Це свідчить про потенціал ШІ як інструменту зменшення професійної нерівності та підвищення загальної ефективності освітніх і наукових систем.

В освітній сфері економічна ефективність ШІ насамперед проявляється через автоматизацію адміністративних та навчально-методичних процесів. Застосування інтелектуальних систем для перевірки завдань, генерації навчальних матеріалів, управління освітніми траєкторіями та аналізу академічної успішності дозволяє значно скоротити витрати часу викладачів на рутинну діяльність. Вивільнений ресурс може бути спрямований на дослідницьку роботу, індивідуальну роботу зі студентами та розвиток інноваційних освітніх методик. Таким чином, відбувається не лише скорочення витрат, але й якісна трансформація структури викладацької праці.

Важливим економічним ефектом є масштабування освітніх послуг. Традиційна модель освіти передбачає лінійну залежність між кількістю здобувачів освіти та витратами на навчання. Водночас використання технологій ШІ та цифрових платформ дозволяє частково розірвати цю залежність, забезпечуючи можливість обслуговування значно більшої кількості студентів без пропорційного збільшення витрат. Це формує ефект економії на масштабі, який є одним із ключових чинників економічної ефективності цифрової трансформації освіти.

Окремого значення набуває вплив ШІ на якість освітнього процесу, який опосередковано визначає довгострокову економічну віддачу через формування людського капіталу. Персоналізовані системи навчання, побудовані на основі алгоритмів штучного інтелекту, дозволяють адаптувати навчальний контент до індивідуальних потреб здобувачів освіти, що

сприяє підвищенню рівня засвоєння знань та зменшенню відсіву студентів. У довгостроковій перспективі це забезпечує підвищення якості робочої сили та зростання продуктивності економіки загалом [2].

У науковій діяльності штучний інтелект виконує функцію інструменту прискорення дослідницьких процесів. Використання алгоритмів машинного навчання, великих мовних моделей та систем автоматизованого аналізу даних дозволяє суттєво скоротити час, необхідний для збору, обробки та інтерпретації наукової інформації. Це підвищує продуктивність дослідників і сприяє збільшенню кількості наукових результатів у розрахунку на одиницю часу та ресурсів. Як свідчать сучасні дослідження, впровадження ШІ у сфері knowledge work призводить до суттєвого зростання ефективності виконання складних аналітичних завдань [1].

Водночас економічна ефективність використання ШІ не є однозначною і супроводжується низкою ризиків та обмежень. Насамперед ідеться про високі початкові витрати на впровадження цифрової інфраструктури, необхідність постійного оновлення програмного забезпечення та підготовки кадрів. Додатковим викликом є ризик посилення цифрової нерівності між закладами освіти та науковими установами, які мають різний рівень доступу до технологій та фінансування. Це може призвести до формування асиметрій у розвитку освітніх систем як на національному, так і на міжнародному рівнях [2].

Окрему групу становлять приховані витрати, пов'язані з використанням ШІ, серед яких витрати на кібербезпеку, захист даних, правове регулювання та етичний контроль. Крім того, у науковій літературі все частіше обговорюється проблема когнітивного аутсорсингу, коли надмірне використання ШІ може призводити до зниження рівня критичного мислення та самостійності у здобутті знань. Це створює потенційні довгострокові економічні ризики, пов'язані зі зниженням якості людського капіталу [3].

Таким чином, економічна ефективність використання технологій штучного інтелекту в освітній та науковій діяльності має комплексний характер і охоплює як короткострокові ефекти підвищення продуктивності та зниження витрат, так і довгострокові структурні зміни у функціонуванні освітніх і наукових систем. З одного боку, ШІ виступає потужним інструментом оптимізації ресурсів та підвищення ефективності, з іншого — потребує значних інвестицій та супроводжується ризиками інституційної нерівності. Отже, економічна доцільність його використання повинна оцінюватися з урахуванням балансу між витратами, короткостроковими вигодами та стратегічними довгостроковими ефектами для розвитку людського капіталу та економіки знань.

Список використаних джерел

1. Noy S., Zhang W. Experimental evidence on the productivity effects of generative artificial intelligence. NBER Working Paper No. 31161. 2023. URL: <https://www.nber.org/papers/w31161>
2. Zawacki-Richter O. et al. Systematic review of research on artificial intelligence applications in higher education. International Journal of Educational Technology in Higher Education. 2019. URL: <https://educationaltechnologyjournal.springeropen.com>
3. Dwivedi Y. K. et al. Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research. International Journal of Information Management. 2021. URL: <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>

Василь АНДРУНИК, доцент кафедри інформаційних систем та мереж, Національний університет «Львівська політехніка», кандидат технічних наук, доцент.

Наталія КАЛЬКА, доцент кафедри практичної психології навчально-науковий інститут управління, психології та безпеки, Львівський державний університет внутрішніх справ, доктор філософії в галузі психології

ДОЦІЛЬНІСТЬ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ЯК ІНСТРУМЕНТУ ВДОСКОНАЛЕННЯ ТЕХНІКИ «БЕЗПЕЧНЕ МІСЦЕ» У ПСИХОТЕРАПЕВТИЧНІЙ ПРАКТИЦІ

Використання техніки «Безпечне місце» неодноразово доводить сою результативність у роботі із подоланням тривожних та стресових станів, симптоматики ПТСР та активації внутрішніх ресурсів. Проте із стрімкою цифровізацією виникає потреба вбудовування у класичну структуру та алгоритм проведення цієї техніки штучного інтелекту. Нами пропонується алгоритм побудови взаємодії, що включає наступні етапи (Рис.1).

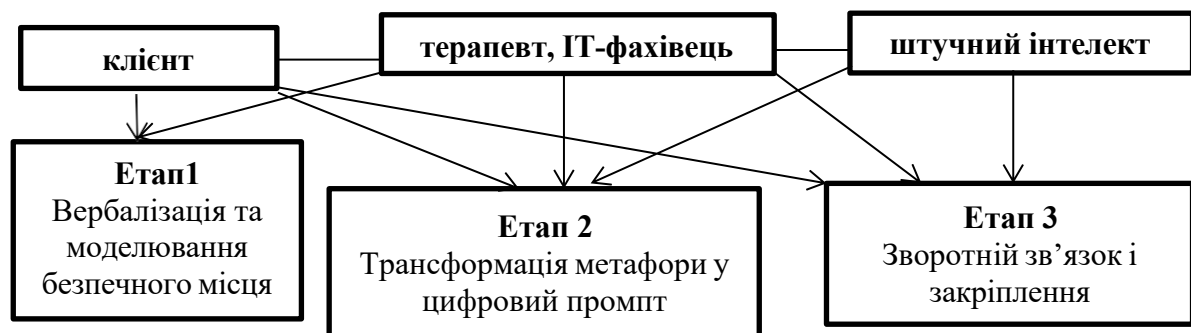


Рисунок 1 – Схема алгоритму реалізації техніки «Безпечне місце» з допомогою ШІ

На першому етапі відбувається безпосередня взаємодія терапевта з клієнтом. Роль терапевта на цьому етапі полягає у мотивації внутрішнього пошуку безпечного місця клієнтом і відповідно вербалізація результату за допомогою метафор, кольорів та образів. Клієнт надає інформацію про візуальний образ безпечного місця, а терапевт спрямовує його за допомогою запитань. На цьому етапі штучний інтелект не задіяний, адже відбувається лише взаємодія «клієнт-терапевт».

Другий етап включає в себе взаємодію у вище згаданій тріаді «клієнт-терапевт-ШІ». Саме на цьому етапі терапевт разом з клієнтом виокремлює ключові слова-дескриптори. Потім терапевт або ІТ-фахівець формує текстовий запит (промпт) і пропонує ШІ згенерувати уявний візуальний образ безпечного місця.

Третій етап пов'язаний із інтерактивною корекцією, верифікацією та стабілізацією. Відбувається споглядання клієнтом та робота із згенерованим зображенням та оцінка своїх відчуттів, емоцій та стану. Якщо певні деталі зображення викликають дискомфорт або не збігаються з внутрішнім відчуттям безпеки, клієнт дає зворотний зв'язок терапевту чи ІТ-фахівцю. Вони коригують промпт, а ШІ адаптує картинку до ідеального, екологічного для клієнта стану. Етап завершується фіксацією індивідуального «Безпечного місця» для подальшої психологічної стабілізації.

Переваги використання ШІ у техніці «Безпечне місце», це пер за все, зменшення когнітивного навантаження, необхідного для роботи уяви. чіткий згенерований візуальний стимул допомагає швидше досягнути стану розслаблення чи стабілізації емоційного стану.

По-друге, у ситуаціях психотравматизації у клієнтів часто переважає стан безпорадності: можливість за допомогою ШІ коригувати візуальне зображення повертає клієнту відчуття контролю і відповідно над своїм емоційним станом

По-третє, під час генерування зображення на запит клієнта ШІ працює як тригер позитивних асоціацій, оскільки контакт із варіантами ми зображення допомагає швидше включатися в процес, пригадувати власні позитивні відчуття і детального розгортати терапевтичну метафору [1].

Попри високу терапевтичну ефективність запропонованого алгоритму, варто окреслити зони ризику, які потребують чіткого контролю з боку терапевта, серед яких можливість некоректної генерації і як наслідок ре-травматизація клієнта (ШІ може генерувати деталі, які будуть прихованими тригерами клієнта і впливатимуть на активізацію гострих стресових станів замість бажаної стабілізації).

Також існує небезпека, що ШІ запропонує занадто яскравий, деталізований, але «чужий» для клієнта естетичний образ як наслідок, клієнт замість того, щоб спиратися на власні внутрішні ресурси, може піддатися впливу «готової картинки» від ШІ, що свою чергу полабить терапевтичний ефект, оскільки згенероване місце не буде повністю конгруентним його внутрішньому стану.

Варто згадати, що використання ШІ може провокувати «ефект цифрової милиці», адже клієнт звикає до використання гаджету із згенерованим зображенням, що може перешкоджати формуванню автономної навички саморегуляції в реальному житті, коли ШІ немає під рукою [2].

Також варто пам'ятати, що використання комерційних або відкритих платформ ШІ передбачає передачу текстових описів на сервери сторонніх компаній. Оскільки описи безпечного місця вразливих категорій клієнтів часто містять глибоко особисту інформацію, виникає небезпека витоку даних або їх використання для навчання нейромереж без належної згоди.

Інтеграція штучного інтелекту у структуру техніки «Безпечне місце» відкриває принципово нові можливості для реалізації цієї техніки, суттєво пришвидшуючи процес емоційної стабілізації та повертаючи клієнту відчуття контролю через цифрову екстерналізацію терапевтичних метафор.

Список використаних джерел

1. Андруник В.А., Калька Н.М., Кривошия М.В., Одинцова Г.Ю., Христук О.Л. Використання VR- середовищ ТА AI згенерованого контенту у психотерапевтичній техніці «Безпечне місце» у роботі з психотравмою. *Наукові перспективи*. № 9(63) 2025 С.1388-1403. DOI: [https://doi.org/10.52058/2708-7530-2025-9\(63\)-1388-1402](https://doi.org/10.52058/2708-7530-2025-9(63)-1388-1402)
2. Dai R., Kim J., Kim K. Safe Place VR: A Guided Imagery-Based Virtual Reality Trauma Stabilization Technique. *IEEE Conference on Virtual Reality and 3D User Interfaces Abstracts and Workshops (VRW)*. P.1460–1461. DOI: <https://doi.org/10.1109/vrw66409.2025.00376>

Олексій СКИЦЬКО, науковий співробітник, Національна академія Служби безпеки України, кандидат технічних наук, старший науковий співробітник

ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В КІБЕРБЕЗПЕЦІ: МОЖЛИВОСТІ, ЗАГРОЗИ ТА ВИКЛИКИ

Вступ. Штучний інтелект (ШІ) став центральним елементом сучасного ландшафту кіберзагроз – водночас потужним інструментом захисту й засобом нападу. За звітом ENISA Threat Landscape 2025, який охопив 4 875 інцидентів за період липень 2024 - червень 2025, на початок 2025 року вже понад 80% спостережуваних соціально-інженерних атак спираються на ШІ для генерації фішингових повідомлень, deepfake-контенту й автоматизованих маніпулятивних кампаній [1]. За даними IBM Cost of a Data Breach Report 2025, середня глобальна вартість витоку даних знизилась до 4,44 млн дол. США (-9% за рік) – переважно завдяки скороченню часу виявлення й реагування за допомогою AI-захисту: екстенсивне впровадження ШІ й автоматизації скорочує життєвий цикл інциденту в середньому на 80 днів і дає економію близько 1,9 млн дол. порівняно з організаціями без ШІ-захисту. Водночас 97% організацій, які зазнали інциденту, пов'язаного з ШІ, не мали належного контролю доступу до своїх ШІ-систем, а 63% не мають політик керування ШІ взагалі [2]. Метою тез є системний огляд застосувань ШІ у кібербезпеці – як захисних, так і наступальних, а також аналіз нових класів загроз, спрямованих на самі ШІ-моделі, і регуляторних рамок, що формуються навколо них.

Захисні застосування ШІ. Сучасні системи кібербезпеки використовують методи машинного й глибокого навчання у багатьох задачах: виявлення зловмисного ПЗ (статичний і динамічний аналіз, behavior-based detection), мережеві системи виявлення вторгнень (IDS/IPS), поведінкова аналітика користувачів і сутностей (UEBA), виявлення фішингу й шахрайства, кореляція подій у SIEM/SOC та автоматизація розслідувань, виявлення вразливостей у вихідному коді, аналіз threat intelligence й обробка великих обсягів OSINT-даних. Систематичний огляд понад 60 досліджень показав, що ML та DL-методи разом із метасвистичними алгоритмами суттєво перевершують традиційні сигнатурні рішення в задачах виявлення зловмисного програмного забезпечення, мережових вторгнень, спаму та фішингу; за окремими архітектурами повідомляється про точність 95-99% (зокрема, CNN-моделі для виявлення вразливостей за датасетом MITRE CVE/CWE сягають близько 98%) [3]. Принципова перевага AI-підходів, здатність детектувати атаки нульового дня, які виходять за межі заздалегідь визначених сигнатур.

Наступальне використання ШІ супротивниками. Зловмисники застосовують ШІ для: автоматизованої генерації фішинг- та spear-phishing-повідомлень із високим рівнем мовної достовірності та персоналізації; створення deepfake-аудіо й відео для атак типу CEO-fraud, шахрайських «голосових інструкцій» і маніпуляцій громадською думкою; розробки поліморфного зловмисного ПЗ, що змінює сигнатури між запусками; обходу CAPTCHA та біометричної верифікації; автоматизованого пошуку експлойтів. ENISA фіксує появу спеціалізованих «зловмисних» ШІ-моделей (зокрема Xanthorox AI), спрямованих на автоматизацію соціальної інженерії та розробки шкідливого коду, а також гібридні інформаційні операції, які поєднують AI-генерований контент із традиційними кібертактиками [1]. За даними IBM, ШІ був задіяний у 16% усіх витоків даних -переважно для фішингу й deepfake-атак, які залишаються одним із найдорожчих типів інцидентів (середня вартість фішинг-витоку складає 4,8 млн дол.) [2].

Атаки на самі ШІ-системи. Окрему категорію загроз становлять атаки на самі моделі. В оновленій таксономії NIST AI 100-2 E2025 (березень 2025) формалізовано чотири основні класи атак для генеративного ШІ – ухилення (evasion), отруєння даних або моделі (poisoning,

включно з clean-label poisoning), порушення приватності (privacy) та зловживання (misuse/abuse), окремо описано prompt injection та витoki тренувальних даних [4]. Структуровану базу знань про реальні тактики й техніки атак на ШІ-системи у форматі ATT&CK-подібної матриці підтримує MITRE ATLAS - на кінець 2025 року вона містить 16 тактик, 84 техніки на основі 42 реальних кейсів [5]. OWASP Top 10 для LLM-застосунків (видання 2025) виокремлює критичні ризики LLM-систем: prompt injection, витік системного пром프트, отруєння даних і моделі, надмірна автономія агентів (excessive agency), уразливості ланцюга постачання тощо [6]. Окремим організаційним викликом залишається shadow AI - несанкціонована робота співробітників з публічними LLM: він був задіяний у 20% витоків даних і додавав у середньому близько 670 тис. дол. до вартості інциденту [2].

Виклики та обмеження. Попри прогрес, AI-системи в кібербезпеці мають низку обмежень: висока частка хибних спрацьовувань на тлі дрейфу концепції; обмежена пояснюваність (explainability) рішень глибоких нейромереж, що ускладнює аудит SOC-розслідувань; галюцинації LLM-копілотів, видумування неіснуючих CVE, неправильна атрибуція індикаторів компрометації, фабрикація посилань на джерела. Як показано в [7], галюцинації є внутрішньою властивістю самого процесу навчання й оцінювання мовних моделей, що накладає принципові обмеження на автономність security-копілотів. Додаткові виклики: залежність якості моделей від обсягу й чистоти навчальних даних, потреба у вимірюванні стійкості до adversarial-вхідних даних, питання приватності. Це вимагає збереження «людини в контурі» (human-in-the-loop) для критичних рішень.

Регуляторний вимір. Базовою методологічною рамкою для управління ризиками ШІ є NIST AI Risk Management Framework (AI RMF 1.0) [8], що визначає функції GOVERN, MAP, MEASURE, MANAGE для всього життєвого циклу системи. У Європейському Союзі набув чинності Регламент 2024/1689 (Акт про ШІ, EU AI Act), який запровадив диференційовані обов'язки залежно від рівня ризику, з особливими вимогами до високоризикових систем - зокрема, у сферах критичної інфраструктури, правоохоронної діяльності, керування кордонами та біометричної ідентифікації [9]. В українському науковому дискурсі системний аналіз загроз і ризиків застосування ШІ викладено в Електронному фаховому науковому виданні «Кібербезпека: освіта, наука, техніка» [10], що формує підґрунтя для національних регуляторних ініціатив і профільних освітніх програм.

Висновки. ШІ кардинально розширює можливості як захисту, так і нападу в кіберпросторі: він знижує середню вартість витоку даних і скорочує час реагування, але водночас створює нові класи загроз, від AI-керованого фішингу й deepfake до атак на самі моделі. Безпечне впровадження ШІ в кібербезпеку потребує: інтеграції принципів security-by-design у життєвий цикл моделей; застосування таксономій NIST AI 100-2 та MITRE ATLAS при проєктуванні моделей загроз; впровадження політик AI-governance і контролю shadow AI; забезпечення пояснюваності та утримання людини в контурі для критичних рішень; узгодження внутрішніх процесів з NIST AI RMF і EU AI Act. Перспективними напрямками є виявлення adversarial-атак у режимі реального часу, безпечні агентні архітектури та національна екосистема оцінювання довіреного ШІ.

Список використаних джерел

1. ENISA Threat Landscape 2025 / European Union Agency for Cybersecurity. October 2025. 87 p. URL: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2025> (дата звернення: 20.05.2026).
2. Cost of a Data Breach Report 2025 / IBM Security; Ponemon Institute. 2025. URL: <https://www.ibm.com/reports/data-breach> (дата звернення: 20.05.2026).
3. Salem A. H., Azzam S. M., Emam O. E., Abohany A. A. Advancing cybersecurity: a comprehensive review of AI-driven detection techniques. Journal of Big Data. 2024. Vol. 11, Art. No. 105. DOI: 10.1186/s40537-024-00957-y.
4. Vassilev A., Oprea A., Fordyce A., Anderson H., Davies X., Hamin M. Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations / National Institute of Standards

and Technology. NIST AI 100-2 E2025. Gaithersburg, MD, 2025. DOI: 10.6028/NIST.AI.100-2e2025.

5. MITRE ATLAS: Adversarial Threat Landscape for Artificial-Intelligence Systems / The MITRE Corporation. URL: <https://atlas.mitre.org/> (дата звернення: 20.05.2026).

6. OWASP Top 10 for Large Language Model Applications 2025 / OWASP Foundation. 2025. URL: <https://genai.owasp.org/resource/owasp-top-10-for-llm-applications-2025> / (дата звернення: 20.05.2026).

7. Kalai A. T., Nachum O., Vempala S. S., Zhang E. Why Language Models Hallucinate. arXiv:2509.04664v1. 2025. URL: <https://arxiv.org/pdf/2509.04664> (дата звернення: 20.05.2026).

8. Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1 / National Institute of Standards and Technology. Gaithersburg, MD, 2023. DOI: 10.6028/NIST.AI.100-1.

9. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union. L, 2024/1689, 12.07.2024. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (дата звернення: 20.05.2026).

10. Загрози та ризики застосування технологій штучного інтелекту. Кібербезпека: освіта, наука, техніка. 2025. URL: <https://csecurity.kubg.edu.ua/index.php/journal/article/view/520> (дата звернення: 20.05.2026).

Роман ШИРШОВ, старший викладач, Національна академія Служби безпеки України

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В РОЗВІДЦІ ЗА ВІДКРИТИМИ ДЖЕРЕЛАМИ (OSINT): МОЖЛИВОСТІ, ІНСТРУМЕНТИ ТА ОБМЕЖЕННЯ

Вступ. Розвідка за відкритими джерелами (Open-Source Intelligence, OSINT) є основою сучасної аналітичної діяльності правоохоронних і безпекових структур: за оцінками, OSINT забезпечує від 80 до 90% усього розвідувального матеріалу західних служб [1]. Стрімке зростання обсягів цифрових даних – соціальних медіа, супутникових знімків, IoT-сенсорів, корпоративних реєстрів, медіа-архівів – зробило традиційні ручні підходи неспроможними обробляти інформаційний потік у режимі реального часу. Світовий ринок OSINT-рішень оцінюється у 14,85 млрд дол. США (2024) із прогнозом зростання до 49,39 млрд дол. до 2029 року (CAGR близько 28%), і значна частина цього зростання припадає саме на AI-підсилені платформи [2]. Систематичний огляд Browne та ін. (2024), що охопив 163 рецензованих публікації за період після 2019 року, фіксує лавиноподібне зростання застосування ШІ на всіх етапах OSINT-циклу [3]. Метою тез є системний огляд застосувань ШІ у задачах OSINT, аналіз ключових інструментальних напрямів і обмежень, а також формулювання вимог до методології безпечного впровадження.

Етапи OSINT-циклу, в яких задіяний ШІ. Класичний цикл - постановка завдання, збір, обробка, аналіз, поширення - сьогодні насичений AI-компонентами на кожному кроці. **Збір** автоматизовано через скрейпінг з LLM-навігацією, моніторинг соціальних мереж і dark web, обробку нестандартних джерел. **Обробка** – нормалізація, дедуплікація, переклад, розпізнавання іменованих сутностей. **Аналіз** – кластеризація, побудова графів зв'язків, виявлення аномалій, прогнозування. **Поширення** – автоматична генерація звітів, інтерактивних дашбордів і брифінгів. Огляд Browne та ін. показує, що 44% досліджень з AI-OSINT походять зі США та Індії; домінують методи NLP, глибокого навчання і графових неймереж [3].

III у роботі з текстовою інформацією. Засоби обробки природної мови (NLP) вирішують задачі розпізнавання іменованих сутностей (NER) – людей, організацій, локацій, торгових марок, ідентифікаторів транспорту; машинного перекладу для мультимовного OSINT; реферування й тематичного моделювання; аналізу тональності та поляризації дискурсу; розв’язання сутностей (entity resolution) для зведення профілів у різних соціальних мережах. Окремий клас задач – LLM-агенти-дослідники, які автоматично виконують багатокрокові запити: декомпонують питання, виконують пошук, синтезують результат. Експериментальне дослідження Karabiyik та ін. (2024) показало, що сучасні генеративні моделі (Gemini, GPT-4) демонструють значущий OSINT-потенціал, але водночас містять вбудовані фільтри приватності, які обмежують доступ до персональних даних [4].

Агентний OSINT та автоматизація розслідувань. Новою хвилею є агентні системи (agentic AI) на основі LLM, що автономно планують і виконують багатокрокові розслідування: декомпонують завдання, обирають інструменти (пошук, скрейпінг, бази даних, API), ітеративно уточнюють гіпотези й самостійно синтезують звіт із посиланнями на джерела. Архітектури з доповненою пошуком генерацією (retrieval-augmented generation, RAG) прив’язують відповіді моделі до перевірених документів, знижуючи рівень галюцинацій. Такі системи інтегруються з класичними OSINT-платформами – графовими аналізаторами зв’язків, інструментами візуалізації та геопросторовими модулями, – перетворюючи аналітика з виконавця рутинних запитів на верифікатора й інтерпретатора. Browne та ін. окремо відзначають інтеграцію наявних OSINT-інструментів з III як ключову, але досі недостатньо опрацьовану дослідницьку прогалину [3]. Водночас зростання автономності агентів підвищує ризики: помилка на одному кроці поширюється ланцюгом міркувань, а надмірні повноваження агента (excessive agency) створюють нові вектори атак.

Візуальний OSINT (IMINT/GEOINT). Машинне навчання радикально розширило можливості роботи з зображеннями та відео: класифікація об’єктів і логотипів, автоматичне розпізнавання військової техніки, типів літальних апаратів і кораблів; геолокація за зображенням (cross-view matching, прив’язка до супутникових знімків) і хронолокація (визначення часу за тінями); виявлення дублікатів і модифікацій (reverse image search); розпізнавання облич; аналіз супутникових знімків (зміна забудови, рух військ, оцінка пошкоджень). Зростання обсягів мультимедіа в соціальних мережах (SOCMINT) робить ручний перегляд тисяч фото й відео непосильним – AI-моделі дозволяють автоматично відсіювати нерелевантний контент і виокремлювати значущі кадри. У роботі з морськими даними AIS і трекінгом транспорту AI-моделі виявляють патерни підозрілої поведінки – «темні» зупинки сигналу, типові схеми обходу санкцій, перевалки в морі. Цей напрям успішно застосовувано у відомих розслідуваннях Bellingcat – від збиття MH17 до сирійського конфлікту й війни в Україні; у 2024 році опубліковано Bellingcat Online Investigation Toolkit, який системно каталогізує AI- та не-AI-інструменти для практиків [5].

Корпоративна розвідка та фінансовий комплаєнс. Одним із найдинамічніших напрямів прикладного AI-OSINT є корпоративна розвідка, due diligence і протидія відмиванню коштів. AI-платформи автоматично агрегують дані з корпоративних реєстрів, судових баз, санкційних списків (OFAC, ЄС, ООН, РНБО), реєстрів кінцевих бенефіціарних власників (UBO) та adverse-media джерел, будуючи графи власності й контролю для виявлення прихованих зв’язків і кінцевих вигодоодержувачів. Так, платформа S&P Global Entity Insights у партнерстві з Quantifind забезпечує скринінг понад 30 млн записів про сутності – контролерів, топ-менеджерів, членів рад і власників – за санкціями та негативними згадками з понад 30 тис. джерел [6]. Машинне навчання істотно знижує частку хибних збігів при санкційному скринінгу й скорочує час перевірки бенефіціарної структури з десятків годин до кількох; предиктивна аналітика дозволяє переходити від статичної перевірки відповідності до безперервного моніторингу зміни ризик-профілю контрагента. Цей напрям критично важливий для виявлення схем обходу санкцій – зокрема через ланцюги підставних компаній.

Інструментальна екосистема. Ринок пропонує широкий спектр AI-підсиленних OSINT-платформ: комплексні розвідувальні рішення з інтегрованими GEOINT- і SOCMINT-модулями, інструменти аналізу й візуалізації графів зв'язків (наприклад, Maltego), платформи корпоративної розвідки та перевірки бенефіціарних власників (зокрема Sayari, YouControl), агрегатори цифрових слідів за ідентифікаторами (OSINT Industries). Систематизовані каталоги таких інструментів – як-от Bellingcat Online Investigation Toolkit – спрощують підбір рішень під конкретне завдання [5]. Стійкою тенденцією є вбудовування LLM-асистентів безпосередньо в робочі процеси аналітика: для автоматичного реферування зібраних даних, формування запитів і чернеток звітів. Ефективність екосистеми визначається не окремим інструментом, а зрілістю методології та якістю інтеграції джерел і верифікації результатів.

Виявлення дезінформації та deepfake. Генеративний ШІ створив нову категорію викликів: фабриковані зображення, deepfake-відео, синтетичні голоси, AI-генеровані «свідчення». ENISA фіксує, що на початок 2025 року вже понад 80% спостережуваних соціально-інженерних кампаній спираються на ШІ для генерації контенту; з'явилися «зловмисні» моделі (Xanthorox AI) для масової автоматизації таких атак [7]. Контрзаходи поєднують ШІ-класифікатори згенерованого контенту, аналіз метаданих, верифікацію джерел і ланцюгів поширення. Показовим є український досвід: платформи Osavul і Mantis Analytics, створені у перші місяці повномасштабного вторгнення, застосовують LLM- та NLP-моделі для виявлення нарративів, координованої неавтентичної поведінки (coordinated inauthentic behaviour) і ботомереж у режимі, близькому до реального часу. Модулі CommSecure (виявлення нарративів) і CIB Guard (аналіз координації акаунтів) дозволяють викривати масштабні російські інформаційні операції на ранніх стадіях; ключова перевага - швидкість, за якої правдива інформація встигає досягти аудиторії раніше за ворожий нарратив [8]. Український досвід засвідчує, що ефективна протидія AI-дезінформації потребує поєднання трьох компонентів - запобігання, виявлення і реагування [8].

Виклики та обмеження. Попри прогрес, ШІ в OSINT має низку принципових обмежень. **По-перше** – галюцинації LLM-агентів: вигадані факти, неіснуючі джерела, фальшиві посилання, помилкова атрибуція цитат; як показано в [9], це не дефект конкретної моделі, а внутрішня властивість самого процесу навчання й оцінювання мовних моделей. **По-друге** - adversarial-загрози: отруєння публічних даних, prompt injection через скомпрометовані веб-сторінки, evasion-атаки на класифікатори – формалізовано в таксономії NIST AI 100-2 E2025 [10]. **По-третє** – операційна безпека: запити до публічних LLM-сервісів самі по собі є витком оперативної інформації аналітика. **По-четверте** – правові й етичні рамки: робота з персональними даними, біометрією й обличчями має узгоджуватись з GDPR, EU AI Act (з диференційованими обов'язками для високоризикових систем) [11] та національним законодавством. Систематичний огляд Ghioni та ін. формалізує управлінсько-етичні (GELSI) обмеження AI-OSINT і виявляє суттєві прогалини в нормативному регулюванні [1]. Український контекст ризиків застосування ШІ в безпекових задачах розглянуто у [12].

Висновки. ШІ якісно змінив можливості OSINT – від швидкості й масштабу збору до глибини аналізу - і фактично став обов'язковим компонентом сучасної аналітичної практики в безпековій, корпоративній та журналістській сферах. Водночас він не замінює аналітика: критичне мислення, верифікація, контекстне розуміння й операційна обережність залишаються прерогативою людини. Безпечне впровадження ШІ в OSINT потребує: збереження «людини в контурі» для всіх верифікаційних рішень; систематичної перевірки джерел, отриманих через LLM, проти первинних даних; розмежування публічних і приватних AI-сервісів за чутливістю запитів; інтеграції методології NIST AI RMF та таксономій NIST AI 100-2 / MITRE ATLAS у процеси оцінювання інструментарію; формування національних освітніх стандартів з AI-OSINT і протидії синтетичній дезінформації. Перспективними напрямками є розвиток пояснюваних і верифікованих агентних архітектур, методів виявлення синтетичного контенту та національної екосистеми підготовки аналітиків.

Список використаних джерел

1. Ghioni R., Taddeo M., Floridi L. Open source intelligence and AI: a systematic review of the GELSI literature. *AI & Society*. 2023. DOI: 10.1007/s00146-023-01628-x.
2. The OSINT Market: AI Integration Shaping Intelligence in 2025 / SpecialEurasia. March 2025. URL: <https://www.specialeurasia.com/2025/03/24/osint-market-ai-integration/> (дата звернення: 20.05.2026).
3. Browne T. O., Abedin M., Chowdhury M. J. M. A systematic review on research utilising artificial intelligence for open source intelligence (OSINT) applications. *International Journal of Information Security*. 2024. Vol. 23, Issue 4. P. 2911–2938. DOI: 10.1007/s10207-024-00868-2.
4. Karabiyik U., Ozer A., Lima V. Generative AI and Open-Source Intelligence: Exploring Capabilities and Privacy Implications / Purdue University, CERIAS Annual Security Symposium 2024. 2024. URL: <https://www.cerias.purdue.edu/symposium/index.php/posters/year/2024/R00-I3A> (дата звернення: 20.05.2026).
5. Wild J. Bellingcat's Online Investigation Toolkit / Bellingcat. 2024. URL: <https://bellingcat.gitbook.io/toolkit> (дата звернення: 20.05.2026).
6. S&P Global Market Intelligence Enhances Entity Insights Solution with Addition of Sanctions and Adverse Media Screening Data / S&P Global Market Intelligence. 24.01.2024. URL: <https://press.spglobal.com/2024-01-24-S-P-Global-Market-Intelligence-Enhances-Entity-Insights-Solution-with-Addition-of-Sanctions-and-Adverse-Media-Screening-Data> (дата звернення: 20.05.2026).
7. ENISA Threat Landscape 2025 / European Union Agency for Cybersecurity. October 2025. 87 p. URL: <https://www.enisa.europa.eu/publications/enisa-threat-landscape-2025> (дата звернення: 20.05.2026).
8. Marushchak A., Petrov S., Khoperiya A. Countering AI-powered disinformation through national regulation: learning from the case of Ukraine. *Frontiers in Artificial Intelligence*. 2025. Vol. 7, Art. 1474034. DOI: 10.3389/frai.2024.1474034.
9. Kalai A. T., Nachum O., Vempala S. S., Zhang E. Why Language Models Hallucinate. arXiv:2509.04664v1. 2025. URL: <https://arxiv.org/pdf/2509.04664> (дата звернення: 20.05.2026).
10. Vassilev A., Oprea A., Fordyce A., Anderson H., Davies X., Hamin M. Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations / National Institute of Standards and Technology. NIST AI 100-2 E2025. Gaithersburg, MD, 2025. DOI: 10.6028/NIST.AI.100-2e2025.
11. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union. L, 2024/1689, 12.07.2024. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (дата звернення: 20.05.2026).
12. Загрози та ризики застосування технологій штучного інтелекту. Кібербезпека: освіта, наука, техніка. 2025. URL: <https://csecurity.kubg.edu.ua/index.php/journal/article/view/520> (дата звернення: 20.05.2026).

Євгенія ЛИТВИНЕНКО, асистент кафедри кримінально-правових дисциплін та судочинства ННІ права, Сумський державний університет, доктор філософії в галузі права

ВИКОРИСТАННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ У ДІЯЛЬНОСТІ ПРАВООХОРОННИХ ОРГАНІВ УКРАЇНИ: ВИКЛИКИ ПРАВОВОГО РЕГУЛЮВАННЯ ТА ГАРАНТІЇ ПРАВ ЛЮДИНИ

У сучасних умовах цифрової трансформації публічного управління та сектору безпеки використання систем штучного інтелекту у діяльності правоохоронних органів набуває особливої актуальності, оскільки такі технології дають змогу обробляти великі масиви даних, виявляти закономірності, автоматизувати рутинні процеси та підвищувати ефективність аналітичної роботи. У європейському підході штучний інтелект розглядається як технологія, що має впроваджуватися на засадах безпечності, людиноцентричності та поваги до основоположних прав, а Регламент (ЄС) 2024/1689 [1] формує ризик-орієнтовану модель правового регулювання таких систем. Для правоохоронної сфери це має особливе значення, адже саме тут використання алгоритмічних рішень безпосередньо пов'язане з приватністю, недискримінацією, справедливим процесом та межами державного втручання у права людини.

Штучний інтелект уже застосовується або може застосовуватися у правоохоронній діяльності для кримінального аналізу, обробки великих і складних наборів даних, автоматизованого аналізу текстів, відеоспостереження, розпізнавання образів, цифрової криміналістики та підтримки прийняття рішень. У доповіді Europol наголошено, що інструменти ШІ здатні підвищити результативність аналізу інформації, виявлення підозрілих зв'язків, прогнозування ресурсних потреб та роботи з неструктурованими даними. Водночас такі технології не є нейтральними: вони можуть відтворювати вбудовані в дані упередження, створювати ризики помилкової ідентифікації, посилювати непрозорість управлінських і процесуальних рішень та ускладнювати реалізацію права особи на ефективний захист [2].

Особливу увагу в контексті правового регулювання привертає застосування ШІ у сферах, де можливе істотне втручання у фундаментальні права. Європейська комісія прямо вказує, що AI Act забороняє окремі практики неприйнятної ризику, зокрема певні форми дистанційної біометричної ідентифікації у публічно доступних місцях для правоохоронних цілей, а також відносить низку правоохоронних застосувань ШІ до високоризикових. Дотаким високоризиковим систем належать, зокрема, рішення, що можуть впливати на оцінку доказів, біометричну ідентифікацію, категоризацію осіб та інші інструменти, здатні впливати на реалізацію основоположних прав. Для них встановлено суворі вимоги щодо якості даних, документації, людського нагляду, кібербезпеки, точності, простежуваності та оцінки ризиків.

Для України ця проблематика є особливо важливою з огляду на поступове наближення національного законодавства до стандартів Європейського Союзу та необхідність одночасного забезпечення ефективності правоохоронної діяльності й дотримання конституційних прав людини. Чинний Закон України «Про захист персональних даних» [3] визначає, що правове регулювання обробки персональних даних спрямоване на захист фундаментальних прав і свобод людини, насамперед права на невтручання в особисте і сімейне життя у зв'язку з обробкою персональних даних. Однак наявні норми не охоплюють усіх специфічних ризиків, що виникають у зв'язку з використанням автономних або напіваавтономних алгоритмічних систем у правоохоронній практиці, зокрема ризиків непрозорості моделей, складності пояснення результату та проблем розподілу відповідальності між розробником, постачальником і користувачем системи.

Окрему небезпеку становлять ризики дискримінації та порушення принципу рівності. Рада Європи у Рекомендації CM/Rec(2020)1 наголошує, що держави повинні забезпечити, аби проєктування, розроблення та застосування алгоритмічних систем здійснювалися з дотриманням прав людини, принципів законності, прозорості, передбачуваності, підзвітності

та контролю. У документі підкреслюється, що алгоритмічні системи можуть не лише відтворювати, а й посилювати вже наявні соціальні дисбаланси, а тому держава має створювати ефективні нормативні й наглядові механізми для запобігання порушенням прав людини. Це положення є особливо релевантним для правоохоронної сфери, де похибка системи або викривлення даних може призвести до незаконного профілювання, необґрунтованої підозри чи надмірного спостереження за окремими групами осіб [4].

Не менш важливим є забезпечення процесуальних гарантій у випадках використання ШІ під час прийняття рішень у кримінальному провадженні чи оперативно-службовій діяльності. У доповіді Europol наголошується, що підзвітність, прозорість і пояснюваність мають принципове значення не лише для етичного застосування ШІ, а й для того, щоб результати, отримані за допомогою таких систем, витримували перевірку в суді та не порушували право на справедливий судовий розгляд. Саме тому людина не повинна бути усунута з процесу ухвалення значущих рішень, а використання алгоритмічних інструментів має мати допоміжний, а не остаточний характер у питаннях, що істотно впливають на права, свободи та законні інтереси особи. У цьому контексті доцільно запроваджувати обов'язкову оцінку впливу на права людини, внутрішній та зовнішній аудит систем, фіксацію логіки функціонування інструменту, а також право особи на отримання інформації про застосування автоматизованих засобів у конкретній справі.

Перспективи впровадження ШІ у діяльність правоохоронних органів України пов'язані не лише з технологічними можливостями, а й із формуванням належної моделі правового врядування. Така модель повинна спиратися на принципи законності, пропорційності, цільового обмеження, мінімізації даних, людського контролю та ефективного засобу юридичного захисту. Доцільним видається розроблення спеціальних нормативних положень щодо використання високоризикових систем ШІ у правоохоронній сфері, гармонізація національних стандартів із європейськими підходами, а також підготовка працівників правоохоронних органів до відповідального використання таких технологій. Лише за умови поєднання інноваційності та правових гарантій штучний інтелект може стати інструментом підвищення ефективності правоохоронної діяльності без порушення прав людини та принципів демократичної правової держави.

Список використаних джерел

1. 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). Official Journal of the European Union. 2024. URL: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
2. Europol. AI and policing: The benefits and challenges of artificial intelligence for law enforcement. Luxembourg: Publications Office of the European Union, 2024. 56 p. URL: <https://www.europol.europa.eu/cms/sites/default/files/documents/AI-and-policing.pdf>
3. Про захист персональних даних: Закон України від 01.06.2010 № 2297-VI. *Відомості Верховної Ради України*. 2010. № 34. ст. 481.
4. Recommendation CM/Rec(2020)1 of the Committee of Ministers to member States on the human rights impacts of algorithmic systems. Council of Europe, 8 April 2020. URL: <https://hcad.am/en/recommendation-coe-2/>

Андрій САВЧУК, заступник начальника кафедри тактико-спеціальних дисциплін, Національна академія Державної прикордонної служби України імені Богдана Хмельницького

ІНТЕГРАЦІЯ ШТУЧНОГО ІНТЕЛЕКТУ В АЛГОРИТМИ СОРТУВАННЯ ПОРАНЕНИХ НА ПОЛІ БОЮ: ПЕРСПЕКТИВИ, ОБМЕЖЕННЯ ТА ЕТИЧНІ ДИЛЕМИ ПРИЙНЯТТЯ КЛІНІЧНИХ РІШЕНЬ

Інтеграція технологій штучного інтелекту в алгоритми сортування поранених на полі бою постає як системна відповідь на трансформацію характеру сучасних збройних конфліктів, що супроводжуються високою інтенсивністю уражень, домінуванням комбінованої та множинної травми, а також суттєвим дефіцитом часу і ресурсів для прийняття клінічних рішень. Існуючі протоколи медичного сортування розроблені в межах концепцій Tactical Combat Casualty Care (TCCC) та масових санітарних втрат MASS CASUALTY, що працюють за певними стандартами, не завжди встигають підлаштуватися під динамічну обстановку сучасного бою і їм бракує «гнучкості». Тому вони потребують постійного перегляду і внесення певних змін. А зробити це можна, залучивши інтелектуальні системи підтримки прийняття рішень — цифрові інструменти, які допомагають аналізувати ситуацію, прогнозувати ризики та підказувати оптимальні дії.

З позицій тактичної медицини, процес сортування це не просто механічний вибір, а складна задача, де рішення приймається в екстремальних умовах, коли інформації бракує, а часу майже немає.

Штучний інтелект, аналізуючи усю наявну інформацію про пораненого (життєві показники (пульс, тиск, дихання), як саме сталася травма, скільки часу минуло від моменту ураження, результати швидкого огляду, дані від сенсорів чи приладів, тощо) створює модель-прогноз, яка показує, як може розвиватися травма — чи стан погіршиться, які ускладнення можливі, що треба зробити першочергово. Це відкриває можливість змінити підхід до сортування поранених, орієнтований не лише на поточний стан пораненого, а й на прогноз його подальшого стану.

Одним із ключових напрямів є впровадження систем раннього попередження критичних станів, що базуються на аналізі прихованих змін у фізіологічних показниках людини (наприклад, пульс, тиск, дихання). Алгоритми здатні ідентифікувати доклінічні ознаки геморагічного шоку, гіпоксії або травматичного ураження центральної нервової системи задовго до проявів симптомів, що принципово змінює логіку медичного втручання. В умовах обмежених ресурсів це дозволяє здійснювати більш раціональний розподіл сил і засобів, зосереджуючи їх на поранених із найбільшою ймовірністю позитивного результату при своєчасному втручанні.

Перспективною є також інтеграція штучного інтелекту з носимими біосенсорними платформами, які забезпечують безперервний моніторинг життєвих функцій у реальному часі. Це створює передумови для формування автоматизованого сортування поранених. Додатково, застосування технологій комп'ютерного зору дозволяє здійснювати первинну оцінку характеру зовнішніх ушкоджень, ідентифікацію джерел кровотечі та навіть приблизну оцінку об'єму крововтрати, що є критично важливим у перші хвилини після ураження.

Окремого значення набуває використання алгоритмів оптимізації в системах медичної евакуації. Штучний інтелект здатен враховувати не лише медичні параметри, але й тактичну обстановку, доступність евакуаційних маршрутів, ризики повторного ураження та часові обмеження «золотої години». Це дозволяє формувати адаптивні логістичні моделі, які мінімізують час доставки поранених до відповідного рівня медичної допомоги та підвищують загальну ефективність системи медичного забезпечення.

Водночас, попри значний потенціал, впровадження штучного інтелекту в тактичну медицину супроводжується низкою фундаментальних обмежень. Передусім, це проблема якості навчальних даних. Бойова травма відрізняється від цивільних травматологічних кейсів, на яких здебільшого тренуються існуючі моделі. Це створює ризик зниження точності прогнозів у реальних бойових умовах.

Технічні аспекти також відіграють критичну роль. Функціонування систем штучного інтелекту в умовах радіоелектронної протидії, обмеженого доступу до енергоресурсів та нестабільного зв'язку потребує створення автономних, енергоефективних і стійких до зовнішніх впливів рішень. Крім того, необхідно враховувати кібербезпекові ризики, пов'язані з можливістю втручання в алгоритми або підміни даних, що може призвести до катастрофічних наслідків.

Узагальнюючи, слід зазначити, що інтеграція штучного інтелекту в алгоритми сортування поранених на полі бою має потенціал стати якісним проривом у системі медичного забезпечення військ. Однак її ефективна реалізація можлива лише за умов міждисциплінарної синергії, що поєднує клінічну експертизу, інженерні рішення, етичні стандарти та нормативно-правове регулювання. Концептуально доцільним є формування гібридної моделі, у якій штучний інтелект виступає інструментом підсилення клінічного мислення, а остаточне рішення залишається за медичним фахівцем. Такий підхід дозволяє зберегти баланс між технологічною ефективністю та гуманістичними засадами медицини, що є принципово важливим у контексті збереження життя в умовах війни.

Дмитро КОБИЛИНСЬКИЙ, курсант 2-го курсу ФПФПКП НПУ, Дніпровський державний університет внутрішніх справ.

Науковий керівник:

Едуард РИЖКОВ, професор кафедри інформаційних технологій, Дніпровський державний університет внутрішніх справ, кандидат юридичних наук, професор

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У ДІЯЛЬНОСТІ ПІДРОЗДІЛІВ СТРАТЕГІЧНИХ РОЗСЛІДУВАНЬ

Сучасний етап розвитку суспільства характеризується стрімкою цифровізацією, що суттєво трансформує підходи до здійснення правоохоронної діяльності. В умовах зростання рівня організованої злочинності, поширення корупційних схем, кіберзагроз та транснаціональних кримінальних мереж підрозділи стратегічних розслідувань потребують нових інструментів аналітичної та оперативної роботи. Одним із таких інструментів є штучний інтелект, який здатний підвищити ефективність виявлення, документування та попередження складних кримінальних правопорушень.

Штучний інтелект (ШІ) у правоохоронній сфері розглядається як сукупність алгоритмічних рішень, що забезпечують автоматизовану обробку великих масивів даних, виявлення закономірностей, прогнозування ризиків та підтримку прийняття управлінських рішень. Для підрозділів стратегічних розслідувань, діяльність яких спрямована на протидію організованій злочинності та корупції, особливого значення набуває можливість аналізу великих обсягів фінансової, комунікаційної та цифрової інформації в режимі реального часу.

Конкретними різновидами ШІ, які знаходять застосування в цих підрозділах, є системи комп'ютерного зору для відеоаналітики, алгоритми обробки природної мови для роботи з текстовими матеріалами, нейронні мережі для виявлення прихованих зв'язків та предиктивні моделі для прогнозування кримінальних ризиків. Таке поєднання технологій забезпечує комплексний підхід до виконання завдань стратегічних розслідувань.

Станом на сьогодні кримінальне процесуальне законодавство України не містить легального визначення штучного інтелекту, однак відповідне тлумачення закріплено в Європейській етичній хартії щодо використання штучного інтелекту в судових системах та їхньому середовищі, де штучний інтелект трактується як сукупність наукових підходів, теоретичних положень і технологічних рішень, спрямованих на відтворення когнітивних функцій людини за допомогою машинних систем [1].

На національному рівні зміст цього поняття розкрито в Концепції розвитку штучного інтелекту в Україні, відповідно до якої штучний інтелект визначається як впорядкована система інформаційних технологій, що дає змогу виконувати складні багаторівневі завдання на основі застосування наукових методів, дослідницьких підходів і алгоритмів обробки інформації, отриманої або сформованої під час функціонування таких систем, а також передбачає формування й використання власних баз знань, моделей прийняття рішень та механізмів досягнення визначених цілей [2].

Реалізація положень зазначеної Концепції запланована до 2030 року та охоплює, зокрема, сферу кримінальної юстиції, що передбачає подальший розвиток уже функціонуючих цифрових інструментів правосуддя, таких як Єдина судова інформаційно-телекомунікаційна система, Електронний суд і Єдиний реєстр досудових розслідувань, упровадження консультаційних сервісів на базі штучного інтелекту для розширення доступу до правової допомоги, використання аналітичних можливостей штучного інтелекту з метою запобігання суспільно небезпечним явищам, визначення ефективних заходів ресоціалізації засуджених на основі аналізу відповідних даних, а також можливість ухвалення судових рішень у справах незначної складності за згодою сторін із використанням результатів інтелектуального аналізу даних [2].

Одним із ключових напрямів застосування штучного інтелекту є аналітика великих даних. Сучасні кримінальні структури використовують складні фінансові механізми, офшорні юрисдикції, криптовалютні транзакції та мережеві комунікації для маскування незаконної діяльності. Алгоритми машинного навчання здатні ідентифікувати підозрілі транзакції, встановлювати зв'язки між суб'єктами, виявляти приховані мережеві структури та прогнозувати ймовірні напрями розвитку злочинної активності. Це суттєво скорочує час аналітичної роботи та підвищує її точність.

Важливим аспектом є використання штучного інтелекту у сфері кримінального аналізу. Упродовж останніх років питання інтеграції штучного інтелекту в кримінальні розслідування набуло значної актуальності у багатьох державах, оскільки такі технології здатні істотно підвищити ефективність процесів збирання, обробки та оцінки доказової інформації, забезпечуючи більш об'єктивний і системний підхід до досудового розслідування, водночас актуалізуючи проблематику захисту приватності, дотримання прав людини та етичних стандартів.

Інтелектуальні системи спроможні працювати з великими масивами даних, включаючи відеозаписи (із застосуванням відеоаналітики для автоматичного виявлення об'єктів, розпізнавання дій, ідентифікації номерних знаків та відстеження руху осіб), текстову інформацію, матеріали із соціальних мереж та інші цифрові джерела. Алгоритми машинного навчання дозволяють автоматизовано виявляти закономірності й приховані взаємозв'язки, що можуть свідчити про підготовку або вчинення кримінальних правопорушень. Окремим потужним напрямом є біометрична ідентифікація на основі ШІ, яка охоплює розпізнавання облич, райдужної оболонки ока, відбитків пальців та голосових характеристик. Завдяки цим технологіям підрозділи стратегічних розслідувань отримують змогу оперативно встановлювати осіб, що перетинають кордони, беруть участь у незаконних зібраннях або

фігурують в оперативних матеріалах, що значно прискорює розшукову роботу. Загалом такий інструментарій суттєво оптимізує діяльність правоохоронних органів у частині виявлення потенційних ризиків і комплексного аналізу криміногенної обстановки, оскільки дає змогу формувати моделі, які враховують не лише очевидні чинники, а й багатовимірні кореляції між різними параметрами, підвищуючи точність прогнозування та ускладнюючи можливості протидії з боку правопорушників. У цьому контексті використання штучного інтелекту відкриває нові можливості для реалізації превентивної функції держави, забезпечуючи оперативне реагування на трансформації кримінальної ситуації та підвищення рівня безпеки [3].

Як наголошує Кисельов А.О., належна організація та тактично виважене здійснення кримінального аналізу сприяє скороченню часових витрат працівників оперативних і слідчих підрозділів Національної поліції під час виконання покладених на них завдань, що безпосередньо впливає на підвищення результативності діяльності у сфері запобігання та протидії злочинності [4; 5, с. 149-152].

Крім того, застосування інтелектуальних аналітичних систем у межах кримінального аналізу дає змогу формувати прогностичні моделі щодо ймовірності вчинення правопорушень, спираючись на історичні дані та сукупність релевантних факторів, що дозволяє визначати потенційно вразливі території або соціальні контексти та завчасно впроваджувати превентивні заходи, спрямовані на недопущення злочинів [3].

Перспективним напрямом є розробка та впровадження предиктивного штучного інтелекту (ПШІ), який ґрунтується на математичних моделях прогнозування та алгоритмах машинного навчання. Аргументами на користь активного розвитку ПШІ є можливість істотно підвищити точність визначення потенційних зон ризику, завчасно прогнозувати активізацію злочинних груп і на цій основі раціонально розподіляти ресурси підрозділів стратегічних розслідувань. Крім того, такий інструмент забезпечує стратегічне планування з урахуванням не лише історичних даних, а й динамічної оперативної обстановки, що робить правоохоронну діяльність проактивною, а не лише реактивною. У такий спосіб штучний інтелект виступає потужним інструментом стратегічного управління.

У цьому ж контексті, крізь призму OSINT-діяльності, слід розглядати методологію SOCTA (Serious and Organised Crime Threat Assessment), яку Європол використовує для оцінювання загроз організованої злочинності. Оскільки SOCTA передбачає багаторівневий аналіз колосальних масивів інформації, саме поєднання методів розвідки на основі відкритих джерел та ШІ-алгоритмів дозволяє автоматизувати збір, класифікацію та первинну обробку даних. Це критично підвищує продуктивність аналітичних підрозділів. Застосування штучного інтелекту для аналізу OSINT-масивів при підготовці оцінок SOCTA нівелює вплив суб'єктивних чинників, пришвидшує виявлення нових кримінальних трендів та дозволяє формувати аргументовані прогнози щодо розвитку транснаціональних злочинних мереж. [6, с. 61; 7, с. 460].

Разом із тим впровадження технологій штучного інтелекту у діяльність підрозділів стратегічних розслідувань пов'язане з низкою правових обмежень. Застосування алгоритмічних систем повинно здійснюватися з дотриманням принципів законності, пропорційності та поваги до прав людини. Особливої уваги потребують питання захисту персональних даних, запобігання дискримінаційним алгоритмічним рішенням та забезпечення прозорості використання автоматизованих систем [8].

Важливим є також питання відповідальності за результати застосування штучного інтелекту. Алгоритмічні рішення не можуть підміняти собою процесуальні повноваження слідчих чи оперативних працівників. Штучний інтелект повинен виконувати допоміжну функцію, надаючи аналітичну підтримку та рекомендації, тоді як остаточні рішення мають прийматися уповноваженими посадовими особами.

Організаційний аспект впровадження штучного інтелекту передбачає необхідність підготовки кадрів, здатних ефективно використовувати цифрові інструменти. Працівники підрозділів стратегічних розслідувань повинні володіти базовими знаннями у сфері аналізу

даних, розуміти принципи роботи алгоритмів машинного навчання та оцінювати ризики їх застосування. Це потребує модернізації системи професійної підготовки та впровадження міждисциплінарних освітніх програм. Водночас сам штучний інтелект може використовуватися як дієвий інструмент для патентної діяльності та підготовки кадрів: інструменти збору інформації, імітаційні тренажери з елементами ШІ дозволяють моделювати складні оперативні ситуації, відпрацьовувати навички аналізу великих даних і прийняття рішень в умовах, максимально наближених до реальних. Це підвищує якість навчання співробітників стратегічних підрозділів, скорочує період адаптації до роботи з сучасними технологіями та сприяє більш ефективному впровадженню ШІ у повсякденну правоохоронну діяльність [9].

Таким чином, використання штучного інтелекту у діяльності підрозділів стратегічних розслідувань є перспективним напрямом модернізації правоохоронної системи. Застосування алгоритмічних рішень сприяє підвищенню якості аналітичної роботи, оптимізації управлінських процесів та ефективності протидії організований злочинності. Водночас впровадження таких технологій повинно здійснюватися з урахуванням правових гарантій, етичних стандартів та необхідності збереження провідної ролі людини у прийнятті рішень.

Список використаних джерел

1. European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and their environment. Adopted at the 31st plenary meeting of the CEPEJ (Strasbourg, 3-4 December 2018) URL: <https://rm.coe.int/ethical-charter-en-for-publication-4-december2018/16808f699c>
2. Концепція розвитку штучного інтелекту в Україні: Розпорядження Кабінету Міністрів України від 2 груд. 2020 р. № 1556-р. URL: <https://www.kmu.gov.ua/npas/proshvalennya-koncepciyi-rozvitku-shtuchnogo-intelektu-v-ukrayini-s21220>
3. Топчій В. В. Штучний інтелект у кримінальному законодавстві України як вид інтелектуального продукту. *Міжнародний юридичний вісник: актуальні проблеми сучасності (теорія та практика)*. 2018. № 1-2 (10-11). С. 156-164.
4. Кисельов А.О. Особливості здійснення кримінального аналізу на прикладі складання аналітичного звіту (за матеріалами УКА ГУНП в Дніпропетровській області). *Бюлетень з обміну досвідом роботи* №228/2021. РВВ ДНДІ. К.. 2021. С. 63-67.
5. Kyselov A.O. Combating illegal amber mining: peculiarities of conflict resolution. *Naukovyi Visnyk Natsionalnoho Hirnychoho Universytetu*, №2. Dnipro: 2019. P. 146-152.
6. Рижков Е.В. Перспективи комплексного використання автоматизованих систем (фреймворків) та моделей штучного інтелекту в процесі OSINT / Е.В. Рижков / Інформаційно-аналітичне забезпечення діяльності органів сектору безпеки і оборони України, наук.-практ. конф.: матер. Всеукр. науково-практ. конф. (м. Львів, 26 грудн. 2025 р.). / упорядник: Т.В. Магеровська. Львів: Львівськ. держ. ун-т внутр. справ, 2026. – С. 59-62.
7. Рижков Е.В. Перспективи використання штучного інтелекту при реалізації технології OSINT / Е.В. Рижков/ Міжнародна та національна безпека: теоретичні і прикладні аспекти : матеріали IX Міжнар. наук.-практ. конф. (м. Дніпро, 21 бер. 2025 р.) ; у 2-х ч. Дніпро : Дніпров. держ. ун-т внутр. справ, 2025. Ч. II. С. 459-461.
8. Рижков Е.В., Виноградова К.В. Використання штучного інтелекту підрозділами стратегічних розслідувань Національної поліції України / Е.В. Рижков, К.В. Виноградова / Штучний інтелект і безпека, наук.-практ. конф. Інституту проблем моделювання в енергетиці ім. Г.Є. Пухова Національної академії наук України, Інституту проблем реєстрації інформації Національної академії наук України : матеріали, 19-21 листопада 2024 р. Київ : ІПМЕ ім. Г.Є. Пухова НАН України, ІПРІ НАН України, 2024. С. 95-97.
9. Рижков Е.В. Інноваційні методи біометричної ідентифікації за клавіатурним почерком у системі заходів протидії організований злочинності (за матеріалами патенту № 130372) // Актуальні питання діяльності Національної поліції як суб'єкту протидії організований злочинності : Матер. Всеукр. наук.-практ. Конф. (в авторській редакції), (м. Кропивницький, 02 квітня 2026 року). Кропивницький, 2026. С. 286-290.

Анатолій ТІХОНОВ, здобувач спеціальності D8 «Право», Мелітопольський державний педагогічний університет імені Богдана Хмельницького, експерт Національного агентства із забезпечення якості вищої освіти.

Науковий керівник:

Ольга УСТЮЖАНІНОВА, доцент кафедри права, Мелітопольський державний педагогічний університет імені Богдана Хмельницького, кандидат юридичних наук

АКАДЕМІЧНА ДОБРОЧЕСНІСТЬ В УМОВАХ ВИКОРИСТАННЯ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ: ПРАВОВІ ПРОБЛЕМИ ТА ПЕРСПЕКТИВИ РЕГУЛЮВАННЯ

Стрімкий розвиток технологій генеративного штучного інтелекту (далі – ШІ) суттєво трансформував освітній та науковий простір. Великі мовні моделі здатні генерувати академічні тексти, що за формою не поступаються результатам людської інтелектуальної діяльності. Це зумовило виникнення принципово нових правових проблем у сфері академічної доброчесності, де усталені правові категорії виявилися недостатніми для регулювання нових суспільних відносин, а традиційні механізми контролю – малоефективними.

У класичному розумінні, закріпленому статтею 42 Закону України «Про освіту» від 05.09.2017 р. № 2145-VIII [1], академічний плагіат визначається як оприлюднення (частково або повністю) наукових (творчих) результатів, отриманих іншими особами, як результатів власного дослідження. Однак використання генеративного ШІ не вкладається в це визначення: здобувач освіти не привласнює працю іншої фізичної особи в юридичному сенсі, оскільки ШІ не є суб'єктом авторського права. Натомість відбувається привласнення результату функціонування алгоритму. Механізм створення контенту кардинально відмінний: якщо плагіат передбачає копіювання або перефразування чужої праці, то генеративний ШІ формує текст на основі статистичних патернів, без жодного акту запозичення у конкретної фізичної особи. Саме ця відмінність зумовила необхідність окремого правового регулювання.

Знаковою відповіддю держави на зазначені виклики стало ухвалення Закону України «Про академічну доброчесність» від 18.12.2025 р. № 4742-IX [2]. Вперше на законодавчому рівні запроваджено окремий вид академічної недоброчесності – «недоброчесне використання результатів штучного інтелекту», що розмежовується з плагіатом як самостійне правопорушення, прирівняне за серйозністю до фальсифікації. Це свідчить про перехід до функціонального розуміння доброчесності, де важлива не лише відсутність «крадіжки» у людини, а й наявність реального пізнавального процесу в самого здобувача. Згідно зі статтею 29 Закону № 4742-IX [2], особи, які використовують засоби ШІ при підготовці наукових та навчальних робіт, зобов'язані декларувати такий факт із зазначенням конкретних інструментів, завдань, для яких вони застосовувалися, та способів верифікації отриманих даних. Такий підхід, відомий у міжнародній практиці як «AI Transparency», дозволяє зберегти довіру до академічного процесу без тотальної заборони технологій.

З точки зору авторського права принципове значення має стаття 33 Закону України «Про авторське право і суміжні права» від 01.12.2022 р. № 2811-IX [3], яка закріплює правовий статус об'єктів, створених комп'ютерними програмами без безпосередньої участі фізичної особи. Такі об'єкти отримують особливий правовий режим (*suí generis*): вони позбавлені особистих немайнових прав (оскільки ШІ не є автором), а майнові права належать власнику програми або законному користувачу строком на 25 років. Простого натискання кнопки для виникнення авторства недостатньо – участь людини у творчому процесі має бути безпосередньою та змістовною.

Отже, результати генеративного ШІ в Україні мають обмежений юридичний захист, що цілком корелює з академічним принципом: науковий ступінь або диплом має здобуватися реальною людською інтелектуальною працею.

Університетська автономія надає закладам вищої освіти право деталізувати положення Закону № 4742-IX [2] через власні етичні кодекси та академічні політики. Аналіз чинних університетських регламентів дозволяє виокремити три зони допустимості використання ШІ в академічній роботі. До допустимої допомоги відносяться: пошук та систематизація джерел з обов'язковою їх перевіркою, стилістична корекція готового тексту, структурування ідей, переклад з подальшим редагуванням. Умовно допустимою є генерація ілюстративних матеріалів, написання фрагментів коду, аналіз великих масивів даних, підготовка окремих перехідних конструкцій з обов'язковим маркуванням. Недопустимим є: повна генерація змістовної частини роботи, фабрикація або підміна джерел, подання результатів ШІ як власних під час будь-яких контрольних заходів. Така градація – не формальність, а механізм збереження головного освітнього результату: розвитку критичного мислення та фахових компетентностей здобувача.

Окремої уваги заслуговує проблема технічного виявлення ШІ-згенерованих текстів. Сучасні системи детекції базуються на аналізі «перплексії» та «вибуховості» (burstiness) тексту, однак їхня практична надійність викликає серйозні наукові та правові застереження. Дослідження Girau, Roe та Espiritu переконливо доводять, що детектори ШІ систематично продукують хибнопозитивні результати, особливо щодо текстів, написаних особами, для яких мова написання не є рідною: граматично правильні, стилістично виважені конструкції статистично нагадують ШІ-генерований текст [4, с. 2, 7]. Автори доходять висновку, що застосування таких інструментів як єдиного доказу є несправедливим і потенційно дискримінаційним, кваліфікуючи можливі наслідки для невинних здобувачів як «академічну смертну кару» [4, с. 8]. Деер та ін. зафіксували критичне падіння точності детекторів після застосування методів «гуманізації» тексту – з 70,3 % до 4,6 %, а також підтверджують, що хибнопозитивні спрацювання прямо порушують принципи справедливості та належної правової процедури [5, с. 16–19].

Зазначені дані мають принципове значення для правозастосовної практики. Застосування результатів роботи детектора ШІ як єдиного і достатнього доказу академічного порушення є неприйнятним з точки зору принципів справедливості та належної правової процедури. Рішення про академічну недоброчесність повинна ухвалювати компетентна людина за результатами комплексної оцінки всіх обставин, а не автоматизована система. Цей підхід відповідає вимогам Регламенту (ЄС) 2024/1689 Європейського Парламенту та Ради (EU AI Act) [6], який прямо відносить освітні системи автоматизованого оцінювання та моніторингу під час іспитів до сфер високого ризику і вимагає обов'язкового людського нагляду: жодна система ШІ не може автоматично, без реальної участі людини, ухвалювати рішення, що істотно впливають на права та освітню траєкторію особи.

У міжнародно-правовому вимірі ЮНЕСКО у Рекомендації з етики штучного інтелекту від 23.11.2021 р. [7] та Керівництві з використання генеративного ШІ в освіті й дослідженнях 2023 р. [8] наголошує на необхідності захисту приватності, запобігання цифровому розриву та забезпечення інклюзії. ШІ має розширювати можливості учасників освітнього процесу, а не перетворювати навчання на суцільний цифровий нагляд. Особливо актуальним є застереження цих документів щодо нерівного доступу до технологій: відсутність у частини здобувачів навичок роботи з ШІ не повинна перетворюватися на конкурентний недолік, а наявність таких навичок – автоматично викликати підозру. Ці засади набувають характеру нових цифрових прав людини в сфері освіти [7; 8].

Перспективи вдосконалення правового регулювання у цій сфері охоплюють кілька напрямів. Доцільно розробити методичні рекомендації щодо порядку декларування використання ШІ та критеріїв оцінювання самостійності виконання завдань; врегулювати статус «Prompt Engineering» як нової форми інтелектуальної діяльності; закріпити процесуальні гарантії для здобувачів у справах про академічну недоброчесність (право

ознайомлюватися з методологією детекції, оспорювати висновки автоматизованих систем, право на усне слухання). Визначення переліку видів робіт, при підготовці яких використання ШІ є допустимим або неприпустимим за замовчуванням, усуне правову невизначеність і забезпечить рівні умови для всіх учасників освітнього процесу.

З урахуванням викладеного, можна сформулювати такі висновки. По-перше, ухвалення Закону № 4742-IX знаменує перехід від заборонної до транспарентної моделі регулювання ШІ в академічному середовищі: не карати за сам факт використання технологій, а вимагати прозорості та відповідальності. По-друге, технічні засоби детекції ШІ-контенту не можуть слугувати самодостатньою підставою для юридичних санкцій через їхню доведену ненадійність та дискримінаційний потенціал щодо окремих категорій здобувачів освіти. По-третє, ефективне регулювання вимагає не лише вдосконалення законодавства, а й системної перебудови освітніх практик: запровадження форматів оцінювання, резистентних до автоматизації (усні захисти, дискусії в режимі реального часу, аналіз реальних кейсів), цілеспрямованого розвитку критичної АІ-грамотності здобувачів і викладачів, а також проєктування педагогічних завдань, де ШІ виступає інструментом підтримки мислення, а не його заміником. Зрештою, академічна доброчесність в умовах цифрової трансформації освіти – це насамперед питання довіри, чесності та відповідальності, а не лише технічного контролю.

Список використаних джерел

1. Про освіту : Закон України від 05.09.2017 р. № 2145-VIII (зі змінами та доп.). URL: <https://zakon.rada.gov.ua/laws/show/2145-19> (дата звернення: 15.05.2026).
2. Про академічну доброчесність : Закон України від 18.12.2025 р. № 4742-IX. URL: <https://zakon.rada.gov.ua/laws/show/4742-ix> (дата звернення: 15.05.2026).
3. Про авторське право і суміжні права : Закон України від 01.12.2022 р. № 2811-IX. URL: <https://zakon.rada.gov.ua/laws/show/2811-20> (дата звернення: 15.05.2026).
4. Giray L., Roe J., Espiritu J. D. AI writing detectors are ineffective, unreliable and harmful. English Teaching: Practice & Critique. 2026. Epub ahead of print. DOI: <https://doi.org/10.1108/ETPC-07-2025-0155>.
5. Deep P. D., Edgington W. D., Ghosh N., Rahaman M. S. Evaluating the Effectiveness and Ethical Implications of AI Detection Tools in Higher Education. Information. 2025. Vol. 16. Art. 905. DOI: <https://doi.org/10.3390/info16100905>.
6. Регламент (ЄС) 2024/1689 Європейського Парламенту та Ради від 13.06.2024 про встановлення гармонізованих правил щодо штучного інтелекту (Artificial Intelligence Act). OJ L, 2024/1689, 12.07.2024. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (дата звернення: 15.05.2026).
7. UNESCO. Recommendation on the Ethics of Artificial Intelligence (UNESCO Doc. 41 C/73), adopted on 23 November 2021. Paris : UNESCO, 2021. URL: <https://unesdoc.unesco.org/ark:/48223/pf0000381137> (дата звернення: 15.05.2026).
8. UNESCO. Guidance for generative AI in education and research. Paris : UNESCO, 2023. 45 p. URL: <https://unesdoc.unesco.org/ark:/48223/pf000038669> 3 (дата звернення: 15.05.2026).

Ірина БАЗИЛЯК, здобувач ступеня вищої освіти бакалавра навчально-наукового експертно-криміналістичного інституту, Національна академія внутрішніх справ.

Науковий керівник:

Олена ВОРОБЕЙ, професорка кафедри криміналістичного забезпечення та судових експертиз навчально-наукового експертно-криміналістичного інституту, Національна академія внутрішніх справ, кандидат юридичних наук, доцент

ЗАСТОСУВАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ В OSINT-РОЗВІДЦІ

OSINT (open source intelligence – розвідка відкритих баз даних) можна визначити як діяльність із пошуку, збору та аналізу інформації, отриманої із загальнодоступних каналів, що не потребує застосування таємних методів збору [1, с. 7]. Ця методологія спрямована на ідентифікацію даних про фізичних та юридичних осіб, верифікацію фактів і протидію дезінформації. До джерел даних належать соціальні мережі, пошукові системи, публічні реєстри, супутникові знімки, форуми, новини та інші. Основними напрямками застосування такого виду розвідки є сфера національної безпеки та оборони (моніторинг загроз, пошук людей), розслідувальна журналістика (викриття фейків, верифікація контенту), а також кібербезпека (пошук витоків конфіденційних даних).

Зважаючи на стрімкі темпи цифрової трансформації, за яких обсяг інформації у світі подвоюється кожні два роки [2], традиційні методи OSINT, що покладаються виключно на людський ресурс, втрачають свою ефективність. За таких умов штучний інтелект (далі – ШІ) перетворює розвідку на автоматизовану систему. Завдяки здатності обробляти великі обсяги даних, аналізувати складні інформаційні структури та знаходити приховані зв'язки, ШІ здатен забезпечити якісно новий рівень відкритої розвідки.

Інструменти ШІ суттєво розширюють можливості розвідки, допомагаючи ідентифікувати людей чи локації, виявляти аномалії в медіапросторі та оперативно збирати дані в реальному часі. Насамперед ШІ повністю перебирає на себе рутинну роботу – від цілодобового моніторингу соцмереж і форумів до відстеження найменших змін на вебсторінках. Разом із цим, аналізуючи історичні дані, алгоритми забезпечують здатність передбачати ймовірний розвиток подій, що допомагає вчасно запобігати конфліктам та гнучко вибудовувати стратегії. Паралельно з цим технології NLP (Natural Language Processing) дозволяють здійснювати глибокий аналіз текстів, миттєво виділяючи головні ідеї з тисяч документів і помічати приховані контексти. На додачу, ШІ ефективно працює з візуальними даними, виконуючи автоматичне розпізнавання облич та об'єктів, точну геолокацію за фото, а також виявлення цифрових маніпуляцій і фейків у відеоматеріалах [3].

Практична реалізація розглянутих можливостей ШІ в OSINT-розвідці забезпечується спеціалізованим програмним забезпеченням. Зокрема, для аналізу соціальних мереж та кібербезпеки активно використовується Maltego. Це потужний інструмент для відображення зв'язків між об'єктами, особами та організаціями, IP-адресами або доменами. Його унікальність полягає в тому, що він автоматично знаходить, а потім структурує закономірності в даних. Datamining найкраще підходить для моніторингу змін даних у режимі реального часу. Інструмент використовує алгоритм, який аналізує величезні обсяги інформації та надсилає сповіщення у разі виявлення загрози, що дозволяє швидше реагувати на небезпеки та прогнозувати кібератаки. Платформа Skopenow найкраще підходить для проведення розслідувань у соціальних мережах та перевірки особи. Тут використовуються алгоритми машинного навчання для розпізнавання важливих даних із соцмереж: імена, місцезнаходження, зв'язки. Це допомагає автоматизувати процес профілювання, пов'язуючи

різних онлайн-персон з однією особою. Високу ефективність демонструє сервіс Ovis, основною функцією якого є розпізнавання об'єктів у великих масивах відео та фотографій. Інструмент CarNet.AI здатен визначити модель, покоління та марку автомобілів, випущених з 1995 року з точністю в 97%. Він може покращити фотографії навіть низької якості та зіставити характеристики автомобіля з понад 3100 моделями в базі даних [4].

Попри очевидні технологічні переваги та широкий спектр інструментарію, впровадження ІІІ в процеси розвідки на основі відкритих джерел має певні проблеми. Великі мовні моделі здатні генерувати хибні або викривлені кореляції, приймаючи випадкові збіги за закономірності. Без перевірки людиною-аналітиком такі дані можуть призвести до дезінформації на стратегічному рівні. Розвиток технологій створення дипфейків та пропагандистського контенту (ІІІСО) бот-мережами значно ускладнюють процес первинної верифікації. Також лишаються відкритими правові та етичні аспекти. Автоматизований збір персональних даних користувачів змушує аналітиків балансувати між правом на приватність та безпекою держави. У правовому полі такі дії вимагають відповідного нормативного регулювання, адже будь-яке розслідування повинно залишатися в межах закону.

ІІІ докорінно трансформував методику роботи з відкритими даними. В умовах «інформаційного вибуху» та втрати традиційними методами відкритої розвідки своєї ефективності, різноманітні інструменти дозволяють автоматизувати рутину, пришвидшити обробку інформації та виявляти приховані закономірності [5]. Використання спеціалізованого програмного забезпечення на основі ІІІ може в рази підвищити ефективність роботи аналітиків, військовослужбовців, журналістів та фахівців з кібербезпеки, мінімізуючи рутинне навантаження та забезпечуючи високу швидкість реагування на загрози.

Водночас слід зазначити, що технології ІІІ є дещо обмеженими з огляду наявності ризиків алгоритмічних помилок, та залишаються виключно інструментом. Ключовим фактором успіху сучасних OSINT-досліджень є концепція синергії людини та технологій: ІІІ забезпечує швидкість і масштаб обробки даних, тоді як людина – критичне мислення, перевірку та ухвалення рішень.

Список використаних джерел

1. OSINT при розслідуванні кримінальних правопорушень : підручник / О. О. Торбас. Одеса : Видавництво «Юридика», 2024. 180 с.
2. Інформаційний потоп: який ковчег нам будувати? *ITC*. веб-сайт. URL: <https://itc.ua/ua/blogs/informatsijnyj-potop-yakyj-kovcheg-dopomozhe/amp/>
3. Штучний інтелект в OSINT: Як покращити розслідування у 2024 році. *Infolight.ua*. веб-сайт. URL: <https://infolight.ua/2024/12/01/shtuchnyj-intelekt-v-osint-yak-pokrashhyty-rozsliduvannya-u-2024-rotsi/>
4. OSINT AI: How to Optimize Your Investigation in 2024. *Molfar*. Веб-сайт. URL: <https://molfar.com/news-posts/osint-ai-how-to-optimize-your-investigation-in-2024>
5. Шевчук В. М. Використання технологій штучного інтелекту в OSINT як напрям підвищення ефективності розслідування воєнних злочинів / В. Шевчук // Роль OSINT-досліджень у підвищенні рівня національної безпеки України : матеріали круглого столу (м. Львів, 7 трав. 2025 р.) / уклад.: І. О. Ревак. Львів, 2025. С. 239–243.

Максим КОРШУНОВ, студент 2-го курсу, спеціальність 121 «Інженерія програмного забезпечення», група ПК-22, Дрогобицький механіко-технологічний фаховий коледж.

Науковий керівник:

Оксана БЕРЕГУЛЯК, старший викладач циклової комісії природничо-математичних дисциплін, Дрогобицький механіко-технологічний фаховий коледж

ВИКОРИСТАННЯ СИСТЕМ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ АВТОМАТИЗОВАНОГО АНАЛІЗУ ВРАЗЛИВОСТЕЙ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ

В умовах стрімкого розвитку інформаційних технологій та масштабної цифровізації суспільства питання забезпечення надійності та безпеки програмного забезпечення набуває критичного значення. Зростання кількості кібератак, складність сучасних архітектур та використання сторонніх бібліотек створюють нові вектори загроз. Традиційні методи тестування та пошуку вразливостей, такі як статичний (SAST) та динамічний (DAST) аналіз коду, часто виявляються недостатньо ефективними через велику кількість хибнопозитивних спрацьовувань та нездатність виявляти складні логічні помилки [1]. У цьому контексті інтеграція систем штучного інтелекту (ШІ) та машинного навчання відкриває нові перспективи для автоматизації аналізу вразливостей.

Штучний інтелект у сфері кібербезпеки та інженерії програмного забезпечення здатний аналізувати великі масиви коду, виявляти патерни, які вказують на потенційні уразливості, та пропонувати оптимальні шляхи їх усунення. На відміну від класичних аналізаторів, які спираються на жорстко задані правила, моделі на базі глибокого навчання (Deep Learning) та великі мовні моделі (LLM) здатні розуміти контекст виконання програми [2].

Основні напрямки застосування ШІ для аналізу коду:

- **Інтелектуальний статичний аналіз:** Використання нейромереж для класифікації ділянок коду на безпечні та потенційно вразливі на основі попередньо навчених наборів даних (баз CVE).
- **Автоматична генерація експлойтів та патчів:** Системи ШІ здатні не лише знаходити помилки, але й генерувати фрагменти безпечного коду для їх виправлення.
- **Аналіз поведінкових аномалій:** Під час динамічного тестування алгоритми машинного навчання аналізують поведінку програми під час виконання.

Порівняльний аналіз традиційних та інтелектуальних систем пошуку вразливостей

Критерій	Традиційні SAST/DAST системи	Аналізатори на базі ШІ
Принцип роботи	Сигнатурний аналіз, жорсткі правила	Семантичний аналіз, розпізнавання патернів
Рівень хибнопозитивних спрацьовувань	Високий (потребує ручної перевірки)	Низький/Середній
Здатність до адаптації	Низька (потребує ручного оновлення)	Висока (навчається на нових вразливостях)

Проте впровадження систем штучного інтелекту в правоохоронній та безпековій діяльності має супроводжуватися урахуванням специфічних ризиків. По-перше, існує проблема "отруєння даних", коли зломисники можуть навмисно модифікувати навчальні вибірки для того, щоб модель ігнорувала певні типи вразливостей. По-друге, моделі ШІ самі по собі можуть бути вразливими до атак (Adversarial attacks) [3].

Висновки. Використання можливостей систем штучного інтелекту для аналізу вразливостей програмного забезпечення є перспективним напрямком, що дозволяє суттєво підвищити рівень кібербезпеки. Вони автоматизують рутинні процеси тестування, знижують навантаження на інженерів з безпеки та дозволяють виявляти складні логічні помилки. Однак для успішного використання таких систем у правоохоронній діяльності необхідно забезпечити прозорість алгоритмів прийняття рішень та створити надійні механізми захисту.

Список використаних джерел

1. Карпенко О. В. Штучний інтелект у кібербезпеці: сучасний стан та перспективи розвитку. *Інформаційні технології та захист інформації*. 2023. № 2. С. 45–52.
2. Buczak A. L., Guven E. A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection. *IEEE Communications Surveys & Tutorials*. 2016. Vol. 18, no. 2. P. 1153–1176.
3. Мельник В. А. Використання нейронних мереж для статичного аналізу коду. *Науковий вісник кібербезпеки*. 2024. Вип. 1. С. 112–119.

Антон ЛИСЕНКО, курсант ННІ 2-25-104
Пд, Харківський національний університет
внутрішніх справ, рядовий поліції.

Науковий керівник:

Олександр СКЛЯР, старший викладач
кафедри тактико-спеціальної підготовки,
Харківський національний університет
внутрішніх справ, капітан поліції, доктор
філософії

ШТУЧНИЙ ІНТЕЛЕКТ У ПРАВООХОРОННІЙ ДІЯЛЬНОСТІ: МОЖЛИВОСТІ ТА РИЗИКИ ЗАСТОСУВАННЯ

Використання штучного інтелекту в правоохоронній діяльності є одним із найбільш перспективних і водночас ризикованих напрямів цифрової трансформації держави. У правоохоронній діяльності штучний інтелект (далі ШІ) використовують для швидкої обробки великої кількості інформації, пошуку закономірностей, аналізу даних і допомоги працівникам поліції під час ухвалення рішень. У сучасних умовах ШІ може застосовуватися не як заміна слідчого, оперативного працівника, прокурора чи суду, а як допоміжний інструмент. Він підвищує швидкість аналізу інформації, покращує координацію підрозділів і дає змогу ефективніше реагувати на злочинність. При цьому використання ШІ повинно відповідати вимогам законодавства, не порушувати права людини та перебувати під контролем уповноважених органів.

Одним із найбільш практичних напрямів застосування ШІ в правоохоронній діяльності є аналітична обробка даних. Щодня правоохоронні органи опрацьовують велику кількість інформації: матеріали справ, відео, бази даних, повідомлення в соціальних мережах та інші джерела. ШІ може допомагати швидше аналізувати ці дані та знаходити зв'язки між подіями

й особами. Наприклад, алгоритми можуть допомагати у виявленні серійних правопорушень, шахрайських схем, кіберзлочинів, незаконних фінансових операцій, підроблених документів або повторюваних моделей поведінки правопорушників. INTERPOL також наголошує, що використання ІІІ в поліції повинно відповідати правам людини та етичним принципам [1, с. 91].

Одним із поширених напрямів використання ІІІ є відеоспостереження та розпізнавання об'єктів. Системи комп'ютерного зору можуть аналізувати відео з камер спостереження, розпізнавати номерні знаки транспортних засобів, фіксувати підозрілі предмети, виявляти залишені речі в громадських місцях, допомагати у пошуку зниклих осіб або встановленні маршруту руху певного транспортного засобу.

Однак найбільш дискусійним є використання технологій розпізнавання облич. З одного боку, вони можуть прискорити розшук осіб, ідентифікацію підозрюваних або потерпілих, а також допомогти в умовах масових подій чи надзвичайних ситуацій. Проте такі технології можуть помилятися та порушувати право людини на приватність, гарантоване Конституцією України. Саме тому в Європейському Союзі використання систем дистанційної біометричної ідентифікації в режимі реального часу у публічних місцях для правоохоронних цілей розглядається як особливо чутливе й допускається лише за суворих умов, зокрема за наявності спеціального дозволу та з урахуванням принципів необхідності й пропорційності [2, с. 261].

Окремий напрям становить прогнозна аналітика злочинності, тобто використання алгоритмів для визначення територій, часу або ситуацій із підвищеним ризиком вчинення правопорушень. Наприклад, система може аналізувати статистику попередніх правопорушень, соціально-економічні фактори, повторюваність подій, сезонність, транспортну активність і на основі цього пропонувати посилення патрулювання в певному районі. Проте такий підхід має суттєві ризики. Якщо система штучного інтелекту навчається на неповних, неточних або упереджених даних, вона може повторювати помилки, які вже існують у правоохоронній системі. У результаті поліція може частіше звертати увагу на окремі райони, соціальні групи чи певні категорії осіб навіть без достатніх підстав. Це створює ризик несправедливого ставлення та порушення прав людини. Тому рішення не повинні прийматися лише на основі прогнозів системи. Вона може бути лише інструментом планування ресурсів, але не самостійною підставою для обмеження прав особи, проведення обшуку, затримання або повідомлення про підозру [3, с. 46]. У кримінальному процесі ІІІ може застосовуватися для систематизації матеріалів провадження, пошуку релевантних документів, аналізу великих масивів електронних доказів, розшифрування аудіо- та відеозаписів, виявлення збігів у показаннях, підготовки аналітичних довідок.

Це особливо важливо у складних кримінальних провадженнях, пов'язаних з організованою злочинністю, корупцією, кіберзлочинами, фінансовими махінаціями або воєнними злочинами, де обсяг доказової інформації може бути надзвичайно великим. Водночас ІІІ не може самостійно визнавати особу винною, оцінювати докази замість суду чи підміняти процесуальну діяльність слідчого. Відповідно до принципів кримінального процесу, доказ має оцінюватися людиною – уповноваженою службовою особою або судом – з урахуванням належності, допустимості, достовірності та достатності. ІІІ може допомагати в аналізі даних, але остаточне рішення повинна приймати людина.

Особливу увагу потрібно приділити захисту персональних даних. Правоохоронна діяльність часто передбачає обробку чутливої інформації: біометричних даних, даних про місце перебування, комунікації, фінансові операції, соціальні зв'язки, судимість, стан здоров'я або інші персональні відомості. Використання ІІІ також створює ризик витоку даних або їх неправильного використання. Тому використання ІІІ в правоохоронних органах повинно чітко регулюватися законом і застосовуватися лише в межах визначених повноважень. У цьому контексті важливим є підхід Ради Європи, яка у Рамковій конвенції про штучний інтелект наголошує на тому, що системи ІІІ повинні використовуватися з дотриманням прав людини та вимог законодавства [4, с. 180].

Для України питання використання ІІІ в правоохоронній діяльності має особливе значення через умови воєнного стану, зростання кіберзагроз, необхідність документування

воєнних злочинів, протидію колабораційній діяльності, шахрайству, незаконному обігу зброї, дезінформації та іншим загрозам безпеці. ШІ може використовуватися для фіксації доказів, обробки супутникових знімків і пошуку інформації у відкритих джерелах. Також він допомагає виявляти інформаційні операції, аналізувати повідомлення громадян та встановлювати зв'язок між різними подіями. Однак саме в умовах війни ризик надмірного втручання держави в права людини є особливо високим, тому застосування таких технологій має супроводжуватися парламентським, судовим, прокурорським, відомчим і громадським контролем. Україна вже рухається до формування правового регулювання ШІ: Міністерство цифрової трансформації презентувало дорожню карту регулювання штучного інтелекту, яка передбачає поступову підготовку до майбутнього законодавства, наближеного до європейського AI Act.

Серед основних проблем використання ШІ можна виділити помилки систем, неточне розпізнавання людей і можливе порушення прав людини. Наприклад, якщо система помилково вказує на особу як на ймовірного правопорушника, виникає питання: хто відповідатиме за наслідки – розробник програми, орган, який її закупив, оператор системи чи посадова особа, яка використала результат? Тому остаточне рішення завжди повинна приймати людина. Кінцевий висновок про застосування примусових заходів, повідомлення про підозру, затримання, проведення слідчої дії або передачу справи до суду не може ухвалюватися виключно алгоритмом. ШІ має надавати допоміжну інформацію, а не створювати автоматизовану форму кримінального переслідування.

Отже, використання штучного інтелекту в правоохоронній діяльності має подвійний характер. З одного боку, ШІ здатний істотно підвищити ефективність боротьби зі злочинністю, пришвидшити аналіз доказів, оптимізувати розподіл ресурсів, покращити розшук осіб, посилити кібербезпеку та допомогти у документуванні складних правопорушень. З іншого боку, без належного правового регулювання він може стати інструментом необґрунтованого контролю, дискримінації, масового стеження та порушення основоположних прав людини. Тому використання ШІ у правоохоронній сфері має бути не лише ефективним, а й законним та безпечним для прав людини.

Список використаних джерел

1. Муравська Ю. Перспективи використання технологій штучного інтелекту у юриспруденції та правоохоронній діяльності. *Актуальні проблеми правознавства*. 2024. С. 90-97. URL: <https://appj.wunu.edu.ua/index.php/appj/article/view/1890/1939> (дата звернення 11.05.2026)
2. Вікторчук М. В., А. С. Багатко. Вітчизняний та міжнародний досвід використання технологій штучного інтелекту в правоохоронній діяльності. *Науковий вісник Ужгородського національного університету*. Серія: Право. № 86. 2024. С. 257-263. URL: <https://visnyk-juris-uzhnu.com/wp-content/uploads/2025/01/41-2.pdf> (дата звернення 11.05.2026).
3. Величко Д. В., В. Ю. Калугін. Перспективи використання штучного інтелекту в правоохоронній діяльності. Рекомендовано до друку та поширення через мережу Інтернет Вченою радою Львівського державного університету внутрішніх справ (протокол № 14 від 25 березня 2025 року). 2025. С. 46. URL: https://if.uu.edu.ua/wp-content/uploads/2025/05/14_03_2025.pdf (дата звернення 13.05.2026).
4. Троцький О. О. Правове регулювання та потенційні ризики застосування штучного інтелекту в правоохоронній діяльності. Штучний інтелект у правовій практиці: межі. 2024. С. 180. URL: https://www.lvduvs.edu.ua/uk/karta-dokumentiv/category/391-conference-2024-magazine.html?download=5479%3Ashtuchnyi-intelekt-u-pravovii-praktytsi-mezhi-ta-mozhlyvosti-materialy-kruhloho-stolu-15-bereznia-2024-roku&utm_source=chatgpt.com (дата звернення 17.05.2026).

Дар'я КАРАСЬОВА, здобувач ступеня вищої освіти бакалавра ННІ права та психології, Національна академія внутрішніх справ.

Науковий керівник:

Каріна БОРТУН, доцент кафедри ділової української та іноземних мов ННІ права та психології, Національна академія внутрішніх справ, кандидат філологічних наук, доцент

ШТУЧНИЙ ІНТЕЛЕКТ У ЮРИДИЧНІЙ ПРАКТИЦІ: МЕЖІ ВІДПОВІДАЛЬНОСТІ ТА ПЕРСПЕКТИВИ РОЗВИТКУ

Стрімкий розвиток штучного інтелекту (далі – ШІ) визначає новий етап цифрової трансформації суспільства, суттєво впливаючи на механізми функціонування державних інституцій, економіки та системи правосуддя. У правовій сфері ШІ поступово переходить від допоміжного аналітичного інструменту до активного елемента ухвалення рішень, що зумовлює необхідність переосмислення меж його застосування та відповідальності за результати його діяльності. Сучасні інтелектуальні системи, побудовані на основі алгоритмів машинного навчання, здатні здійснювати пошук, обробку та аналіз великих обсягів правової інформації, моделювати судові рішення, оцінювати ризики та навіть формулювати правові висновки [1, с. 231].

Водночас широке впровадження таких технологій породжує низку теоретичних, етичних і практичних проблем. Зокрема, виникає необхідність забезпечення прозорості алгоритмів прийняття рішень, дотримання принципів конфіденційності, уникнення алгоритмічної дискримінації та визначення меж юридичної відповідальності між розробниками, операторами і користувачами систем ШІ. Особливу актуальність проблема набуває в контексті трансформації професійної діяльності юристів, коли традиційні підходи до аналізу правових норм, складання документів чи прогнозування судових результатів змінюються під впливом автоматизованих технологій.

Отже, у сучасних умовах постає завдання всебічного дослідження потенціалу і ризиків використання ШІ в юриспруденції, формування концептуальних і нормативних засад його правового регулювання, а також розроблення етичних стандартів, що гарантуватимуть збереження принципів верховенства права та прав людини в умовах технологічного прогресу. Питання використання штучного інтелекту у юриспруденції активно досліджується як вітчизняними, так і зарубіжними вченими. Серед українських дослідників особливу увагу приділяли проблемам правового регулювання та впровадження ШІ у правничу практику Н. Бааджи, П. Білик, М. Граб, І. Гелецька, С. Наumenко, І. Скляренко, В. Сметюх, М. Томашівський, Р. Чанишев, М. Шовдра, С. Цигульський тощо. Значна частина досліджень присвячена аналізу перспектив автоматизації рутинних юридичних операцій, інтеграції мовних моделей і LegalTech-платформ у процес підготовки документів, оцінки ризиків та прогнозування судових рішень.

З початку 2020-х років питання правового врегулювання штучного інтелекту вийшло на міжнародну та регіональну політичну повістку. У квітні 2021 року Європейський Союз розпочав роботу над проектом нормативного акта, відомого як «Artificial Intelligence Act», метою якого стало створення єдиних правил для застосування і контролю ШІ на території країн-членів. Паралельно на глобальному рівні в липні 2023 року відбулося перше засідання ООН, присвячене проблематиці ШІ; ідею заснування міжнародного органу для узгодження світових стандартів підтримав Генеральний секретар Організації Об'єднаних Націй. Надалі, 21 березня 2024 року, Генеральна Асамблея ООН ухвалила резолюцію, що підкреслює необхідність розвивати безпечні та надійні системи ШІ в інтересах сталого розвитку, а також

закликає дбати про включення країн, що розвиваються, у процес цифрової модернізації шляхом розширення доступу до інфраструктури й технологій [2].

У європейському вимірі регулювання досягло практичної форми: у 2024 році Європейський парламент ухвалив «Закон про штучний інтелект», який згодом отримав схвалення Ради ЄС. Цей правовий акт має інтегральний характер і вперше намагається системно впорядкувати суспільні відносини, що виникають у зв'язку із застосуванням ШІ. Серед ключових елементів закону – категоризація систем штучного інтелекту за рівнем ризику заподіяння шкоди та чітке визначення поняття «система штучного інтелекту»: програмне забезпечення, яке для заданих людиною цілей здатне генерувати результати (контент, прогнози, рекомендації або рішення), що можуть впливати як на фізичне, так і на віртуальне середовище [3-4].

Отже, стрімкий розвиток технологій штучного інтелекту зумовлює суттєву трансформацію сучасної юридичної сфери, відкриваючи нові можливості для автоматизації правових процесів, аналізу нормативної бази та прогнозування судових рішень [5-6]. Водночас інтеграція ШІ у правову практику супроводжується низкою складних викликів, пов'язаних із забезпеченням прозорості алгоритмів, дотриманням прав людини, недопущенням дискримінації та визначенням юридичної відповідальності за результати діяльності інтелектуальних систем. Міжнародний та європейський досвід свідчить про активне формування нормативних механізмів контролю й регулювання ШІ, що актуалізує необхідність адаптації українського законодавства до нових технологічних реалій.

Список використаних джерел

1. Скляренко І.В. Перспективи застосування штучного інтелекту в цивільному судочинстві. *Науковий вісник Ужгородського Національного Університету, Серія ПРАВО*. 2024. №81. Ч. 1. С. 231-237. DOI: <https://doi.org/10.24144/2307-3322.2024.81.1.36>
2. Цигульський С.М. Проблеми правового регулювання поняття штучного інтелекту. Штучний інтелект у правовій практиці: межі та можливості: збірник тез круглого столу (14 березня 2025 року) / упор. О. О. Барабаш. Львів: ЛьвДУВС, 2025. С. 211-213.
3. Гелецька І.О., Шовдра М.М. Визначення поняття «штучного інтелекту» та його місце у системі цивільного законодавства України. *Галицькі студії: юридичні науки*. 2024. № 6. С. 13 –19. DOI: https://doi.org/10.32782/galician_studies/law-2024-6-2
4. Artificial intelligence act. a9-0188/2023. URL: https://www.europarl.europa.eu/doceo/document/A-9-2023-0188-AM-808808_EN.pdf (дата звернення: 12.10.2025).
5. Карасьова Д.І. Штучний інтелект у юридичній лінгвістиці: сучасні виклики для юриста. *Глухівські читання – 2025. Актуальні питання суспільних та гуманітарних наук: збірник матеріалів XV міжнародної науково-практичної інтернет-конференції (10-12 листопада 2025 р.)* / Глухівський НПУ ім. О. Довженка. Глухів, 2025. С. 1235 – 1238.
6. Бортун К.О. Шляхи інтеграції юридичної освіти в європейський освітній простір: упровадження штучного інтелекту. *Наука і правоохорона: наук. журн.* Київ: Нац. акад. внутр. справ, 2025. Вип. 4 (70). С. 71 – 79. DOI: <https://doi.org/10.33270/05257004.7>

Анастасія ЯКОВЕЦЬ, курсантка 2-го курсу н. гр. 207 ННІ поліцейської діяльності, Національна академія внутрішніх справ

ФЕНОМЕН «КОГНІТИВНОЇ ЛІНІ» В УМОВАХ РОЗВИТКУ ШТУЧНОГО ІНТЕЛЕКТУ

Стрімкий розвиток штучного інтелекту став одним із наймасштабніших технологічних прогресів ХХІ століття. Системи штучного інтелекту вже сьогодні використовуються у різних

сферах людського життя: навчанні, професійній діяльності, державному управлінні, правоохоронній сфері, психології, соціології та навіть у повсякденному побуті. Водночас попри позитивні можливості використання таких комп'ютерних технологій виникають і серйозні психологічні та соціальні наслідки, які з кожним днем все частіше стають предметом наукових досліджень, адже технічний процес не стоїть на місці, а постійно перебуває на стадії розвитку.

Однією з проблем постає феномен «когнітивної лінії», який в процесі розвитку штучного інтелекту набуває дискусійного значення. Йдеться, в першу чергу, про поступове зниження здатності людини до самостійного аналізу, критичного осмислення інформації, логічного міркування та прийняття рішень через надмірне делегування інтелектуальних функцій мозку людини системам штучного інтелекту. Когнітивне розвантаження, яке досліджувалося Ріско та Гілбертом, передбачає використання зовнішніх інструментів для зменшення когнітивного навантаження на робочу пам'ять людини [1, с. 676]. Хоча це може звільнити когнітивні ресурси, це також може призвести до зниження когнітивної залученості та розвитку навичок.

Сьогодні щонайменше кожний третій дедалі частіше передає штучному інтелекту не лише технічні функції, а й процеси, які раніше вимагали активної розумової людської діяльності, включаючи як пошук інформації, аналіз джерел, формування висновків, створення текстів, прийняття рішень, так навіть і емоційну підтримку.

У науковій літературі це явище описується через поняття *cognitive offloading* — «когнітивне розвантаження», тобто перенесення частини мисленнєвих функцій людини на зовнішні цифрові системи, включаючи і штучний інтелект [3, с. 4]. Небезпека полягає у тому, що мозок людини поступово звикає до спрощеного механізму отримання готових відповідей. Якщо раніше людина була змушена аналізувати великі обсяги інформації, порівнювати різні позиції та критично оцінювати джерела, то сьогодні достатньо сформулювати короткий запит до AI-системи. З кожним днем людина все дедалі рідше користується своєю когнітивною функцією, а мозок втратив здатність критично мислити, що буквально поставило його на паузу.

На мою думку, найбільше це проявляється серед молоді: школярів та студентів. Використання штучного інтелекту значно спрощує написання контрольних робіт, рефератів, есе, курсових робіт та навіть наукових текстів. Проте разом із цим знижується рівень самостійного осмислення матеріалу. Людина за секунду отримує готову структуровану відповідь і часто сприймає її як власне знання/рішення, хоча процес глибокого розуміння не відбувається взагалі. Здатність сформулювати питання штучному інтелекту взагалі не гарантує активізації нейронних зв'язків та когнітивних процесів [4, с. 19].

У психології це називають «ілюзією компетентності» або ефектом Даннінга-Крюгера — станом, коли користувач переоцінює власний рівень знань лише через доступ до швидкої інформації, не розуміючи помилок [2, с. 2226]

Водночас вплив штучного інтелекту виходить далеко за межі лише когнітивної сфери. технології штучного інтелекту починають змінювати не лише сприймання інформації, а й психіку людини, її соціальну поведінку та навіть способи комунікації. Через постійну взаємодію з цифровими системами людина поступово адаптується до спрощеного формату сприйняття інформації: короткі відповіді, швидкі висновки, мінімальна концентрація уваги. Це негативно впливає на здатність до тривалого аналізу, читання та розуміння великих текстів та осмислення складних проблем.

До того ж, надмірне використання штучного інтелекту може сприяти розвитку тривожності та емоційного виснаження людини [5, с. 6]. Адже, постійний потік інформації і повідомлень створює для психіки стан стимуляції, через що мозок перебуває у стані хронічного когнітивного навантаження. Це, в свою чергу, може призводити до проблем зі сном, підвищеної дратівливості, втрати концентрації та емоційної нестабільності.

Проте головна небезпека сучасного світу полягає саме у зміні моделі мислення людини. Інформаційна культура дедалі більше орієнтується не на пошук істини, а на швидкість отримання відповіді. У таких умовах людина поступово звикає до інтелектуального

споживання без глибокого осмислення. Це може мати довгострокові наслідки, включаючи зниження рівня освіти, ослаблення логічного мислення, падіння інформаційної стійкості суспільства, підвищення сприймання дезінформації.

Таким чином, феномен «когнітивної лінії» є не просто психологічною особливістю сучасної людини, а масштабною соціальною трансформацією. Штучний інтелект поступово перебирає на себе аналітичні функції, що створює ризик втрати людиною навички самостійного мислення. Головний виклик полягає не у самому існуванні штучного інтелекту, а у тому, чи залишиться людина активним суб'єктом мислення, а не пасивним споживачем готових відповідей.

Список використаних джерел

1. Risko, E.F.; Gilbert, S.J. Cognitive Offloading. *Trends in Cognitive Science*. 2016. Vol 20. № 9. P. 676–688.
2. Чуйко Г. В., Чаплак Я. В. Ефект Даннінга-Крюгера: чому невігласи завищують свою самооцінку. *Журнал «Наукові інновації та передові технології»*. 2025. № 2. С. 2221-2232.
3. Gerlich M. AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking. *Societies*. 2025. Vol 15. № 6. P. 1-28.
4. Лукашова Т., Друшляк М. Штучний інтелект як засіб розвитку критичного мислення майбутніх учителів математики. *Фізико-математична освіта*. 2023. Том 38. № 5. С. 18-25.
5. Zhai, C., Wibowo, S., Li, L.D. The effects of over-reliance on AI dialogue systems on students' cognitive abilities: a systematic review. *Smart Learning Environments*. 2024. Vol. 11. P. 1-37.

Юрій ДОЛГІХ, слухач інституту стратегічних комунікацій, Національний університет оборони України.

Науковий керівник:

Сергій НЕХАЄНКО, доцент кафедри військової психології, Національний університет оборони України, кандидат педагогічних наук

ЗАГРОЗИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ ПСИХОЛОГІЧНОЇ ПІДТРИМКИ ВІЙСЬКОВОСЛУЖБОВЦІВ В УМОВАХ ВІДСІЧІ ЗБРОЙНІЙ АГРЕСІЇ РОСІЇ ПРОТИ УКРАЇНИ

Розв'язана російською федерацією проти України 24 лютого 2022 року широкомасштабна збройна агресія висунула підвищені вимоги до психологічної підтримки персоналу. Крім фізичного знищення особового складу вогневими засобами існують також потенційні загрози, пов'язані з використанням штучного інтелекту (далі – ШІ) для психологічної підтримки, які можуть бути використані з метою підвищення мотивації до виконання завдань військовослужбовцями Збройних Сил України (далі – ЗСУ), Національної гвардії України (далі – НГУ) та інших складових сил оборони (далі – ІССО) під час відсічі збройній агресії росії проти України. Обговорюються технічні обмеження ШІ як засобу психологічної підтримки, проблеми кібербезпеки, етичні та юридичні аспекти, а також можливі негативні впливи на ефективність виконання службово-бойових завдань.

Розглянемо можливість використання онлайн-сервісів із ШІ для мобільних телефонів з операційними системами Android/iOS як інструмента психологічної підтримки на прикладі інтегрованого в месенджер «Viber» ШІ-коуча (далі – Pi AI).

Коуч (від англ. coach – тренер) – це спеціаліст, який допомагає людям досягати особистих чи професійних цілей, розкривати потенціал та знаходити власні рішення, ставлячи запитання та використовуючи спеціальні техніки, фокусуючись на фактичному стані клієнта. У нашому випадку Pi AI – це віртуальний помічник, який може бути поруч із військовослужбовцем незалежно від пори року, часу доби та наявності фахівців психологічної підтримки.

Pi AI був інтегрований у месенджер «Viber» для забезпечення емоційної підтримки та стимулювання користувачів платформи [14]. Він створений компанією Inflection AI та представлений як онлайн-помічник, призначений для надання рекомендацій та мотиваційних настанов у повсякденних сценаріях.

Однак, незважаючи на його широке визнання, впровадження Pi AI для мотивації військовослужбовців у контексті виконання завдань із відсічі збройній агресії пов'язане з істотними ризиками. Pi AI не функціонує як професійний психотерапевтичний інструмент і не адаптований до екстремальних умов, подібних до бойових дій. Його алгоритми моделюють емпатію, але переважно генерують поверхові та стереотипні відповіді, які ігнорують складні емоційні динаміки [6].

У підрозділах НГУ, де особовий склад стикається з посттравматичним стресовим розладом, емоційним виснаженням чи кризами мотивації, Pi AI може пропонувати невідповідні рекомендації, що потенційно можуть погіршити ситуацію [4]. Емпіричні дослідження демонструють, що чат-боти на базі ШІ, подібні до Pi AI, неефективні в управлінні тяжкими психічними станами з імовірністю шкідливих порад у 20% випадків. Наприклад, вони можуть акцентувати негативні елементи або не ідентифікувати індикатори суїцидальних тенденцій, завдаючи шкоду [6]. Більше того, Pi AI не враховує специфіку військового контексту, включаючи культурні, вікові чи гендерні фактори, що впливають на мотиваційні процеси [5]. Під час збройних конфліктів мотивація до виконання завдань часто ґрунтується на колективних ідеалах, таких як патріотизм чи групової єдності, а не на індивідуалізованих порадах.

Застосування Pi AI може формувати ізольовані «емоційні сфери», де алгоритм адаптується до користувача, але не сприяє розвитку справжньої резиліентності, необхідної в складних умовах. Якби алгоритми не були прописані в програмі, ШІ не може замінити кваліфікованих фахівців [5]. Месенджер «Viber», як середовище для функціонування Pi AI, демонструє значну вразливість в умовах протистояння у кіберпросторі. Російські кіберзлочинці, зокрема угруповання UAC-0184 (Hive0156), систематично використовують месенджер «Viber» для фішингових операцій проти Збройних Сил України та державних інституцій [7; 13]. Вони розповсюджують шкідливе програмне забезпечення через архіви ZIP, маскуючи їх під офіційні файли, що веде до інсталяції на телефони Remcos RAT (Remote Control and Surveillance) – програм для шпигунства та дистанційного контролю за пристроєм [8].

У 2026 році зареєстровано нові атаки, де месенджер «Viber» є середовищем для спрямованих операцій проти сил безпеки і оборони України [1; 8]. Проблеми конфіденційності даних у Pi AI також критичні: діалоги можуть використовуватися для вдосконалення моделей або передаватися зовнішнім сторонам, створюючи ризик витоку конфіденційної інформації про оперативні завдання чи психологічний стан військовослужбовців [10]. У контексті конфлікту, коли росія застосовує ШІ для поширення дезінформації та кіберінцидентів, використання месенджера «Viber» для мотивації може відкрити шляхи для маніпуляцій, наприклад, через діпфейки (від англ. deepfake; об'єднання слів deep learning («глибоке навчання») та fake («підробка»)) чи імітаційні чат-боти [2; 9].

Росія заборонила месенджер «Viber» на своїй території, але активно використовує його для кіберінцидентів проти України, що підкреслює ненадійність цієї платформи [7]. Використання ШІ для мотивації військовослужбовців суперечить і етичним принципам. Збройний конфлікт в Україні ілюструє небезпеку використання ШІ в силах безпеки та оборони України: системні упередження можуть спричинити й ескалацію насильства [3; 14]. ШІ,

позбавлений людського нагляду, може непрямо стимулювати деструктивні дії, нехтуючи нормами пропорційності в міжнародному гуманітарному праві [9]. З етичної перспективи мотивація за допомогою Рі AI сприяє дегуманізації конфлікту, перетворюючи військових на «оптимізованих» суб'єктів без урахування психологічних загроз, таких як залежність чи ерозія моральної відповідальності [5]. Дуже важливо, що неправильно сформовані запити до Рі AI можуть знизити авторитет командування [12]. Крім того, брак нормативного регулювання ІІІ у військовій сфері робить таке застосування потенційно протиправним [14]. Замість використання Рі AI у ЗСУ та ІССО доцільно впроваджувати захищені платформи під людським контролем, уникати комерційних месенджерів та проводити експертизу ІІІ перед його інтеграцією [11].

Таким чином, застосування Рі AI у месенджері «Viber» для мотивації військовослужбовців у значній частині переважає ризиками над перевагами: від психологічної шкоди та кіберзагроз до етичних порушень. Виклики сьогодення вимагають надійних та спеціалізованих інструментів, а не використання універсальних чат-ботів. Недооцінка цих факторів може знизити якість виконання службово-бойових завдань та призвести до необґрунтованих втрат на полі бою. З огляду на викладене необхідно якнайшвидше сформувати нормативно-правову базу для використання ІІІ для ЗСУ, НГУ та ІССО, враховуючи наявні обмеження та ризики.

Список використаних джерел

1. Censor.net. Чи була причиною блекауту кібератака? [електронний ресурс]. 2026. URL: <https://censor.net/ua/blogs/3598315/chy-bula-prychynoyu-blekautu-kiberataka> (дата звернення: 07.05.2026).
2. Detector Media. Завдяки алгоритмам штучного інтелекту російські боти тепер пишуть українською [електронний ресурс]. 2024. URL: <https://detector.media/infospace/article/228349/2024-06-17-zavdyaky-algorytmam-shtuchnogo-intelektu-rosiyski-boty-teper-pyshut-ukrainskoyu> (дата звернення: 07.05.2026).
3. Suspilne. Штучний інтелект у творчості та на війні: можливості, ризики та авторське право [електронний ресурс]. 2023. URL: <https://suspilne.media/culture/426576-stucnij-intelekt-u-tvorcosti-ta-na-vijni-pro-mozlivosti-riziki-ta-avtorske-pravo> (дата звернення: 07.05.2026).
4. Ukrinform. Штучний інтелект і нова війна обману: Уроки з України [електронний ресурс]. 2025. URL: <https://www.ukrinform.ua/rubric-ato/4051095-stucnij-intelekt-i-nova-vijna-obmanu-uroki-z-ukraini.html> (дата звернення: 07.05.2026).
5. Bajbouj M., Kemna S. Ethical Implications of AI in Warfare // Journal of Military Ethics. 2023. № 2. С. 45–67.
6. BMC Psychology. Effectiveness of AI Chatbots in Mental Health Crises. 2024. № 3. С. 112–130.
7. CERT-UA. Phishing Attacks via Viber Targeting Ukrainian Military. Ukrainian Cyber Emergency Response Team [електронний ресурс]. 2024. URL: <https://cert.gov.ua/> (дата звернення: 07.05.2026).
8. Cybersecurity and Infrastructure Security Agency (CISA). Remcos RAT Campaigns in Ukraine [електронний ресурс]. 2026. URL: <https://www.cisa.gov/> (дата звернення: 07.05.2026).
9. Human Rights Watch. AI and International Humanitarian Law in Ukraine Conflict. 2023. 89 с.
10. Inflection AI. Pi AI Documentation on Privacy and Data Usage. 2024. 56 с.
11. NATO. AI in Psychological Resilience Programs for Armed Forces [електронний ресурс]. 2024. URL: <https://www.nato.int/> (дата звернення: 07.05.2026).
12. Ukrainian Ministry of Defense. Psychological Support Systems for NGU. 2025. 210 с.
13. Ukrainska Pravda. Number of cyberattacks on Ukraine increased by 70% in past year [електронний ресурс]. 2025. URL: <https://www.pravda.com.ua/eng/news/2025/01/09/7492671> (дата звернення: 07.05.2026).
14. United Nations. Political Declaration on Responsible Military Use of Artificial Intelligence and Autonomy [електронний ресурс]. 2023. URL: <https://www.un.org/en/documents/political-declaration-responsible-military-use-ai> (дата звернення: 07.05.2026)

Дарія КОВЕНЬКО, студентка н. гр. 112_СПД
ННІ права та психології, Національна академія
внутрішніх справ.

Науковий керівник:

Микола ШВЕЦЬ, викладач кафедри
інформаційних технологій ННІ права та
психології, Національна академія внутрішніх справ

ТЕХНОЛОГІЇ ШТУЧНОГО ІНТЕЛЕКТУ В УКРАЇНСЬКИХ ІНФОРМАЦІЙНИХ СИСТЕМАХ

У сучасному світі технології штучного інтелекту є одним із головних напрямів розвитку інформаційних систем. Використання ШІ дозволяє автоматизувати обробку даних, прогнозувати події, оптимізувати процеси управління та підвищувати ефективність роботи організацій. Для України розвиток штучного інтелекту є важливим елементом цифрової трансформації держави та інтеграції у світовий інформаційний простір [1].

Українські інформаційні системи поступово впроваджують інтелектуальні алгоритми машинного навчання, комп'ютерного аналізу даних і автоматизованої підтримки прийняття рішень. Особливо активно ШІ використовується у державному та оборонному секторах, медицині, фінансових сервісах, освіті та сфері кібербезпеки.

Успішна інтеграція ШІ в інформаційні системи неможлива без чіткого правового поля. Україна активно рухається в напрямку євроінтеграції, орієнтуючись на європейський Акт про штучний інтелект (*EU AI Act*) [4]. Міністерство цифрової трансформації України розробило «Концепцію розвитку штучного інтелекту в Україні» та дорожню карту регулювання ШІ. Державна стратегія базується на принципі «bottom-up» (знизу вгору), що передбачає спочатку надання бізнесу та розробникам м'яких інструментів (рекомендацій, білих книг), а вже потім — ухвалення жорсткого законодавства. Це стимулює інновації в ІС, зберігаючи баланс між безпекою та технологічним прогресом [3].

Одним із найуспішніших прикладів цифрової трансформації України є державна система цифрових послуг «Дія». Платформа забезпечує доступ громадян до електронних документів, адміністративних послуг та цифрової взаємодії з державою. У 2024 р. система Дії розширилася завдяки **Дія.AI** — першому державному штучному інтелекту, який допомагає користувачам знаходити й замовляти послуги просто у чаті на порталі diia.gov.ua [2]

Іншим прикладом є система аналітики державних закупівель **Prozorro**, де алгоритми аналізу даних використовуються для виявлення ризикових тендерів та можливих корупційних схем. Prozorro стала одним із найуспішніших ІТ-проектів українського уряду й отримала визнання міжнародних донорів і антикорупційних організацій. Система стала моделлю для інших країн: на її основі створено MTender у Молдові та аналогічні рішення у Казахстані, Таджикистані й Узбекистані. У період воєнного стану Prozorro продовжує роботу, зокрема у форматі **Prozorro Market**, який забезпечує швидкі закупівлі для цивільних і військових потреб.

Для моніторингу відкритих державних даних застосовується система **Opendatabot**, яка використовує автоматизований аналіз великих масивів інформації та інтелектуальні алгоритми пошуку. Opendatabot об'єднує інформацію з численних державних баз, включно з ЄДР, судовими реєстрами, реєстром боржників та податковими даними. Користувачі можуть відстежувати зміни у компаніях, судові справи, податкові ризики чи санкції, отримуючи сповіщення через месенджери або електронну пошту.

Використання ШІ у державному секторі дозволяє: автоматизувати документообіг; здійснювати аналітику великих масивів даних; прогнозувати ризики; підвищувати ефективність державного управління; покращувати кіберзахист інформаційних систем.

В умовах повномасштабної агресії **інформаційні системи військового призначення** стали головним споживачем ШІ-технологій [5]:

- **Геоінформаційні системи (ГІС):** ШІ аналізує супутникові знімки та дані з БПЛА для автоматичного виявлення техніки та змін у рельєфі.
- **Системи розпізнавання звуку:** Платформи акустичного моніторингу використовують нейромережі для ідентифікації траєкторій польоту ракет та дронів-камікадзе.
- **Кіберзахист:** ІС Центрів реагування на кіберінциденти автоматично блокують DDoS-атаки за допомогою аномалійного аналізу трафіку.

Український **банківський сектор** активно впроваджує інтелектуальні технології для автоматизації фінансових операцій та підвищення рівня безпеки. Одним із прикладів є цифровий банк **monobank**, де використовуються алгоритми машинного навчання для аналізу фінансових операцій, виявлення шахрайства, автоматизованої підтримки клієнтів, персоналізованих фінансових рекомендацій [6].

Основними перевагами використання штучного інтелекту в українських інформаційних системах є:

- швидка обробка великих обсягів інформації;
- автоматизація рутинних процесів;
- зменшення людського фактору;
- підвищення точності прогнозування;
- покращення якості державних та комерційних послуг.

Завдяки інтелектуальним технологіям організації можуть ефективніше використовувати ресурси та швидше адаптуватися до сучасних викликів.

Попри значні переваги, впровадження ШІ в Україні супроводжується певними труднощами: недостатнім рівнем фінансування; потребою у кваліфікованих спеціалістах; ризиками витоку персональних даних; кіберзагрозами; необхідністю вдосконалення нормативно-правової бази.

Для оптимізації процесу впровадження ШІ в інформаційні системи України доцільно вжити такі кроки, як **створення локальних моделей** (розробка україномовних великих мовних моделей (LLM) для мінімізації помилок перекладу та покращення роботи державних чат-ботів) та **гібридних систем** (використання архітектури, де ШІ виступає асистентом (Decision Support System), а фінальне рішення залишається за людиною-оператором).

Використання штучного інтелекту в інформаційних системах України є не просто технологічним трендом, а необхідною умовою для виживання та розвитку держави. Україна вже має унікальні кейси застосування ШІ у сферах е-урядування та оборони. Подальший успіх залежить від гармонізації законодавства з нормами ЄС, підвищення рівня кібербезпеки інформаційних систем та фінансування вітчизняних науково-практичних розробок у галузі машинного навчання.

Список використаних джерел

1. Про схвалення Концепції розвитку штучного інтелекту в Україні: Розпорядження Кабінету Міністрів України від 2 грудня 2020 р. № 1556-р. *Урядовий кур'єр*. 2020.
2. Федоров М. А. Цифрова трансформація та роль технологій штучного інтелекту в державній системі «Дія». *Вчені записки ТНУ імені В.І. Вернадського. Серія: Технічні науки*. 2023. Т. 34 (73), № 2. С. 45–51.
3. Міністерство цифрової трансформації України. Дорожня карта регулювання штучного інтелекту в Україні. Офіційний сайт. URL: thedigital.gov.ua (дата звернення: 15.05.2026).
4. European Parliament. Artificial Intelligence Act. Briefing: EU Legislation in Progress. 2024. URL: european.eu (дата звернення: 12.05.2026).
5. Рижков О. В., Шевченко А. І. Застосування нейромережевих технологій в інформаційних системах моніторингу критичної інфраструктури. *Кібербезпека та комп'ютерна інженерія*. 2024. № 1 (11). С. 12–19.
6. Снігур В. М. Проблеми кібербезпеки при інтеграції AI-сервісів в банківські інформаційні системи України. *Сучасні інформаційні технології*. 2025. Т. 12, № 4. С. 89–95.

Володимир ШИПП, курсант ННІ поліцейської діяльності, Національна академія внутрішніх справ.

Науковий керівник:

Ірина БОТНАРЕНКО, старший науковий співробітник науково-дослідної лабораторії з проблем протидії злочинності ННІ поліцейської діяльності, Національна академія внутрішніх справ, кандидат юридичних наук

ШТУЧНИЙ ІНТЕЛЕКТ У СУДОЧИНСТВІ ТА ДІЯЛЬНОСТІ ПОЛІЦІЇ: ІННОВАЦІЙНІ МОЖЛИВОСТІ ТА РИЗИКИ ДЛЯ ПРАВ ЛЮДИНИ

Штучний інтелект (далі – ШІ) сьогодні розглядається як одна з ключових технологій, що здатні радикально трансформувати публічне управління, судочинство та правоохоронну діяльність. Перспективи автоматизації відповідних процесів включають підвищення ефективності ухвалення рішень, оптимізацію ресурсів, оперативність аналізу даних та створення передових інструментів прогнозування ризиків і загроз. Разом із тим широке впровадження ШІ супроводжується серйозними етичними викликами, зокрема можливостями дискримінації, порушенням приватності, непрозорістю алгоритмічних рішень і ризиком зловживань у високочутливих сферах суспільного життя.

Європейський Союз, який розробив один із найамбіційніших у світі підходів до регулювання ШІ через так званий AI Act, зосереджує увагу саме на поєднанні інновацій і захисту основоположних прав людини [1]. В рамках цього регламенту вже заборонено найризикованіші застосування штучного інтелекту, як-то соціальний скоринг, непрозорі системи прогнозу поліції та автоматичне розпізнавання обличь без чітких гарантій прав та свобод, а для систем з високим ризиком накладаються посилені вимоги щодо прозорості, безпеки та контролю, вводяться суворіші стандарти [2]. Таким чином, перспективи автоматизації у публічному секторі розглядаються ЄС як потенційно корисні за умови суворого дотримання етичних та правових стандартів, що мають бути імplementовані як на рівні розробників, так і на рівні органів влади.

ШІ дедалі активніше інтегрується в діяльність правоохоронних органів Європейського Союзу, зокрема в системи, що підтримують аналіз даних, розпізнавання обличь, прогнозу поліцію та автоматизоване виявлення загроз. За даними правових досліджень, застосування ШІ у кримінальному процесі ЄС включає використання різноманітних технологій для оптимізації розслідувань, підвищення оперативності контролю за злочинністю та аналізу великих обсягів доказів, що створює нові можливості для боротьби зі складними злочинами [3].

Водночас широке застосування таких технологій викликає серйозні етичні й правові питання, зокрема щодо прозорості алгоритмів, гарантій непорушення фундаментальних прав людини, контролю за безпекою даних і недопущення дискримінації у процесі застосування автоматизованих рішень правоохоронцями. Саме тому Європейський Союз активно працює над розробленням нормативних актів та регуляцій, які б забезпечили баланс між ефективністю боротьби із злочинністю та дотриманням етичних стандартів при використанні ШІ поліцією, судовими та іншими органами кримінальної юстиції.

Використання штучного інтелекту в правоохоронній діяльності супроводжується низкою суттєвих етичних ризиків, серед яких особливу увагу привертають алгоритмічна упередженість, автоматизований профайлінг та потенційна дискримінація окремих соціальних груп. Системи прогнозування злочинності, розпізнавання обличь та аналізу поведінкових моделей можуть відтворювати або навіть посилювати вже наявні соціальні стереотипи, якщо алгоритми навчаються на статистичних даних, що містять історичні викривлення. У такому випадку виникає ризик непропорційного контролю щодо окремих категорій населення, що суперечить принципам рівності та недискримінації.

Європейський Союз у 2024 році ухвалив Регламент про штучний інтелект (AI Act), який визначає використання систем біометричної ідентифікації в режимі реального часу правоохоронними органами як високоризикове та передбачає суворі обмеження щодо їх застосування, зокрема обов'язковість правових підстав, пропорційності та судового контролю [1]. Такий підхід демонструє прагнення забезпечити баланс між потребами публічної безпеки та захистом фундаментальних прав людини, закріплених у праві Європейського Союзу.

Отже, перспективи впровадження штучного інтелекту в діяльність поліції безпосередньо пов'язані з необхідністю формування чітких процедур аудиту алгоритмів, запровадження механізмів прозорості прийняття рішень та гарантування права особи на оскарження автоматизованих рішень. Без належних правових і етичних запобіжників автоматизація правоохоронної діяльності може призвести до системних порушень прав людини та підриву довіри до державних інституцій.

В Україні інтеграція технологій ШІ у правозастосування, зокрема в судочинстві та правоохоронній діяльності, розглядається як перспективний напрям модернізації правової системи, що має потенціал суттєво підвищити її ефективність. Використання ШІ під час аналізу нормативних актів, автоматизації рутинних процесів, прогнозування результатів розгляду справ та підтримки процедур ухвалення рішень може знизити навантаження на суддів і правоохоронців, пришвидшити доступ до правосуддя та зменшити рівень суб'єктивних помилок [4].

Водночас українські дослідники звертають увагу на важливість чітких правових і етичних обмежень при впровадженні ШІ у правозастосовну практику. Це включає забезпечення прозорості алгоритмів, гарантування контролю людини над критичними рішеннями, захисту персональних даних, недопущення алгоритмічної дискримінації та створення механізмів відповідальності за помилки автоматизованих систем. Без врахування цих аспектів використання ШІ може призвести до порушення принципів справедливості й рівності перед законом, що підриває довіру до інститутів правосуддя та правоохоронних органів. Таким чином, українська перспектива застосування ШІ в правозастосуванні вимагає поєднання інноваційних можливостей із суворими етичними й правовими запобіжниками, що гарантують дотримання фундаментальних прав людини.

Список використаних джерел

1. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*. 2024. L 1689. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (дата звернення: 25.02.2026).
2. Europe sets benchmark for rest of the world with landmark AI laws/ Reuters. 21.05.2024. URL: <https://www.reuters.com/world/europe/eu-countries-back-landmark-artificial-intelligence-rules-2024-05-21/> (дата звернення: 25.02.2026).
3. Chernychenko I. Artificial intelligence in criminal procedure: current legal regulations in the EU. *Visegrad Journal on Human Rights*. 2025. DOI: <https://doi.org/10.61345/1339-7915.2025.3.1> (дата звернення: 25.02.2026).
4. Gutsalyuk M., Klymenko-Panchenko O. Implementation of artificial intelligence in judicial activities: international and Ukrainian experience. *Air and Space Law Journal*. 2025. Vol. 2. Pp. 167–175. URL: <https://doi.org/10.18372/2307-9061.75.20230> (дата звернення: 25.02.2026).

Вікторія ЗАХАРОВА, здобувач першого (бакалаврського) рівня вищої освіти, Львівський державний університет внутрішніх справ.

Науковий керівник:

Роман ДЕМКІВ, доцент кафедри загальноправових дисциплін, Львівський державний університет внутрішніх справ, кандидат юридичних наук, доцент

ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ ПІД ЧАС ЗДІЙСНЕННЯ ПРАВОСУДДЯ: ПЕРСПЕКТИВИ ТА РИЗИКИ

Використання штучного інтелекту (далі – ШІ) під час здійснення правосуддя є одним із найактуальніших напрямів розвитку сучасної правової системи. Стрімка цифровізація суспільства, зростання обсягів інформації та потреба в оперативному й об’єктивному розгляді справ зумовлюють інтерес до впровадження інтелектуальних технологій у діяльність судів. ШІ розглядається як комплекс програмних рішень, здатних аналізувати великі масиви даних, виявляти закономірності, прогнозувати результати та підтримувати ухвалення рішень.

У сфері правосуддя ці можливості відкривають нові перспективи для підвищення ефективності судочинства, одночасно породжуючи низку правових, етичних та соціальних ризиків.

Однією з головних перспектив застосування ШІ у судовій системі є автоматизація рутинних процесів. Судді та працівники апарату судів щодня опрацьовують значну кількість документів: позовних заяв, клопотань, судових рішень, доказових матеріалів. Алгоритми штучного інтелекту можуть сприяти суду та учасникам справи у їх процесуальній діяльності, головним чином, складаючи проекти процесуальних документів та пропонуючи редакційні правки у них, висловлюючи рекомендації щодо застосування законодавства з урахуванням практики Європейського суду з прав людини, Верховного Суду та інших джерел права, здійснюючи логічну та арифметичну перевірку обґрунтованості вимог та заперечень учасників справи [1, с. 63].

Важливою перевагою ШІ є можливість швидкого аналізу судової практики. Системи машинного навчання здатні опрацьовувати тисячі рішень і виявляти тенденції, що сприяє уніфікації судової практики. Закордонний досвід (Китай з концепцією «розумного суду», системи Prometea в Аргентині, VICTOR у Бразилії, OLGA у Німеччині, COMPAS у США) демонструє ефективність таких інструментів для попереднього аналізу справ, оцінки ризиків та прогнозування результатів, однак свідчить про відсутність єдиної моделі впровадження та використання ШІ у правосудді. Юрисдикції використовують чотири основні підходи: від цільового чи загального регулювання та створення власних AI-систем до повної або тимчасової заборони їх застосування у правосудді [2, с. 379].

Штучний інтелект може також сприяти розширенню доступу до правосуддя. Чат-боти та електронні консультанти здатні надавати громадянам базову правову інформацію, допомагати у підготовці документів, пояснювати порядок звернення до суду та відстеження стану справи. Для осіб, які не мають можливості скористатися послугами адвоката, такі сервіси можуть стати важливим інструментом реалізації права на судовий захист. Крім того, автоматизовані системи можуть забезпечувати переклад документів, адаптацію інформації для осіб з інвалідністю та інші сервіси, що роблять судову систему більш інклюзивною.

Отже, використання штучного інтелекту під час здійснення правосуддя відкриває значні можливості для підвищення ефективності, оперативності та доступності судочинства. Разом із тим існують суттєві ризики, пов’язані з упередженістю алгоритмів, непрозорістю їхньої роботи, захистом персональних даних, невизначеністю відповідальності та загрозою втрати людського виміру правосуддя. Використання ШІ також несе ризик надмірної формалізації.

Судове рішення є результатом не лише сукупності доказів, моральних аспектів, принципів справедливості. Алгоритм не здатний повною мірою врахувати контекст конкретної справи. Тому передача ключових функцій штучному інтелекту суперечить самій природі правосуддя та могла б призвести до знеособлення процесу [3, с. 236].

Іншим серйозним ризиком є недостатня прозорість алгоритмів. Багато моделей штучного інтелекту функціонують як «чорні скриньки», тобто користувач не може повною мірою зрозуміти логіку отриманого результату. Для правосуддя, яке ґрунтується на принципах мотивованості, гласності та можливості оскарження рішень, така непрозорість є проблематичною. Якщо суддя використовує рекомендації алгоритму, але не може пояснити, як саме вони були сформовані, це може підірвати довіру до судової влади та поставити під сумнів справедливість судового розгляду.

Суттєве значення має й питання відповідальності за помилки штучного інтелекту. У разі неправильного прогнозу або некоректної рекомендації виникає питання, хто повинен нести юридичну відповідальність: розробник програмного забезпечення, судова адміністрація, суддя чи інший суб'єкт. Чинне законодавство більшості держав поки не містить чітких механізмів розподілу відповідальності за наслідки використання ШІ у правосудді, що створює правову невизначеність.

Не менш важливим є ризик порушення права на приватність і захист персональних даних. Для ефективної роботи системи штучного інтелекту потребують доступу до великого обсягу інформації, зокрема персональних даних учасників процесу, відомостей про судимості, фінансовий стан, сімейні обставини тощо. Недостатній рівень кібербезпеки або неналежне використання таких даних може призвести до витоку конфіденційної інформації та порушення прав людини.

Міжнародні організації наголошують на необхідності принципу «людина під контролем». Зокрема, Європейська комісія з ефективності правосуддя (CEPEJ) та Акт ЄС про штучний інтелект підкреслюють, що ШІ може використовуватися лише як допоміжний інструмент, а остаточне рішення повинен ухвалювати суддя-людина.

Таким чином, підсумовуючи, можна зазначити, що для України питання використання штучного інтелекту є особливо актуальним у контексті реформування судової системи та розвитку цифрових сервісів. Перспективними напрямками є автоматизація рутинних процесів, розподіл справ, контроль процесуальних строків та допомога суддям у підготовці рішень.

Список використаних джерел

1. Маймур Ф. Ф. Використання штучного інтелекту при здійсненні правосуддя: проблеми та перспективи. *Право та державне управління*. 2025. № 1. С. 61–67.
2. Бессонов О. В. Зарубіжний досвід використання штучного інтелекту в правосудді. *Науковий вісник Ужгородського національного університету. Серія: Право*. 2025. Вип. 88. С. 375–380.
3. Складенко І. В. Перспективи застосування штучного інтелекту в цивільному судочинстві. *Науковий вісник Ужгородського національного університету. Серія: Право*. 2024. Вип. 81. С. 231–236.

Владислава ІВАНОВА, студентка 1-го курсу бакалаврату юридичного факультету, Хмельницький університет управління та права імені Леоніда Юзькова.

Науковий керівник:

Катерина ВЛАДОВСЬКА, доцент кафедри конституційного, адміністративного та фінансового права, Хмельницький університет управління та права імені Леоніда Юзькова, доктор філософії в галузі права

КОНСТИТУЦІЙНО-ПРАВОВІ АСПЕКТИ ЦИФРОВІЗАЦІЇ ПУБЛІЧНОЇ ВЛАДИ: МЕЖІ АВТОМАТИЗАЦІЇ УПРАВЛІНСЬКИХ ПРОЦЕСІВ

Цифровізація публічної влади в Україні давно вийшла за межі технологічного експерименту і перетворилася на повноцінний напрям державної політики. Сьогодні мільйони громадян щодня взаємодіють з державою через портал «Дія», електронні реєстри та автоматизовані сервіси. Разом із тим стрімке впровадження технологій штучного інтелекту (далі в тексті - ШІ) та алгоритмізованих систем у сферу прийняття управлінських рішень ставить питання – де знаходиться конституційна межа між допустимою автоматизацією та неприпустимим делегуванням владних повноважень алгоритму? Вирішення цього питання визначає актуальність дослідження та його мету.

Нормативну основу державної стратегії у сфері ШІ закладає Концепція розвитку штучного інтелекту в Україні, схвалена Кабінетом Міністрів у грудні 2020 року [1]. Документ передбачає широкий горизонт застосування ШІ в публічному секторі - від автоматизованої обробки адміністративних запитів до підтримки правоохоронної та судової діяльності. Водночас Концепція оминає питання конституційних меж такого застосування, зосереджуючись переважно на технологічних та економічних перевагах цифрової трансформації. Це створює значну прогалину між декларованими цілями та реальними правовими гарантіями для громадян.

Вихідним конституційним принципом для розуміння меж автоматизації є стаття 19 Конституції України: «Органи державної влади та органи місцевого самоврядування, їх посадові особи зобов'язані діяти лише на підставі, в межах повноважень та у спосіб, що передбачені Конституцією та законами України» [2]. Ця норма встановлює так звану нормативно закріплену модель повноважень публічної влади - держава може робити лише те, на що є пряма законодавча підстава. Перенесення цього принципу у цифровий вимір означає таке, що жодна автоматизована система не може здійснювати управлінські функції без чіткого законодавчого регулювання, яке б визначало її повноваження, допустимі методи роботи та правові наслідки її рішень. Поки що спеціального закону про застосування ШІ у публічному управлінні в Україні немає, отже, значна частина нинішньої практики автоматизованих рішень де-факто функціонує в умовах правової невизначеності.

Однією з ключових конституційних проблем у контексті цифровізації є захист приватності. Стаття 32 Конституції гарантує кожному недоторканість особистого і сімейного життя та прямо забороняє збір, зберігання й поширення конфіденційної інформації про особу без її згоди, крім випадків, визначених законом [2]. Автоматизовані державні системи акумулюють колосальні масиви персональних даних: відомості про доходи, стан здоров'я, місце проживання, пересування, соціальні зв'язки тощо. Проблема полягає не лише в можливості витоку даних, хоча це і є реальним ризиком. Значно суттєвішою є загроза алгоритмічного профілювання особи, яке полягає в автоматичному формуванні детального «портрета» громадянина на основі даних з різних реєстрів, що надалі є підставою для прийняття рішень про надання чи відмову у послугах, соціальних виплатах або адміністративних дозволах.

Реальна практика свідчить, підтверджує, що окреслені ризики є реальними загрозами, які потребують негайного нормативного реагування. Поради з відповідального використання ІІІ публічними службовцями, видані Міністерством цифрової трансформації у 2025 році, прямо застерігають від введення персональних та чутливих даних до сторонніх ІІІ-платформ, оскільки це «може поставити під загрозу не лише персональні дані, а й інформаційні системи, сервери, бази даних органів влади, де працює публічний службовець, що врешті спричинить ризики для безпеки держави» [3, с. 2]. Той факт, що держава змушена видавати подібні застереження, сам по собі красномовно свідчить про масштаб проблеми.

Принцип рівності, закріплений у статті 24 Конституції України, набуває принципово нового виміру в умовах алгоритмізованого управління [2]. Існує таке явище як «алгоритмічна упередженість» (algorithmic bias), за якої цифрові системи на основі історичних даних репродукують та посилюють соціальні нерівності. За наявності прихованого гендерного чи соціального дисбалансу в інформаційних масивах автоматизований розподіл витрат або оцінка правосуб'єктності здійснюється за дискримінаційними патернами. При цьому довести факт такої дискримінації практично унеможлиблюється, оскільки алгоритм, як правило, не пояснює логіку свого вибору. Як наголошують у своїй роботі Сухан І.С. та Сенько А.І., в процесі цифровізації «в деяких сферах ще є певні обмеження або проблеми з адаптацією для окремих категорій громадян», що підтверджується ризиками алгоритмічної упередженості, за якої автоматизовані системи репродукують соціальні нерівності [4, с. 87].

Конституційне право на відшкодування шкоди, завданої незаконними рішеннями органів державної влади, закріплено у статті 56 Конституції України [2]. У контексті автоматизації це право стикається з принципово новою юридичною проблемою - відсутністю очевидного суб'єкта відповідальності. Коли рішення про відмову в адміністративній послугі, призначенні виплати чи застосуванні санкції приймає алгоритм, питання про те, хто саме відповідає за наслідки помилки, залишається відкритим. Можливими варіантами є: держава як замовник системи; орган, що її впровадив; розробник програмного забезпечення; конкретна посадова особа, яка не здійснила належного людського нагляду за алгоритмічним рішенням. Чинне законодавство не дає однозначної відповіді на жодне з цих питань, що на практиці фактично позбавляє громадянина ефективного механізму захисту прав.

Шопіна І.М. характеризує наявний стан справ як системну прогалину, де правове регулювання використання ІІІ органами виконавчої влади характеризується «переважанням рекомендаційних механізмів» над обов'язковими нормами, а відповідальність за помилки алгоритму фактично не визначена [5, с. 296]. Це означає, що попри наявність конституційної гарантії статті 56, механізм її практичної реалізації в умовах автоматизованого управління наразі відсутній.

Особливу увагу в контексті меж автоматизації заслуговує стаття 64 Конституції України, яка встановлює, що конституційні права і свободи людини не можуть бути обмежені, крім випадків, прямо передбачених Конституцією [2]. Небезпека алгоритмізованих систем полягає в тому, що вони здатні фактично обмежувати права людини без формальної заборони, наприклад через алгоритмічне виключення зі сфери доступу до послуг, помилкову ідентифікацію як особи з ризиковим профілем або автоматичне блокування облікового запису. Таке «непомітне» обмеження прав є особливо небезпечним саме тому, що воно не фіксується як свідоме владне рішення і значно ускладнює доступ до правового захисту, тому що людина може навіть не усвідомлювати, що її права порушено.

У цьому контексті принципового значення набуває питання цифрової доступності. Частина громадян: літні люди, особи з інвалідністю, мешканці районів із поганим покриттям інтернету - об'єктивно мають обмежені можливості для використання цифрових сервісів. Якщо держава переводить надання адміністративних послуг виключно в електронний формат, не забезпечуючи альтернативних офлайн-механізмів, вона фактично створює непряму дискримінацію за ознакою цифрової компетентності - що також суперечить статтям 24 та 64 Конституції.

Найбільш категоричне конституційне обмеження щодо автоматизації закріплено у статті 124 Конституції України: правосуддя в Україні здійснюється виключно судами [2]. Ця норма є абсолютною, що не допускає ніяких винятків. У правовій доктрині це означає, що жодна система ШІ не може виносити судові рішення, навіть якщо технологічно вона є здатною до цього. ШІ може виконувати корисні допоміжні функції у судочинстві: аналізувати судову практику, готувати проєкти процесуальних документів, виявляти релевантні прецеденти, здійснювати лінгвістичний аналіз доказів. Однак рішення по суті справи, оцінка доказів та визначення правових наслідків для конкретної особи є виключною прерогативою судді як носія судової влади. Будь-яке відступлення від цього принципу несумісне з конституційним ладом України.

Узагальнюючи конституційно-правовий аналіз, можна запропонувати триступеневу модель меж автоматизації управлінських процесів у сфері публічної влади. Перший рівень - абсолютні заборони, тобто функції, які за жодних умов не можуть бути повністю делеговані алгоритму. До них належать: здійснення правосуддя (ст. 124 Конституції); прийняття рішень, що безпосередньо обмежують конституційні права і свободи (ст. 64); несанкціонована масова обробка чутливих персональних даних (ст. 32). Другий рівень - умовна допустимість: автоматизація можлива лише за наявності спеціального законодавчого регулювання, гарантованого права на оскарження рішення, доступу до пояснення логіки алгоритму та реального людського нагляду. Цей рівень охоплює, зокрема, автоматизоване нарахування соціальних виплат, визначення ризикових профілів, ухвалення адміністративних санкцій. Третій рівень - технічна підтримка рішень: ШІ виступає лише допоміжним інструментом, тоді як остаточне рішення у будь-якому разі приймає уповноважена посадова особа. Саме ця модель, на нашу думку, здатна забезпечити баланс між технологічним прогресом і конституційними гарантіями прав людини.

Отже, цифровізація публічної влади зумовлює трансформацію владних відносин і актуалізує потребу в правовому осмисленні меж автоматизації. При створенні спеціального закону про застосування ШІ в публічному управлінні, потрібно обов'язково базуватися на неухильному дотриманні конституційних засад. Закон має чітко розмежовувати допустиму автоматизацію владних функцій, визначати юридичну відповідальність за алгоритмічні помилки, забезпечувати право особи на оскарження та обґрунтування рішення системи, й гарантувати альтернативні способи отримання послуг.

Список використаних джерел

1. Про схвалення Концепції розвитку штучного інтелекту в Україні: розпорядження Кабінету Міністрів України від 02.12.2020 р. № 1556-р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-p#Text>
2. Конституція України від 28.06.1996 р. № 254к/96-ВР. URL: <https://zakon.rada.gov.ua/laws/show/254к/96-вр#Text>
3. Поради з відповідального використання штучного інтелекту публічними службовцями. Міністерство цифрової трансформації України. 2025. URL: <https://storage.thedigital.gov.ua/files/f/bf/a9595e0dcd238ab2b3602909107aabbf9.pdf>
4. Сухан І.С., Сенько А.І. Електронне урядування в Україні: правові аспекти та перспективи розвитку. *Науковий вісник Ужгородського національного університету*. 2024. № 86. Ч. 4. С. 84–88. URL: <https://visnyk-juris-uzhnu.com/wp-content/uploads/2025/01/15-3.pdf>
5. Шопіна І.М. Інформаційна безпека використання штучного інтелекту органами виконавчої влади. *Аналітично-порівняльне правознавство*. 2026. № 2. Ч. 2. С. 296–302. URL: <https://app-journal.in.ua/wp-content/uploads/2026/04/49-1.pdf>

Альбіна НАЗАРЕНКО, курсант, Національна академія Державної прикордонної служби України імені Богдана Хмельницького.

Науковий керівник:

Андрій САВЧУК, заступник начальника кафедри тактико-спеціальних дисциплін, Національна академія Державної прикордонної служби України імені Богдана Хмельницького

ШТУЧНИЙ ІНТЕЛЕКТ ЯК ІНСТРУМЕНТ НАВЧАННЯ І САМООСВІТИ: ТЕХНОЛОГІЇ, МОЖЛИВОСТІ ТА РИЗИКИ

Останніми роками штучний інтелект (ШІ) перетворився з предмета наукової фантастики на буденний інструмент, яким активно послуговуються як студенти, так і викладачі. Темп цих змін настільки стрімкий, що система освіти ще не встигла виробити усталену відповідь на виклики, які постають перед нею.

Це вимагає якомога скорішого осмислення того, як саме ШІ вписується в освітній процес — де він справді допомагає, а де непомітно підміняє те, заради чого навчання взагалі існує [1, с. 14].

Серед найбільш поширених сценаріїв використання ШІ в освіті — системи адаптивного навчання, чат-боти-тьютори та інструменти автоматизованої перевірки знань. Платформи на кшталт Khan Academy Khanmigo або Duolingo Max застосовують великі мовні моделі для персоналізації навчальних траєкторій: вони відслідковують помилки, коригують складність завдань і дають миттєвий зворотний зв'язок. Дослідження показують, що інтелектуальні навчальні системи за ефективністю наближаються до індивідуального навчання з живим тьютором і суттєво перевершують традиційний лекційний формат, — насамперед завдяки можливості підлаштовуватися під темп і прогалини кожного учня [2, с. 198]. Такі дані змушують переосмислювати роль викладача — не як єдиного носія знань, а як модератора навчального середовища.

З технологічного погляду сучасні освітні ШІ-рішення спираються на кілька ключових архітектур. Великі мовні моделі (Large Language Models, LLM) забезпечують генерацію пояснень і діалогову взаємодію; системи з доповненим пошуком (Retrieval-Augmented Generation, RAG) дозволяють прив'язати відповіді моделі до конкретного навчального контенту — підручника, курсу, бази знань установи — замість спирання виключно на попередньо навчені параметри. Це суттєво знижує ризик галюцинацій і підвищує предметну точність. Окремим напрямом є мультимодальні моделі, здатні аналізувати не лише текст, а й схеми, графіки, рукописні нотатки, що відкриває нові можливості для роботи з технічними та природничими дисциплінами. Поєднання цих підходів дає змогу будувати освітні середовища, які реагують на контекст конкретного навчального закладу, а не пропонують універсальні, знеособлені відповіді [7, с. 215].

У контексті самоосвіти ШІ відкриває можливості, які ще десять років тому здавалися нереальними. Людина може опанувати нову мову, технічну дисципліну чи наукову галузь, спираючись виключно на доступні онлайн-інструменти з інтегрованим ШІ. Генеративні моделі здатні пояснювати складні концепції у зручний для конкретного користувача спосіб, будувати індивідуальні плани навчання та імітувати діалог із фахівцем. Це особливо значимо для людей, які з різних причин позбавлені доступу до якісної формальної освіти. За даними ЮНЕСКО, станом на 2023 рік понад 244 мільйони дітей у світі не мають доступу до школи, і саме ШІ-рішення розглядаються як один із механізмів подолання цього розриву [3, с. 8].

Водночас некритичне і надмірне покладання на ШІ несе суттєві педагогічні ризики. Найбільш очевидний — витіснення власного мислення: коли студент замість формування власної аргументації звертається до чат-бота за готовою відповіддю, він фактично відмовляється від тієї когнітивної роботи, яка і є навчанням. Це явище вже отримало назву

"когнітивного аутсорсингу" і активно досліджується у педагогічній психології [4, с. 112]. Крім того, мовні моделі не позбавлені так званих галюцинацій — генерації правдоподібно виглядаючої, але фактично хибної інформації. Для студента, який ще не має достатньої предметної компетентності, щоб розпізнати помилку, це може стати джерелом системних хибних уявлень.

Питання академічної доброчесності в епоху генеративного ШІ набуло особливої гостроти. Якщо раніше плагіат стосувався переважно некоректного запозичення текстів, то тепер постала принципово нова проблема: студент може подати роботу, яка формально не скопійована, але й не написана ним самостійно. Частина університетів — зокрема, Массачусетський технологічний інститут та низка провідних європейських закладів — вже змінили свої підходи до оцінювання, переходячи до усних захистів, проєктних робіт та завдань, орієнтованих на процес, а не лише результат [5, с. 33]. В Україні це питання також поступово знаходить своє відображення у нормативних документах закладів вищої освіти, хоча єдиного підходу ще не сформовано.

Особливої уваги заслуговує досвід упровадження ШІ в українську освіту в умовах воєнного часу. Масштабне переведення закладів на дистанційний і змішаний формати навчання прискорило цифровізацію, яка за мирних умов тривала б значно довше. Водночас це навчання відбувається в умовах нерівномірного доступу до інтернету, перебоїв з електропостачанням і психологічного навантаження на всіх учасників освітнього процесу. За таких обставин ШІ-інструменти можуть виконувати роль не лише навчального асистента, а й середовища для самостійної роботи у будь-який зручний момент, незалежно від розкладу та місця перебування. Разом із тим українські заклади освіти лише починають формувати інституційні підходи до регулювання використання ШІ — і тут важливо не повторювати помилок тих країн, що вже пройшли цей шлях, а критично адаптувати їхній досвід до власних реалій [6, с. 12].

Перспективним напрямом є не заборона ШІ, а формування культури його відповідального використання — так само, як колись калькулятор не скасував математику, а лише змістив акцент із рутинних обчислень на розуміння закономірностей. Викладачі, які усвідомлено вбудовують ШІ-інструменти в освітній процес — наприклад, пропонують студентам критично оцінювати відповіді моделі, виявляти її помилки або порівнювати різні підходи до розв'язання задачі, — фактично розвивають навички критичного мислення вищого порядку [8, с. 89]. Саме ці навички в умовах інформаційного перевантаження стають ключовою компетентністю сучасної людини.

Отже, штучний інтелект у сфері освіти — це не однозначне благо і не загроза, яку варто уникати. Це дзеркало, яке відображає якість наших педагогічних цілей: якщо ми навчаємо запам'ятовувати — ШІ робить це ефективніше за нас; якщо ми навчаємо думати — він стає інструментом, а не заміником. Ключове завдання сьогодні полягає в тому, щоб освітня спільнота не пасивно реагувала на технологічні зміни, а свідомо формувала середовище, в якому ШІ підсилює людський інтелект, а не підміняє його.

Список використаних джерел

1. Selwyn N. Education and Technology: Key Issues and Debates. 2nd ed. London : Bloomsbury Academic, 2016. 232 p.
2. VanLehn K. The Relative Effectiveness of Human Tutoring, Intelligent Tutoring Systems, and Other Tutoring Systems. *Educational Psychologist*. 2011. Vol. 46, № 4. P. 197–221. DOI: 10.1080/00461520.2011.611369.
3. UNESCO. Global Education Monitoring Report 2023: Technology in Education – A Tool on Whose Terms? Paris : UNESCO Publishing, 2023. 416 p. URL: <https://www.unesco.org/gem-report/en/2023> (дата звернення: 14.03.2025).
4. Gerlich M. AI Tools in Society: Impacts on Cognitive Offloading and the Future of Critical Thinking. *Societies*. 2025. Vol. 15, № 1. P. 1–15. DOI: 10.3390/soc15010006.

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 ШАПОШНІКОВ, 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

великими обсягами інформації, виявляти приховані закономірності та адаптуватися до нових умов функціонування мережевого середовища, ці технології активно впроваджуються у сучасні системи кібербезпеки. Особливу цінність для прогнозування кібератак становить аналіз мережевого трафіку, який містить відомості про характеристики мережевих з'єднань, часові особливості передавання даних, поведінку користувачів та потенційні аномалії, що можуть свідчити про підготовку або реалізацію кібератак [3].

Наукові дослідження свідчать, що застосування методів машинного навчання дозволяє не лише підвищити ефективність виявлення вторгнень, а й забезпечити прогнозування кіберзагроз на ранніх стадіях їх виникнення. Такий підхід дає змогу завчасно ідентифікувати потенційно небезпечні тенденції в мережевому середовищі та мінімізувати ризики реалізації атак ще до настання негативних наслідків [4].

Класичні алгоритми машинного навчання в задачах прогнозування кібератак. Серед найбільш поширених підходів до аналізу мережевого трафіку особливе місце займають класичні алгоритми машинного навчання. Їх популярність пояснюється відносною простотою реалізації, невисокими вимогами до обчислювальних ресурсів та можливістю інтерпретації отриманих результатів.

Одним із найвідоміших алгоритмів є метод опорних векторів (Support Vector Machine), який забезпечує ефективне розділення об'єктів різних класів шляхом побудови оптимальної гіперплощини. У системах кібербезпеки цей підхід використовується для класифікації мережевого трафіку на безпечний та потенційно небезпечний. Разом із тим ефективність методу значною мірою залежить від правильного вибору параметрів моделі та функції ядра [1].

Не менш поширеним є алгоритм k-найближчих сусідів, який базується на визначенні схожості між об'єктами у просторі ознак. Завдяки цьому він успішно застосовується для виявлення аномальної мережевої активності. Однак зі збільшенням обсягів даних та кількості ознак ефективність алгоритму поступово знижується, що обмежує його використання в масштабних мережевих системах [9].

Наївний байєсівський класифікатор також залишається популярним інструментом аналізу кіберзагроз. Його перевагами є висока швидкість роботи та простота впровадження. Водночас базове припущення про незалежність ознак часто не відповідає реальним характеристикам мережевого трафіку, що негативно впливає на точність прогнозування [1].

Загалом класичні методи машинного навчання залишаються актуальними завдяки можливості швидкого розгортання та невисоким вимогам до ресурсів, однак їх ефективність поступається сучасним ансамблевим і глибоким моделям при роботі з великими наборами даних та складними сценаріями кіберзагроз [2].

Ансамблеві моделі як засіб підвищення точності прогнозування. Важливим етапом розвитку інтелектуальних систем кібербезпеки стало впровадження ансамблевих методів машинного навчання. Основна ідея таких підходів полягає у поєднанні кількох базових моделей для формування більш точного та стійкого прогнозу.

Одним із найефективніших ансамблевих алгоритмів є випадковий ліс (Random Forest), який являє собою сукупність дерев рішень, навчених на різних підмножинах даних. Такий підхід дозволяє зменшити ризик перенавчання, підвищити узагальнювальну здатність моделі та забезпечити високу стійкість до шумів у даних [5].

Іншим популярним рішенням є методи градієнтного бустингу, серед яких найбільш відомими є XGBoost та LightGBM. Ці алгоритми будують послідовність моделей, кожна з яких коригує помилки попередньої. Завдяки цьому забезпечується висока точність класифікації навіть у випадках, коли набір даних є незбалансованим, що є характерною особливістю задач кібербезпеки [6].

Дослідження свідчать, що ансамблеві методи забезпечують кращі результати прогнозування порівняно з окремими класифікаторами. Крім того, вони дозволяють досягти оптимального балансу між точністю, швидкодією та стійкістю системи до змін у структурі

мережевого трафіку. Саме тому ансамблеві моделі сьогодні активно використовуються в системах моніторингу мережевої безпеки та автоматизованого реагування на інциденти [5; 6].

Застосування методів глибокого навчання для аналізу мережевого трафіку. Суттєвий прорив у сфері прогнозування кіберзагроз пов'язаний із використанням технологій глибокого навчання. На відміну від класичних алгоритмів, глибокі нейронні мережі здатні автоматично формувати інформативні ознаки на основі сирих даних без необхідності їх попереднього ручного відбору [2].

Одним із найбільш перспективних напрямів є застосування згорткових нейронних мереж (CNN). Такі моделі дозволяють виявляти складні просторові залежності між характеристиками мережевого трафіку та ефективно класифікувати різні типи мережевої активності [8].

Особливий інтерес для прогнозування кібератак становлять рекурентні нейронні мережі (RNN) та їх удосконалені модифікації типу Long Short-Term Memory (LSTM). Головною перевагою таких моделей є здатність враховувати часові залежності між подіями. Оскільки більшість сучасних атак реалізується поетапно та має характерні часові закономірності, використання LSTM дозволяє прогнозувати розвиток атаки ще до її активної фази [7].

Рекурентні мережі здатні аналізувати послідовності мережевих подій, накопичувати інформацію про попередні стани системи та враховувати їх під час формування прогнозу. Завдяки цьому досягається значно вища точність прогнозування порівняно з традиційними алгоритмами [7; 8].

Разом із тим застосування методів глибокого навчання пов'язане з низкою труднощів. Насамперед це висока обчислювальна складність, необхідність використання потужних апаратних ресурсів та потреба у великих обсягах навчальних даних. Крім того, нейронні мережі часто функціонують як так звані «чорні скриньки», що ускладнює інтерпретацію результатів та пояснення прийнятих рішень [2].

Часові моделі та прогнозування кіберзагроз. Важливу роль у прогнозуванні кіберінцидентів відіграють методи аналізу часових рядів. Такі підходи дозволяють досліджувати закономірності зміни мережевої активності у часі та прогнозувати ймовірність виникнення кібератак у майбутньому [3].

До традиційних часових моделей належать статистичні методи, зокрема ARIMA. Їх перевагою є відносна простота побудови та інтерпретації результатів. Водночас вони не завжди здатні адекватно відображати складні нелінійні залежності, характерні для сучасного мережевого трафіку [4].

Значно кращі результати демонструє поєднання часових моделей із методами глибокого навчання. Такі гібридні рішення забезпечують високу точність прогнозування, адаптивність до нових загроз та здатність враховувати складні закономірності розвитку кіберінцидентів. Саме тому вони розглядаються як один із найперспективніших напрямів розвитку систем прогнозування кіберзагроз [3; 4].

Особливості підготовки даних та оцінювання ефективності моделей. Якість прогнозування значною мірою залежить від характеристик наборів даних, які використовуються для навчання моделей. У сучасних дослідженнях найбільш поширеними є датасети NSL-KDD, UNSW-NB15 та CICIDS2017, що містять як нормальну мережеву активність, так і різноманітні типи атак [10].

Перед використанням даних необхідно виконати їх попередню обробку, яка включає очищення від шумів, нормалізацію числових параметрів, кодування категоріальних ознак та балансування класів. Недотримання цих процедур може суттєво погіршити результати прогнозування та призвести до формування упереджених моделей [10].

Для оцінювання якості прогнозних моделей використовуються такі показники, як Accuracy, Precision, Recall та F-measure. Особливе значення для систем кібербезпеки мають показники повноти та F-міри, оскільки вони дозволяють оцінити здатність моделі виявляти максимальну кількість реальних атак при мінімізації кількості хибних спрацювань. Додатково

застосовується аналіз ROC-кривих та матриці помилок, що забезпечує комплексну оцінку ефективності моделей [10].

Висновки. Проведений аналіз свідчить, що використання методів машинного навчання відкриває широкі можливості для створення проактивних систем кібербезпеки, здатних прогнозувати виникнення кіберзагроз на основі аналізу мережевого трафіку. Класичні алгоритми забезпечують простоту реалізації та інтерпретації результатів, ансамблеві методи демонструють високий рівень точності та стабільності, а моделі глибокого навчання дозволяють враховувати складні закономірності й часові залежності між подіями [1; 5; 7].

Найбільш перспективним напрямом розвитку сучасних систем прогнозування кібератак є створення гібридних моделей, які поєднують переваги різних підходів. Такі системи здатні забезпечити високу точність прогнозування, адаптивність до нових типів загроз та прийнятну обчислювальну складність. Подальші дослідження доцільно спрямувати на інтеграцію методів пояснюваного штучного інтелекту, удосконалення механізмів аналізу часових рядів та впровадження адаптивних систем прогнозування кіберзагроз у практику забезпечення інформаційної безпеки [11].

Список використаних джерел

1. Buczak A. L., Guven E. *A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection*. IEEE Communications Surveys & Tutorials, 2016.
2. Berman D. S., Buczak A. L., Chavis J. S., Corbett C. L. *A Survey of Deep Learning Methods for Cyber Security*. Information (MDPI), 2019.
3. Landauer M., Skopik F., Stojanović B., Flatscher A., Ullrich T. *A Review of Time-Series Analysis for Cyber Security Analytics: From Intrusion Detection to Attack Prediction*. International Journal of Information Security, 2025.
4. Werner G., Yang S., McConky K. *Time Series Forecasting of Cyber Attack Intensity*. ACM CISRC, 2017.
5. Sommer R., Paxson V. *Outside the Closed World: On Using Machine Learning for Network Intrusion Detection*. IEEE Symposium on Security and Privacy, 2010. P. 305–316.
6. Chandola V., Banerjee A., Kumar V. *Anomaly Detection: A Survey*. ACM Computing Surveys. 2009. Vol. 41(3). Article 15.
7. Kim G., Lee S., Kim S. *A Novel Hybrid Intrusion Detection Method Integrating Anomaly Detection with Misuse Detection*. Expert Systems with Applications. 2014. Vol. 41(4). P. 1690–1700.
8. Shone N., Ngoc T. N., Phai V. D., Shi Q. *A Deep Learning Approach to Network Intrusion Detection*. IEEE Transactions on Emerging Topics in Computational Intelligence. 2018. Vol. 2(1). P. 41–50.
9. Lin W., Ke S., Tsai C. *CANN: An Intrusion Detection System Based on Combining Cluster Centers and Nearest Neighbors*. Knowledge-Based Systems. 2015. Vol. 78. P. 13–21.
10. Aljawarneh S., Aldwairi M., Yassein M. B. *Anomaly-Based Intrusion Detection System Through Feature Selection Analysis and Building Hybrid Efficient Model*. Journal of Computational Science. 2018. Vol. 25. P. 152–160.
11. Khan S., Lloret J., Kaur K. *A Novel Cyber Attack Detection System Using Deep Learning Techniques*. Electronics. 2020. Vol. 9(10). Article 1569.

Анна РАБІЙЧУК, студентка н. гр. 102 СПД ННІ права та психології, Національна академія внутрішніх справ.

Анна ОМЕЛЯН, студентка н. гр. 102 СПД ННІ права та психології, Національна академія внутрішніх справ

Науковий керівник:

Вадим КУДІНОВ, завідувач кафедри інформаційних технологій ННІ права та психології, Національна академія внутрішніх справ, кандидат фізико-математичних наук, доцент

ІІІ-КОМПЕТЕНТНОСТІ СПІВРОБІТНИКА ПРАВООХОРОННОГО ОРГАНУ У КОНТЕКСТІ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ ПРАВООХОРОННОЇ СИСТЕМИ

Штучний інтелект (ШІ) сьогодні проникає майже в усі сфери життя. Правоохоронна діяльність не виняток. Системи розпізнавання облич, аналітичні платформи, автоматизація документообігу – усе це поступово стає реальністю і для поліції. Але технології самі по собі не працюють. Головне – це людина, яка ними користується. Тому питання про те, якими знаннями та навичками має володіти правоохоронець для роботи з ШІ, стає дуже актуальним.

Мета роботи – визначити структуру ІІІ-компетентностей, необхідних співробітнику правоохоронних органів, та окреслити можливі шляхи їх формування.

Завдання дослідження: проаналізувати сучасні тенденції впровадження ШІ в роботу поліції в різних країнах; виокремити ключові компоненти ІІІ-компетентностей; виявити основні проблеми, що виникають при їх формуванні, та запропонувати шляхи вирішення.

Аналіз зарубіжних джерел показує, що штучний інтелект вже активно використовується в роботі поліції багатьох країн. Розглянемо найбільш показові приклади.

У США Національний інститут поліції реалізує проєкт з використання ШІ для аналізу відео з натільних камер поліцейських. Система оцінює поведінку правоохоронця під час затримань – його мовлення, тон, міміку, жести, а також реакцію затриманого. Отримані дані використовуються не тільки для контролю, але й для навчання курсантів у поліцейських академіях. ШІ тут виступає не стільки контролером, скільки помічником і тренером.

У Європі реалізується проєкт LAW-GAME, який фінансується Європейським Союзом. Це інтерактивна навчальна платформа, що використовує гейміфікацію та штучний інтелект. Поліцейські можуть відпрацьовувати навички ведення допитів, тактику переговорів, розпізнавання ознак підготовки терористичних актів. ШІ моделює поведінку підозрюваних, адаптується до дій поліцейського, створюючи максимально реалістичні сценарії.

Європейське агентство з підготовки працівників правоохоронних органів включило розвиток цифрових навичок, зокрема пов'язаних зі ШІ, до своєї стратегії на 2023-2027 роки. Важливо навчити поліцейських користуватися готовими ШІ-інструментами, сформулювати в них розуміння принципів роботи алгоритмів, їхніх обмежень та етичних ризиків.

У Великій Британії College of Policing випустив методичні рекомендації щодо використання штучного інтелекту в поліцейській діяльності. Вони охоплюють процедури тестування та валідації ШІ-систем перед впровадженням, а також етичні принципи роботи з персональними даними. Особлива увага приділяється проблемі "чорної скриньки" – коли алгоритм приймає рішення, але логіка цього рішення залишається незрозумілою. У таких випадках використання ШІ рекомендується обмежувати або забороняти.

Китайські дослідники пропонують концепцію "алгоритмічної грамотності" поліцейських. Це поняття включає здатність розуміти принципи роботи алгоритмів, критично оцінювати результати їхньої діяльності, усвідомлювати можливі помилки та упередження, а також вміння правильно формулювати запити до ШІ і інтерпретувати отримані відповіді.

Отже, світова практика свідчить: впровадження ШІ в роботу поліції відбувається досить активно, охоплюючи різні напрями – від аналітики до навчання. І паралельно з цим зростає розуміння того, що ключову роль відіграє людина, яка користується технологіями.

На основі аналізу джерел можна виділити кілька ключових блоків ШІ-компетентностей. Важливо підкреслити, що йдеться не про підготовку програмістів, а про формування базової грамотності, яка дозволяє ефективно, безпечно та відповідально використовувати ШІ-інструменти.

Перший блок – технічний. Працівник повинен розуміти базові принципи роботи штучного інтелекту. Не потрібно знати всі математичні моделі, але варто орієнтуватися в тому, що таке машинне навчання, чим воно відрізняється від звичайного програмування, які дані потрібні для навчання ШІ, чому алгоритми можуть помилятися. Крім того, необхідно знайомство з конкретними ШІ-інструментами, які використовуються в правоохоронній діяльності, а також вміння коректно формувати запити до таких систем.

Другий блок – аналітичний. Він є, мабуть, найважливішим. Працівник має вміти критично оцінювати результати роботи ШІ, спираючись на професійний досвід та здоровий глузд. Наприклад, якщо система розпізнавання обличчя із високою впевненістю ідентифікує людину, яка фізично не могла перебувати в цьому місці, співробітник повинен це помітити. До цього блоку також належить розуміння обмежень технологій – що якість роботи ШІ залежить від якості даних, на яких його навчали, і що він може відтворювати та посилювати упередження, наявні в суспільстві.

Третій блок – етично-правовий. Використання ШІ в правоохоронній діяльності пов'язане з обробкою персональних даних, часто дуже чутливих. Працівник повинен знати вимоги законодавства щодо захисту таких даних, розуміти, які дії допустимі, а які – ні. Важливо усвідомлювати проблему алгоритмічної упередженості. Дослідження показують, що системи розпізнавання обличчя мають вищу похибку для людей з темним кольором шкіри або для жінок. Якщо працівник не знає про це, він може несвідомо порушувати права громадян.

Четвертий блок – комунікативний. Працівник, який використовує ШІ, повинен вміти пояснити громадянам, чому і яким чином застосовуються ці технології, чому алгоритм видав саме такий результат, наскільки йому можна довіряти. Якщо поліцейський не може пояснити це в суді, доказ, отриманий за допомогою ШІ, можуть визнати недопустимим.

П'ятий блок – безпековий. Він стосується розуміння ризиків, пов'язаних із використанням неофіційних ШІ-інструментів. Багато співробітників використовують ChatGPT та подібні сервіси для робочих завдань, не усвідомлюючи, що завантажені дані можуть бути використані для навчання моделей і стати доступними третім особам. Працівник має знати внутрішні політики щодо використання ШІ та розуміти відповідальність за їх порушення.

Отже, **ШІ-компетентності** – це комплексне утворення, яке включає технічну, аналітичну, етично-правову, комунікативну та безпекову складові. Вони не зводяться до вузьких технічних навичок, а передбачають здатність критично мислити, діяти в межах закону, відповідати за свої рішення та ефективно взаємодіяти з технологіями.

Формування ШІ-компетентностей стикається з низкою серйозних проблем.

Розглянемо основні.

Перша проблема – швидкість технологічного розвитку. Штучний інтелект змінюється так швидко, що навчальні програми просто не встигають за ним. Те, що було актуальним два роки тому, сьогодні вже застаріло. Оновлення програм – це довгий бюрократичний процес, який може займати місяці або навіть роки.

Виникає розрив між тим, чого вчать, і тим, що потрібно на практиці.

Друга проблема – брак викладачів. Щоб навчати ШІ-компетентностям, потрібні люди, які розуміються і на праві, і на технологіях. А таких фахівців дуже мало. Зазвичай або це юристи, які поверхово знайомі з ІТ, або ІТ-фахівці, які не знають специфіки правоохоронної діяльності. До того ж, викладачі самі потребують підвищення кваліфікації.

Третя проблема – різний рівень підготовки курсантів. Хтось вільно володіє комп'ютером з дитинства, а хтось має лише базові навички. Викладати матеріал для такої різномірної аудиторії дуже складно. Орієнтуєшся на сильних – слабкі не розуміють; спускаєшся до рівня слабких – сильним стає нудно.

Четверта проблема – відсутність стандартів. На сьогодні немає загальноприйнятого переліку ІІІ-компетентностей для правоохоронців, немає єдиних підходів до їх оцінювання. Кожен навчальний заклад діє по-своєму, що призводить до розпорошення зусиль.

П'ята проблема – "тіньове" використання ІІІ. Співробітники, які не мають офіційного доступу до необхідних інструментів, починають використовувати безкоштовні версії ChatGPT та подібних сервісів. Вони завантажують туди документи з персональними даними, іноді навіть секретними. Це створює серйозні загрози інформаційній безпеці.

Що можна запропонувати як шляхи вирішення?

Необхідно переглянути підходи до підготовки викладачів. Замість пошуку готових фахівців, яких майже немає, варто навчати власних – відбирати викладачів-правознавців із базовими ІТ-навичками та систематично підвищувати їхню кваліфікацію. Варто запровадити модульний підхід до навчання. Базовий модуль є обов'язковим для всіх – там дають загальне розуміння ІІІ, його можливостей та ризиків. Потім ідуть спеціалізовані модулі для різних категорій працівників – для слідчих, для патрульних, для оперативників. Це дозволяє врахувати різний рівень підготовки.

Потрібно розробити та запровадити відомчі стандарти ІІІ-компетентностей. У них мають бути чітко прописані вимоги до знань і навичок для різних категорій працівників.

Це стане основою для розробки навчальних програм і системи оцінювання. Слід розробити чітку політику щодо використання ІІІ в правоохоронних органах. Вона має визначати, які ІІІ-інструменти дозволені, за яких умов, які дані можна завантажувати. І, звісно, дати співробітникам офіційний доступ до безпечних інструментів.

Необхідно активно використовувати практико-орієнтовані методи навчання. Теоретичних лекцій недостатньо. Потрібні кейси, симуляції, тренінги, де курсанти можуть попрацювати з реальними або наближеними до реальних ІІІ-системами.

Підсумовуючи результати дослідження, можна зробити кілька **основних висновків**.

По-перше, штучний інтелект вже сьогодні активно впроваджується в правоохоронну діяльність різних країн. Системи розпізнавання облич, аналітичні платформи, автоматизація документообігу, навчальні симулятори – це не майбутнє, а сучасність. Відповідно, питання готовності працівників до роботи з цими технологіями стає цілком практичним.

По-друге, ІІІ-компетентності – це комплексне утворення, яке включає технічну, аналітичну, етично-правову, комунікативну та безпекову складові.

Працівник має не лише розуміти, як працює ІІІ, але й критично оцінювати його результати, усвідомлювати ризики упередженості, діяти в межах закону, вміти пояснити свої дії громадянам, а також відповідально ставитися до безпеки даних.

По-третє, формування ІІІ-компетентностей стикається з серйозними викликами: стрімкий розвиток технологій, брак кваліфікованих викладачів, неоднорідність базової підготовки курсантів, відсутність чітких стандартів, "тіньове" використання ІІІ співробітниками. Однак ці проблеми не є нездоланими.

По-четверте, основними напрямками роботи мають стати: оновлення програм підготовки викладачів, запровадження модульного підходу до навчання, розробка відомчих стандартів ІІІ-компетентностей, створення чітких політик використання ІІІ в правоохоронній діяльності та активне використання практико-орієнтованих методів навчання.

Отже, ІІІ-компетентності є необхідною складовою професійної підготовки сучасного правоохоронця. Їх формування потребує системних зусиль, але без них цифрова трансформація правоохоронної системи залишиться неповною. Чим швидше будуть зроблені перші кроки в цьому напрямку, тим меншим буде розрив між технологічними можливостями та реальною практикою їх використання.

Список використаних джерел

1. Інтерактивний, колаборативний підхід цифрової гейміфікації до ефективного експериментального навчання та прогнозування кримінальних дій. European Commission : [сайт]. URL: <https://cordis.europa.eu/project/id/101021714> (дата звернення: 22.05.2026).
2. Рада директорів приймає нову стратегію CEPOL. CEPOL : [сайт]. URL: <https://www.cepol.europa.eu/newsroom/news/management-board-adopts-new-cepol-strategy> (дата звернення: 22.05.2026).
3. Рекомендації з використання технологій штучного інтелекту в обробці персональних даних: найкращі практики ЄС. EU4DIGITALUA : [сайт]. URL: https://eu4digitalua.eu/wp-content/uploads/2025/01/ai_guidelines_ua.pdf (дата звернення: 22.05.2026).
4. Чень Чжисюань, Хань Чуньмей, Ван Жуншен. Алгоритмічна грамотність поліцейських у цифрову епоху: теоретичне обґрунтування та практична цінність. *Журнал Шаньсійської поліцейської академії*, 2025. No 3. С. 45-52.
5. Стан розвитку штучного інтелекту в поліції. UNITE.AI : [сайт]. URL: <https://www.unite.ai/uk/the-state-of-ai-in-policing/>

Максим ПАЦУЛА, курсант 3-го курсу факультету підготовки фахівців для органів досудового розслідування Національної поліції України, Дніпровський державний університет внутрішніх справ.

Науковий керівник:

Дмитро БОДИРЄВ, старший викладач кафедри вогневої підготовки, Дніпровський державний університет внутрішніх справ

ЗАСТОСУВАННЯ ЦИФРОВИХ ТРЕНАЖЕРІВ ТА VR-ТЕХНОЛОГІЙ У СИСТЕМІ ВОГНЕВОЇ ПІДГОТОВКИ ПОЛІЦЕЙСЬКИХ

Застосування цифрових тренажерів та VR-технологій у системі вогневої підготовки поліцейських є одним із перспективних напрямів модернізації професійної підготовки поліцейських. Традиційна вогнева підготовка, що ґрунтується на вивченні матеріальної частини зброї, правил безпеки, правових підстав застосування вогнепальної зброї та практичному виконанні вправ зі стрільби, залишається базовою складовою навчання. Однак сучасні умови поліцейської діяльності вимагають не лише технічного володіння зброєю, а й здатності швидко оцінювати обстановку, приймати законні та пропорційні рішення, розрізняти рівень загрози, діяти в умовах стресу, обмеженого часу, недостатньої видимості, присутності цивільних осіб або ризику помилкової ідентифікації небезпеки.

Цифрові стрілецькі тренажери забезпечують можливість багаторазового відпрацювання базових і складніших елементів вогневої підготовки без постійного використання бойових набоїв, що має як економічне, так і організаційне значення. У таких системах можуть фіксуватися точність наведення, час реагування, стабільність дій, помилки у прийнятті рішення, рівень концентрації та динаміка результатів конкретного слухача. Це дає змогу перейти від формального оцінювання за принципом «влучив — не влучив» до більш глибокого аналізу професійної готовності поліцейського. Наприклад, інструктор може бачити не лише кінцевий результат вправи, а й те, як саме особа діяла: чи поспішала, чи правильно оцінила умовну загрозу, чи дотримувалася безпечної поведінки, чи не здійснила необґрунтованої дії. Важливо, що цифрові тренажери можуть використовуватися на початкових етапах навчання,

коли курсант або поліцейський ще не має достатнього рівня впевненості, а також на етапах підтримання професійних навичок, коли необхідне регулярне тренування без значних матеріальних витрат [1, с. 370].

VR-технології мають ще ширші можливості, оскільки створюють ефект присутності в умовній службовій ситуації. На відміну від звичайного тирю або стандартного тренажера, віртуальна реальність дозволяє моделювати багатофакторні сценарії: перевірку документів, зупинку транспортного засобу, реагування на домашнє насильство, затримання правопорушника, ситуацію масового скупчення людей, напад на поліцейського або необхідність захисту третьої особи. Цінність таких технологій полягає в тому, що вони формують не механічну навичку застосування зброї, а професійну поведінкову модель. Поліцейський навчається оцінювати контекст: чи є реальна загроза життю і здоров'ю, чи можна застосувати менш суворий захід примусу, чи достатньо усного попередження, чи потрібно змінити позицію, викликати підкріплення або забезпечити евакуацію цивільних. Таким чином, VR-підготовка сприяє розвитку правового мислення, стресостійкості, ситуаційної обізнаності та відповідальності за наслідки службових рішень [2, с. 585].

Конкретний аналіз ефективності цифрових тренажерів і VR-технологій показує, що їх доцільно розглядати не як заміну реальній стрільбі, а як додатковий інструмент підвищення якості вогневої підготовки. Реальна стрільба необхідна для формування практичного відчуття зброї, контролю її фізичних характеристик, звуку пострілу, віддачі та психологічної адаптації до використання вогнепальної зброї. Водночас традиційний тир не завжди дає можливість відпрацювати складні службові ситуації, де головною проблемою є не технічна точність, а правильність рішення. Саме тут цифрові системи мають перевагу: вони дозволяють створювати варіативні сценарії, змінювати поведінку умовних учасників події, аналізувати реакцію поліцейського та проводити детальний розбір помилок після тренування. Наприклад, якщо особа надто швидко переходить до умовного застосування зброї без достатнього оцінювання обставин, це може свідчити про недостатній рівень правової та психологічної готовності. Якщо ж поліцейський надмірно зволікає у ситуації очевидної небезпеки, це також потребує корекції підготовки. Отже, цифрові тренажери дають можливість індивідуалізувати навчання та виявляти слабкі місця не лише у стрілецьких, а й у когнітивних і поведінкових навичках [3, с. 257].

Водночас упровадження VR-технологій у систему вогневої підготовки поліцейських має супроводжуватися чітким методичним, правовим і психологічним забезпеченням. Недостатньо просто закупити обладнання: необхідно розробити навчальні сценарії, що відповідають реальним службовим ситуаціям, українському законодавству, стандартам прав людини та принципам законності, необхідності, пропорційності й обґрунтованості застосування примусу. Також важливо підготувати інструкторів, які зможуть не лише технічно обслуговувати тренажери, а й професійно аналізувати дії слухачів, пояснювати помилки та формувати правильні алгоритми поведінки. Перспективним є поєднання VR-тренувань із психологічною підготовкою, тактичною медициною, комунікативними навичками деескалації конфлікту та аналізом правових наслідків застосування зброї.

Отже, цифрові тренажери та VR-технології можуть істотно підвищити якість вогневої підготовки поліцейських, якщо вони використовуються системно, не формально, а як частина комплексної моделі професійної підготовки, спрямованої на безпечне, законне й відповідальне виконання поліцейських повноважень.

Список використаних джерел

1. Тимофєєв В. П., А. Д. Наточій, Ю. М. Толмачова. Роль БПЛА і мультимедія в підготовці поліцейських до служби в умовах воєнного стану. Науковий вісник Ужгородського національного університету. Серія: Право 3.92. 2025. С. 369-375.
2. Громик І.О. Безпека під час проведення занять з вогневої підготовки. The 1 st International scientific and practical conference "Innovation and development in world science"(November 3-5, 2025) MDPC Publishing, Zurich, Switzerland. 2025. 585 p.

3. Степаненко В. Інтеграція сучасних технологій у процес вогневої підготовки: проблеми, можливості та перспективи для збройних сил України. Актуальні проблеми службово-бойової діяльності сил сектору безпеки і оборони України: матеріали Всеукр. 2025. С. 257.

Михайло ЛИТВИНЕНКО, здобувач вищої освіти 2-го курсу ФПФПКП НПУ, Дніпровський державний університет внутрішніх справ.

Науковий керівник:

Едуард РИЖКОВ, професор кафедри інформаційних технологій, Дніпровський державний університет внутрішніх справ, кандидат юридичних наук, професор

МЕТОДИ АНАЛІТИКИ В ДІЯЛЬНОСТІ ПІДРОЗДІЛІВ КРИМІНАЛЬНОЇ ПОЛІЦІЇ

Ефективність протидії злочинності в сучасних умовах прямо корелює зі здатністю оперативних підрозділів перетворювати окремі дані в цілісну картину у структуровані масиви знань. Методи аналітики в діяльності кримінальної поліції - це не просто технічні прийоми, а методологічна основа для виявлення прихованих закономірностей, прогнозування криміногенних процесів та встановлення особи за їхнім «цифровим» чи «поведінковим» слідом. Сучасна аналітична робота базується на інтеграції декількох ключових методів.

Головна відмінність кримінальної поліції від інших служб полягає в тому, що вона не просто фіксує порушення, а «полює» на професійних злочинців та розплутує складні схеми. Якщо звичайна поліція частіше працює з тим, що вже сталося на очах, то кримінальна - занурюється в деталі, які на перший погляд непомітні. Вона шукає приховані зв'язки, вистежує організаторів і намагається зрозуміти логіку злочинця, щоб випередити його наступний крок. Це робота не стільки силова, скільки інтелектуальна, де головною зброєю є вміння скласти до купи розрізнені факти.

Оперативний аналіз та спосіб мислення від загального до конкретного (Case Analysis). Цей метод є фундаментальним для розкриття конкретних злочинів. Він базується на детальному вивченні матеріалів кримінального провадження через призму послідовності подій у часі. Аналітик не просто фіксує події, а вибудовує часові послідовності (Анасара-схеми), що дозволяє виявити логічні розбіжності в алібі або показах. Важливо розуміти, що оперативний аналіз дозволяє знайти відсутність даних про дії людини в певний час - моменти, де активність об'єкта не зафіксована, що часто вказує на час підготовки або приховування знарядь злочину. Кримінальний аналіз є невід'ємною частиною сучасної правоохоронної діяльності, оскільки він дозволяє систематизувати величезні масиви інформації та забезпечувати прийняття обґрунтованих рішень [1, с. 42].

Візуальний аналіз мережевих зв'язків (Link Analysis). У контексті боротьби з організованою злочинністю цей метод не має альтернатив. Він дозволяє перетворити сотні протоколів допитів та результатів негласних слідчих розшукових дій у зрозумілі графічні моделі. Вузлами такої моделі виступають не лише люди, а й об'єкти: спільні адреси реєстрації, спільні транспортні засоби чи телефонні контакти. Це дає змогу виявити «центральні вузли» (лідерів) та «периферійних виконавців». Логічна перевага методу в тому, що він підсвічує непрямі зв'язки, які не очевидні при звичайному читанні справ. Використання спеціалізованого програмного забезпечення для візуалізації зв'язків дає можливість виявити ключові фігури в злочинних мережах, які раніше залишалися поза полем зору [2, с. 115].

Геопросторовий аналіз та теорія «гарячих точок» (Crime Mapping). Метод базується на постулаті, що злочинність не поширюється хаотично. Використовуючи цифрові карти, аналітики кримінальної поліції проводять кластерний аналіз. Виявлення територіальних центрів тяжіння злочинності (наприклад, поблизу ломбардів або транспортних вузлів) дозволяє не лише ефективно розподіляти сили патрулювання, а й проводити превентивні заходи. Це логічний перехід від реагування на події до управління територіальними ризиками. Використання геоінформаційних систем дозволяє візуалізувати просторовий розподіл злочинності, що є ключовим для раціонального розподілу сил і засобів [3, с. 88].

Аналіз методів вчинення злочинів (Modus Operandi Analysis). Суть цього методу полягає у виявленні унікального «почерку» злочинця або групи. Це аналіз стійких звичок: спосіб зламу, вибір конкретного часу, використання специфічного сленгу під час нападів чи навіть залишені на місці специфічні сліди. У сучасних умовах до цього додається аналіз цифрового почерку (IP-адреси, типи використовуваного шкідливого ПЗ). Поєднання епізодів за методом Modus Operandi дозволяє об'єднувати провадження в різних регіонах в одну серію, що значно підвищує шанси на затримання. Алгоритми порівняння патернів дозволяють виявляти складні закономірності злочинної поведінки, значно прискорюючи процес ідентифікації серійних правопорушників [4, с. 67].

Прогностичний та стратегічний аналіз. Він базується на статистичних методах і дозволяє прогнозувати динаміку злочинності на основі зовнішніх чинників (економічних показників, міграційних процесів). Логіка прогнозування дозволяє кримінальній поліції готуватися до появи нових видів шахрайств або сплесків вуличної злочинності ще до того, як вони набудуть масового характеру. Сучасні методи превентивного аналізу спрямовані на оцінку ризиків та вразливостей, що дозволяє переходити до стратегії активного запобігання злочинності [4, с. 102].

Окрему увагу в контексті технологічного оновлення кримінальної поліції слід приділити перспективам інтеграції штучного інтелекту (AI) у процеси OSINT-розвідки. Гармонійне поєднання штучного інтелекту з технологіями збору даних із відкритих джерел дозволяє не лише автоматизувати рутинний моніторинг соцмереж та форумів, а й значно прискорює ідентифікацію підозрюваних через аналіз великих масивів візуальної та текстової інформації [5, с. 520]. Використання таких інструментів, як комп'ютерний зір для аналізу супутникових знімків чи алгоритмів розпізнавання облич (наприклад, Clearview AI), трансформує OSINT із допоміжного методу на потужний автономний механізм, здатний знаходити приховані взаємозв'язки та фільтрувати інформаційний «шум» в умовах воєнного стану. При цьому, попри високий рівень автоматизації початкових стадій аналізу, ключова роль на завершальному етапі все одно залишається за фахівцем-аналітиком, чия професійна оцінка є необхідною для перевірки інформації на правдивість та їхнього юридичного закріплення [5, с. 521].

Методи аналітики в діяльності кримінальної поліції пройшли еволюцію від простої систематизації карток до складних інтелектуальних систем підтримки прийняття рішень. Комбінація методів аналізу зв'язків, геокартування та вивчення Modus Operandi створює умови, за яких злочинна діяльність втрачає свою анонімність. Для правоохоронної системи це означає перехід на якісно новий рівень, де професіоналізм базується не лише на досвіді, а й на суворому математичному та логічному аналізі даних, що мінімізує помилки та забезпечує невідворотність покарання.

Список використаних джерел

1. Федча І.А. Професійна підготовка аналітиків. Навчальний посібник з основ кримінального аналізу для правоохоронних органів України. К.: МВС, 2023. 256 с.
2. Шепітько В.Ю. Діджиталізація поліцейської діяльності // Сучасні інструменти кримінального аналізу: Збірник матеріалів науково-практичної конференції. Дніпро, 2024. 312 с.
3. Василичук В.І. Технології безпеки в 21 столітті. Методичні рекомендації щодо використання аналітичних систем в оперативно-розшуковій діяльності. Л.: Світ, 2024. 198 с.
4. Колодяжний М.Г. Стратегії запобігання злочинності. Аналітичний огляд світового досвіду впровадження превентивних поліцейських заходів. К.: МВС України, 2023. 144 с.

5. Рижков Е. В. Перспективи використання штучного інтелекту при реалізації технології OSINT / Е. В. Рижков // Міжнародна та національна безпека: теоретичні і прикладні аспекти: матеріали ІХ Міжнар. наук.-практ. конф. (м. Дніпро, 21 березня 2025 р.); у 2-х ч. Ч. II. Дніпро : ДДУВС, 2025. С. 519-521.

Валерія БУНЧУЖНА, курсантка 2-го курсу факультету підготовки фахівців для підрозділів кримінальної поліції Національної поліції України, Дніпровський державний університет внутрішніх справ.

Науковий керівник:

Марина ЛОГІНОВА, викладач кафедри цивільно правових дисциплін, Дніпровський державний університет внутрішніх справ

НІКЧЕМНІ ПРАВОЧИНИ ТА ЇХ ПРАВОВІ НАСЛІДКИ

Нікчемні правочини є одним із ключових інститутів цивільного права України, оскільки вони безпосередньо пов'язані із забезпеченням законності цивільного обороту, захистом прав учасників майнових і немайнових відносин та недопущенням використання правочинів у супереч закону. Відповідно до ст. 202 Цивільного кодексу України (далі – ЦК України), правочином є дія особи, спрямована на набуття, зміну або припинення цивільних прав та обов'язків. Однак не кожен правочин породжує юридичні наслідки. Якщо під час його вчинення порушуються вимоги закону, такий правочин може бути недійсним. Особливе місце серед недійсних правочинів займають саме нікчемні правочини, тобто такі, недійсність яких прямо встановлена законом і які не потребують окремого судового рішення для визнання їх недійсними. Це прямо передбачено ч. 2 ст. 215 ЦК України: нікчемний правочин є недійсним у силу закону, а його визнання судом недійсним не вимагається [1].

Правова природа нікчемного правочину полягає в тому, що він вважається недійсним із моменту його вчинення. Тобто сторони формально можуть укласти договір, підписати документи, передати майно чи кошти, однак з погляду права такий правочин не створює тих юридичних наслідків, на які сторони розраховували. Основна відмінність нікчемного правочину від оспорюваного полягає в тому, що оспорюваний правочин вважається чинним доти, доки суд не визнає його недійсним, а нікчемний є недійсним автоматично. Наприклад, якщо договір укладений із порушенням обов'язкової нотаріальної форми у випадках, коли така форма прямо передбачена законом, або якщо правочин спрямований на порушення публічного порядку, він може бути нікчемним незалежно від того, чи зверталася сторона до суду. Саме тому нікчемність є більш суворою формою правової реакції держави на порушення вимог цивільного законодавства [2, с. 401].

Загальні вимоги до чинності правочину визначені ст. 203 ЦК України. Зокрема, зміст правочину не повинен суперечити ЦК України, іншим актам цивільного законодавства, інтересам держави і суспільства, його моральним засадам; особа, яка вчиняє правочин, повинна мати необхідний обсяг цивільної дієздатності; волевиявлення учасника має бути вільним і відповідати його внутрішній волі; правочин має вчинятися у формі, встановленій законом; він має бути спрямований на реальне настання правових наслідків. Недодержання цих вимог у передбачених законом випадках може бути підставою для нікчемності. Отже, нікчемний правочин – це не просто «помилкова» або «невигідна» угода, а правочин, у якому

порушення є настільки істотним, що закон не визнає за ним юридичної сили від самого початку.

Цивільний кодекс України містить низку конкретних прикладів нікчемних правочинів. Наприклад, нікчемним є правочин, що обмежує можливість фізичної особи мати не заборонені законом цивільні права та обов'язки. Також нікчемними можуть бути правочини, які порушують вимоги щодо форми, якщо закон прямо встановлює такий наслідок. Окрему групу становлять правочини, що порушують публічний порядок. Такі правочини є особливо небезпечними, оскільки вони спрямовані не лише проти приватних інтересів сторін, а й проти загальних засад правопорядку. Наприклад, правочин, метою якого є незаконне заволодіння майном, приховування незаконної діяльності або порушення конституційних прав особи, не може отримати правовий захист. У цьому проявляється охоронна функція цивільного права: закон не дозволяє використовувати форму договору для прикриття незаконної або аморальної поведінки [3, с. 84].

Правові наслідки нікчемного правочину регулюються насамперед ст. 216 ЦК України. Основне правило полягає в тому, що недійсний правочин не створює юридичних наслідків, крім тих, що пов'язані з його недійсністю. Це означає, що сторони повинні бути повернуті у становище, яке існувало до вчинення такого правочину. Такий механізм називається реституцією. За загальним правилом застосовується двостороння реституція: кожна сторона зобов'язана повернути другій стороні все одержане за правочином, а якщо повернення в натурі неможливе – відшкодувати вартість одержаного. Наприклад, якщо за нікчемним договором купівлі-продажу одна сторона передала майно, а інша – гроші, то майно має бути повернуте продавцю, а гроші – покупцю. Якщо ж майно вже знищене, спожите або відчужене третім особам, виникає питання про грошове відшкодування його вартості.

Водночас правові наслідки нікчемного правочину не завжди обмежуються лише реституцією. Якщо внаслідок вчинення такого правочину одній зі сторін або третій особі завдано збитків чи моральної шкоди, може поставати питання про їх відшкодування. Крім того, у певних випадках закон може передбачати спеціальні наслідки недійсності. Наприклад, якщо правочин був спрямований на порушення публічного порядку або мав протиправну мету, наслідки можуть бути суворішими, ніж звичайне повернення сторін у попередній стан. Це пояснюється тим, що цивільне право не повинно створювати вигідних умов для осіб, які свідомо діють всупереч закону. Саме тому суд, розглядаючи спір, пов'язаний із нікчемним правочином, оцінює не лише сам факт формального порушення, а й зміст правочину, мету його укладення, поведінку сторін та реальні наслідки для цивільного обороту. Важливо підкреслити, що хоча нікчемний правочин не потребує визнання його недійсним судом, судова процедура все одно може мати значення. На практиці сторони часто звертаються до суду не для того, щоб «визнати» нікчемний правочин недійсним, а для застосування наслідків його недійсності: повернення майна, стягнення коштів, скасування реєстраційних дій, припинення незаконного володіння тощо. Тобто суд у таких справах не створює факт недійсності, а лише підтверджує його та визначає, які саме правові наслідки мають бути застосовані. У цьому полягає практична складність нікчемних правочинів: з одного боку, вони є недійсними автоматично, але з іншого – для реального захисту порушеного права особа часто змушена звертатися до суду. Судова влада України також звертає увагу, що підставою недійсності правочину є недодержання вимог, установлених законом, саме на момент його вчинення [4, с. 131].

Отже, нікчемні правочини є важливим правовим механізмом захисту законності у сфері цивільних відносин. Їх сутність полягає в тому, що закон не визнає юридичної сили за правочинами, які прямо порушують імперативні норми цивільного законодавства, публічний порядок або основні вимоги до чинності правочину. Основними правовими наслідками нікчемності є відсутність передбачених сторонами юридичних результатів, повернення одержаного за правочином, можливість відшкодування збитків та застосування спеціальних наслідків у випадках, визначених законом. Значення цього інституту полягає не лише у

відновленні майнового становища сторін, а й у запобіганні зловживанню правом, прикриттю незаконних дій формою договору та порушенню стабільності цивільного обороту.

Список використаних джерел

1. Цивільний кодекс України: Закон України від 16.01.2003 № 435-IV. URL: <https://zakon.rada.gov.ua/laws/show/435-15>.
2. Вовк М. З., Юркевич Ю. М. Проблеми тлумачення договорів в Україні в умовах гармонізації з правом європейського союзу. Актуальні питання у сучасній науці. 2024. № 2 (20). С. 394-405.
3. Гречишкін Є.Ю., К.Р. Резворович. Актуальні питання теоретичного визначення недійсних правочинів у межах цивільного права. Матеріали Всеукраїнської науково-практичної онлайн-конференції (ДДУВС, 12.12.2023). 2023. С. 84-85.
4. Цивільне право : навч. практикум. кол. авт. : К Р. Резворович (заг. ред.), М. П. Юніна, А. С. Ділігул та Логінова М.В. Дніпро : Дніпроп. держ. ун-т внутр. справ, 2022. 220 с.

Юрій КОРОЛЬ, курсант 2-го курсу ННІ з підготовки фахівців для підрозділів кримінальної поліції, Львівський державний університет внутрішніх справ.

Науковий керівник:

Віталій КУЧЕР, завідувач кафедри загальноправових дисциплін факультету № 1, Львівський державний університет внутрішніх справ, кандидат юридичних наук, доцент

МОНІТОРИНГ ДОТРИМАННЯ ПРАВ І СВОБОД ЛЮДИНИ В ОРГАНАХ ПОЛІЦІЇ

Проблема захисту прав і свобод людини в діяльності правоохоронних органів залишається одним із ключових викликів для демократичної держави. В Україні, яка прагне до євроінтеграції та утвердження верховенства права, питання належного моніторингу дотримання конституційних прав людини у відносинах із поліцією набуває особливої актуальності. Систематичні скарги на зловживання службовим становищем, незаконне затримання та неналежне поводження із затриманими свідчать про нагальну потребу у дієвих механізмах нагляду та контролю за діяльністю Національної поліції України (далі – НПУ).

Ефективний моніторинг є не лише правовим обов'язком держави, але й необхідною умовою відновлення суспільної довіри до поліцейських інститутів. Відповідно до Конституції України, людина, її життя і здоров'я, честь і гідність, недоторканність і безпека визнаються найвищою соціальною цінністю (ст. 3). Держава відповідає перед людиною за свою діяльність, що безпосередньо покладає на неї обов'язок забезпечити контроль за дотриманням прав громадян у правоохоронній сфері. Конституція України встановлює і конкретні гарантії: ст. 29 закріплює право на особисту недоторканність, ст. 55 – право на захист прав і свобод у судовому порядку, а ст. 101 – функцію парламентського контролю за дотриманням прав людини і основоположних свобод, яку здійснює Уповноважений Верховної Ради України з прав людини (Омбудсмен) [1].

Спеціальне законодавство у сфері поліцейської діяльності конкретизує конституційні положення. Стаття 7 Закону України «Про Національну поліцію» закріплює принцип дотримання прав і свобод людини як основоположний у діяльності поліції, де зазначено, що під час виконання своїх завдань поліція забезпечує дотримання прав і свобод людини, гарантованих Конституцією та законами України, а також міжнародними договорами України, згода на обов'язковість яких надана Верховною Радою України, і сприяє їхній реалізації [2]. Цей принцип є нормативним підґрунтям для побудови всієї системи внутрішнього та зовнішнього контролю за роботою поліції.

Закон також зобов'язує поліцейських негайно повідомляти затриманих осіб про підстави їх затримання та роз'яснювати права (ст. 208 КПК України [3]). Кримінальний процесуальний кодекс України деталізує процедурні гарантії: право на захисника, право зберігати мовчання, право на медичну допомогу та інші, що у сукупності формують стандарт поведінки із затриманими, відхилення від якого є підставою для перевірки та притягнення до відповідальності.

У системі моніторингу дотримання поліцейськими прав та свобод людини ключове місце займає Наказ МВС України від 24.05.2022 № 311 «Про затвердження Інструкції з формування та ведення інформаційної підсистеми «Custody Records» інформаційно-комунікаційної системи «Інформаційний портал Національної поліції України», який до завдань ІП «Custody Records» відносить запобігання незаконному затриманню осіб, удосконалення системи їх захисту від катувань і належного поведінки, а також підвищення стандартів захисту прав поліцейських від можливих неправдивих звинувачень у неправомірних діях [4].

Запровадження ІП «Custody Records» виявилось одним з найефективніших інструментів запобігання незаконному поведінку з особами, позбавленими волі: система забезпечує повну, безперервну та надійну реєстрацію всіх дій, що вчиняються щодо особи від фактичного затримання до звільнення або поміщення в ізолятор тимчасового тримання. Вона включає в себе регулятивні, організаційні та технічні елементи і створює єдиний інформаційний простір, що мінімізує ризик маніпулювання даними та сприяє об'єктивному контролю [5, с. 36].

Інституційна система моніторингу дотримання прав та свобод людини в поліцейській діяльності включає як зовнішні, так і внутрішні механізми. Зовнішній моніторинг дотримання прав та свобод людини здійснює Уповноважений Верховної Ради України з прав людини, який відповідно до Закону України «Про Уповноваженого Верховної Ради України з прав людини» має право безперешкодно відвідувати місця несвободи, перевіряти умови тримання та опитувати осіб, які перебувають у таких місцях, вносити подання щодо усунення порушень [6].

Внутрішній механізм моніторингу дотримання прав та свобод людини в поліцейській діяльності реалізується через інспекторів з дотримання прав та свобод людини - працівників підрозділів дотримання прав людини Національної поліції України, діяльність яких регулюється Інструкцією з організації діяльності підрозділів дотримання прав людини Національної поліції України, затвердженою Наказом МВС України від 16.10.2024 № 699. Відповідно до зазначеної Інструкції такий моніторинг полягає у збиранні та аналізі інформації з метою визначення відповідності дій поліцейських вимогам законодавства щодо дотримання прав затриманих, доставлених до органів (підрозділів) поліції та ізоляторів тимчасового тримання.

Працівники підрозділів дотримання прав людини здійснюють систематичний моніторинг через доступ до записів ІП «Custody Records» щодо кожної затриманої особи, а також через віддалений доступ до ІР-камер системи відеоспостереження, що дозволяє здійснювати безперервний контроль у режимі реального часу.

На рівні центрального органу управління поліції діє Департамент головної інспекції та дотримання прав людини Національної поліції України – самостійний структурний підрозділ, що здійснює контроль за службовою діяльністю органів поліції, реалізовує державну політику у сфері дотримання прав людини під час надання поліцейських послуг, а також проводить службові розслідування та перевірки щодо дисциплінарних проступків [8].

Поєднання зовнішнього парламентського контролю з боку Омбудсмена та внутрішньовідомчого моніторингу дотримання прав людини формує дворівневу систему захисту прав людини в поліцейській діяльності, ефективність якої залежить від налагодженої міжінституційної взаємодії та прозорості процедур.

Таким чином, система моніторингу дотримання прав і свобод людини в органах поліції спирається на розгалужену нормативну базу – від конституційних (ст. 3, 29, 55, 101 Конституції України) та законодавчих засад (ст.ст. 6, 7, 8, 18 Закону України «Про Національну поліцію») до підзаконних відомчих нормативно-правових актів. Водночас наукові дослідження та практика свідчать про необхідність посилення координації між зовнішніми та внутрішніми інституціями, що спрямовані на забезпечення прав і свобод людини. Лише комплексний підхід, дієві інституційні механізми та активне залучення громадянського суспільства, здатен забезпечити реальний, а не декларативний захист прав людини в поліцейській діяльності.

Список використаних джерел

1. Конституція України від 28.06.1996 р. URL: <https://zakon.rada.gov.ua/laws/show/254%D0%BA/96-%D0%B2%D1%80>
2. Закон України «Про Національну поліцію» від 02.07.2015 р. № 580-VIII. URL: <https://zakon.rada.gov.ua/laws/show/580-19>
3. Кримінальний процесуальний кодекс України від 13.04.2012 р. № 4651-VI. URL: <https://zakon.rada.gov.ua/laws/show/4651-17>
4. Про затвердження Інструкції з формування та ведення інформаційної підсистеми «Custody Records» інформаційно-комунікаційної системи «Інформаційний портал Національної поліції України»: Наказ МВС України від 24.05.2022 р. № 311. URL: <https://zakon.rada.gov.ua/laws/show/z0629-22#Text>
5. Байбак А.Г. Досвід програми Custody Records в Україні в контексті правозахисної проблематики. *Науковий вісник Ужгородського Національного Університету. Серія «Право»*. 2025.. Випуск 92: частина 3. С. 29-37.
6. Закон України «Про Уповноваженого Верховної Ради з прав людини» від 23.12.1997 р. № 776/97-ВР. URL: <https://zakon.rada.gov.ua/laws/show/776/97-%D0%B2%D1%80#Text>
7. Інструкція з організації діяльності підрозділів дотримання прав людини, затверджена Наказом МВС України від 16.10.2024 р. № 699. URL: <https://zakon.rada.gov.ua/laws/show/z1589-24#Text>

Діана ХИЖА, курсантка 2-го курсу ННІ з підготовки фахівців для підрозділів кримінальної поліції, Львівський державний університет внутрішніх справ

ПРИНЦИП ЗАКОННОСТІ, ЯК ГАРАНТІЯ ЗАХИСТУ ПРАВ І СВОБОД ЛЮДИНИ В ПОЛІЦЕЙСЬКІЙ ДІЯЛЬНОСТІ

У сучасних умовах розвитку демократичної та правової держави особливого значення набуває питання забезпечення прав і свобод людини у діяльності правоохоронних органів. Одним із найважливіших принципів функціонування Національна поліція України є принцип законності, який виступає основою професійної діяльності поліцейських та гарантією недопущення порушення прав громадян.

Законність – важливий принцип діяльності поліції. Він виражається як в правоохоронній діяльності поліції, так і в тому, що ця діяльність здійснюється на основі чіткого та неухильного дотримання законів і підзаконних нормативних актів. Національна поліція України, охороняючи громадський порядок, захищаючи в межах своїх повноважень життя, здоров'я та свободи громадян, власність від протиправних посягань, тим самим послідовно забезпечує законність в Україні [1, с. 209].

Відповідно до Конституції України органи державної влади та їх посадові особи повинні діяти лише на підставі, у межах повноважень та у спосіб, що передбачені законом [2]. Саме тому діяльність поліції не може виходити за межі законодавства, навіть у випадках необхідності швидкого реагування чи застосування примусових заходів.

Для успішної реалізації принципу законності в діяльності поліції, необхідно, щоб поліцейський чітко усвідомлював, що ні за яких умов не може виконувати злочинні накази вищих посадових осіб. А у разі виконання таких, буде нести відповідальність відповідно до закону [3, с. 71].

Особливість поліцейської діяльності полягає в тому, що поліція має право застосовувати заходи державного примусу, обмежувати окремі права осіб, проводити перевірки, затримання та використовувати спеціальні засоби. Саме тому надзвичайно важливо, щоб усі дії поліцейських були законними, обґрунтованими та пропорційними ситуації. Дотримання принципу законності дозволяє забезпечити баланс між підтриманням публічного порядку та повагою до прав людини.

Принцип законності є фундаментальним принципом, якому підпорядковуються всі сфери діяльності. Сутність принципу визначена у ст. 8 Закону України «Про Національну поліцію», але слід зазначити, що для його ефективної реалізації поліцейський повинен усвідомлювати, що ні за яких обставин він не може порушувати норми закону.

Без дотримання законності діяльність правоохоронних органів буде незаконною, недійсною, її результати не будуть визнаватися ні суспільством, ні державою [4, с. 21].

Оцінюючи діяльність поліції України, громадяни самі можуть визначити, наскільки держава забезпечує їхні законні права та свободи, адже чим менше буде зафіксовано порушень чинного законодавства під час здійснення поліцейськими своїх функцій, тим ефективнішою буде сама поліцейська діяльність [5, с. 29-30].

Також, важливу роль у забезпеченні законності відіграє професійна підготовка працівників поліції. Поліцейські повинні добре знати законодавство, розуміти межі своїх повноважень та вміти діяти відповідно до принципу верховенства права навіть у складних та конфліктних ситуаціях. Крім того, важливе значення має контроль за діяльністю поліції з боку держави, суду та суспільства.

Варто зазначити, що міжнародні стандарти у сфері прав людини також вимагають від правоохоронних органів дотримання принципу законності. Практика Європейський суд з прав людини неодноразово підкреслювала, що будь-яке втручання держави у права людини повинно бути законним, необхідним та обґрунтованим.

На мою думку, принцип законності є не лише юридичною вимогою, а й моральною основою діяльності поліцейського. Працівник поліції повинен усвідомлювати відповідальність за свої дії, адже будь-яке перевищення службових повноважень негативно впливає на авторитет правоохоронної системи та рівень довіри населення до держави.

Особливість поліцейської діяльності обумовлює підвищені вимоги до рівня правової культури, професійної етики та відповідальності працівників поліції. Надання поліцейським владних повноважень і права застосування заходів державного примусу потребує чіткого нормативного регулювання та ефективного контролю за законністю їх реалізації. У цьому аспекті принцип законності виступає не лише правовою категорією, а й важливим механізмом гарантування конституційних прав особи.

Слід зазначити, що реальне утвердження принципу законності у діяльності поліції є можливим лише за умови поєднання належного законодавчого забезпечення, високого рівня професійної підготовки поліцейських, ефективного громадського контролю та

невідворотності юридичної відповідальності за порушення прав людини. Саме такий підхід сприятиме зміцненню авторитету правоохоронної системи, підвищенню рівня правосвідомості суспільства та подальшому розвитку України як демократичної, соціальної і правової держави.

Отже, принцип законності є однією з фундаментальних засад діяльності Національна поліція України та необхідною умовою ефективного забезпечення прав і свобод людини у демократичній державі. Його зміст полягає не лише у формальному дотриманні норм законодавства, а й у практичній реалізації ідей верховенства права, справедливості, правової визначеності та поваги до людської гідності.

Саме через дотримання законності забезпечується недопущення свавілля з боку правоохоронних органів, обмеження незаконного втручання у права громадян та формування належного рівня суспільної довіри до поліції.

Список використаних джерел

1. Бакутін Є.І. Принцип законності в діяльності поліції під час використання технічних засобів фіксації адміністративних правопорушень. *Часопис Київського університету права*. 2020. № 2. С. 208-213. URL: <https://chasprava.com.ua/index.php/journal/article/download/374/355/>
2. Конституція України від 28.06.1996 р. URL: <https://zakon.rada.gov.ua/laws/show/z0629-22#Text>
3. Нестерцова-Собакарь О. В. Реалізація принципу законності в діяльності національної поліції України в умовах воєнного стану. *Актуальні проблеми превентивної діяльності Національної поліції в умовах воєнного стану*: Матеріали Всеукраїнського науково-практичного семінару (ДДУВС, 27.04.2022). С. 70-71.
4. Ніколенко Л. М. Принципи діяльності поліції: Сутність та особливості. *Українська поліцейстика: теорія, законодавство, практика*. 2021. № 1. С. 19-25.
5. Кобзар О. Ф. Поліцейська діяльність в Україні: адміністративно-правовий аспект : монографія. Харків; Дніпропетровськ: Панов, 2015. 316 с.

Наукове видання

Укладачі:

Корнейко Олександр Васильович

Швець Микола Миколайович

**ПРОБЛЕМИ ТА ПЕРСПЕКТИВИ ВИКОРИСТАННЯ
МОЖЛИВОСТЕЙ СИСТЕМ ШТУЧНОГО
ІНТЕЛЕКТУ В ПРАВІ, ПСИХОЛОГІЇ
ТА ПРАВООХОРОННІЙ ДІЯЛЬНОСТІ**

МАТЕРІАЛИ

**Міжнародної науково-практичної конференції
молодих вчених та здобувачів вищої освіти
(м. Київ, 27 травня 2026 року)**

Відповідальний за випуск та комп'ютерна верстка: О. В. Корнейко

Підписано до друку 21.07.2026. Формат 60х84/16. Папір офсетний.

Обл.-вид. арк. 7,5. Ум. друк. арк. 6,97.

Тираж 50 прим.

Національна академія внутрішніх справ.

03035, м. Київ, пл. Солом'янська, 1.

Свідоцтво про внесення суб'єкта видавничої справи до державного реєстру видавців,
виготовників і розповсюджувачів видавничої продукції ДК № 4155 від 13.09.2011.

Друк: ФОП Поліщук О.В.

Свідоцтво суб'єкта видавничої справи ДК № 2142 від 31.03.2015

07400, м. Бровари, вул. Незалежності, 2, кв. 148

тел. (044) 592-13-49