

ЮРИДИЧНА ПСИХОЛОГІЯ; ПОЛІТИЧНА ПСИХОЛОГІЯ

УДК 159.-047.22

DOI <https://doi.org/10.32782/2709-3093/2025.6/16>

Ташматов В.А.

Національний університет «Одеська юридична академія»

Геращенко О.С.

Національний університет «Одеська юридична академія»

Голопотелюк А.О.

Національний університет «Одеська юридична академія»

ШТУЧНИЙ ІНТЕЛЕКТ У СУДОВО-ПСИХОЛОГІЧНІЙ ЕКСПЕРТИЗІ: МОЖЛИВОСТІ, РИЗИКИ ТА ЕТИЧНІ РАМКИ

Стаття присвячена всебічному аналізу можливостей, ризиків та етичних рамок використання технологій штучного інтелекту у сфері судово-психологічної експертизи. У контексті стрімкого розвитку алгоритмів машинного навчання та збільшення обсягів даних, доступних для експертного аналізу, ШІ дедалі частіше розглядається як інструмент, здатний підвищити точність прогнозів, стандартизувати процедури оцінювання та мінімізувати суб'єктивний фактор у діяльності судових експертів – психологів. У статті окреслено ключові напрями застосування ШІ, зокрема в оцінці ризику рецидиву, аналізі поведінкових патернів, виявленні психологічних аномалій, а також у побудові комплексних профілів особи на основі великих масивів кримінологічних та психометричних даних. Розглядаються приклади алгоритмічних систем, які вже використовуються в різних країнах світу для підтримки прийняття рішень у кримінальному процесі.

Разом з тим, в статті наголошується, що впровадження ШІ у таку чутливу сферу, як судово-психологічна експертиза, супроводжується низкою суттєвих ризиків. До них належать алгоритмічні упередження, похибки інтерпретації, залежність результатів від якості та репрезентативності даних, а також потенційна втрата індивідуального підходу до оцінюваної особи. Важливою проблемою є й те, що машинні моделі нерідко відтворюють і підсилюють соціальні диспропорції, помилково інтерпретують культурні чи поведінкові особливості та можуть формувати надмірну довіру до автоматизованих висновків. Окрему увагу приділено необхідності розроблення чітких етичних стандартів, які б регулювали використання ШІ – прозорість алгоритмів, відповідальність за наслідки їх застосування, дотримання прав людини та захист конфіденційності.

Стаття підкреслює, що ефективне інтегрування ШІ у судово-психологічну практику можливе лише за умови міждисциплінарної взаємодії психологів, юристів, фахівців з використання даних. Такий підхід дозволить використовувати технології як інструмент підтримки, а не заміни експерта, забезпечуючи баланс між інноваційністю й відповідальністю. Зроблено висновок, що ШІ потенційно здатен підвищити об'єктивність та точність експертної діяльності, однак потребує ретельного контролю, нормативного регулювання й етичної обережності при застосуванні в системі кримінального правосуддя.

Ключові слова: штучний інтелект, судово-психологічна оцінка, кримінальне правосуддя, ризики рецидиву, алгоритмічні упередження, етика ШІ, експертна діяльність.

Постановка проблеми. Стрімкий розвиток технологій штучного інтелекту (ШІ) зумовив суттєві зміни в підходах до оцінювання поведінки, ризиків та психічних характеристик особи у кримінальному правосудді. Актуальність дослідження визначається зростанням обсягів даних, потребою підвищення об'єктивності експертних висновків та ризиками, що виникають у зв'язку з делегуванням машинним системам функцій, які традиційно здійснювали фахівці-психологи. У світовому науковому дискурсі ШІ розглядається як інструмент підвищення точності прогнозів рецидиву, оцінки небезпечності, виявлення поведінкових патернів і аналізу лінгвістичних чи невербальних ознак у матеріалах провадження [6, 20].

Аналіз останніх досліджень і публікацій. Поняття «штучний інтелект» уперше запропонував Дж. Маккарті у 1950-х роках як науку про створення інтелектуальних машин. Через міждисциплінарність феномену єдиного визначення не існує. Європейська етична хартія (2018) трактує ШІ як сукупність методів, спрямованих на відтворення когнітивних здібностей людини. За ДСТУ 2938-94 ШІ – це здатність систем виконувати функції, що асоціюються з людським інтелектом: логічне мислення, навчання, самовдосконалення. У працях [1, с. 65]. О. Баранова та Є. Мічурина підкреслюється, що ШІ є результатом творчої діяльності людини та реалізується через апаратні або програмні системи, здатні до навчання, прогнозування, роботи з неповними даними й оцінки ризиків [19, 3].

На сучасному етапі ШІ активно застосовується в комунікаціях, медицині, транспорті, обороні та правозастосуванні. У військовій сфері він використовується для аналізу великих масивів даних і прогнозування загроз. У практиці правоохоронних органів – для ідентифікації транспортних засобів, фіксації порушень, аналітики відеоданих, прогнозування криміногенної ситуації. ШІ також застосовується для розшуку депортованих дітей, демонструючи потенціал алгоритмів у складних гуманітарних завданнях [1, с.66].

Світова література засвідчує поступове входження ШІ в кримінальну психіатрію та судову психологію. Огляди в *Forensic Psychiatry and Psychology* (2021–2024) підкреслюють його можливості у скринінгу психічних розладів, аналізі мовлення, автоматизованій оцінці ризику насильства та рецидиву. Системи машинного навчання демонструють точність прогнозування, співставну з експертними оцінками, але породжують ризики алгоритмічної упередже-

ності, непрозорості рішень і порушення прав людини [25].

Юридичні дослідження фіксують нормативну невизначеність та відсутність процедурного контролю впровадження ШІ у кримінальному провадженні. Законодавство не встигає за технологічними змінами, що створює ризики неправомірного використання автоматизованих оцінок судом чи слідством. Це зумовлює потребу системного аналізу можливостей, обмежень і ризиків таких технологій, а також механізмів їх відповідального застосування.

Постановка завдання. Формулювання цілей статті є систематичний аналіз можливостей, ризиків та етичних рамок використання ШІ у судово-психологічній оцінці. Завдання включають огляд моделей ШІ, що застосовуються для оцінки ризику та діагностики, визначення ключових правових і етичних проблем і формулювання рекомендацій щодо відповідального впровадження цих технологій у судово-психологічну практику.

Виклад основного матеріалу. Штучний інтелект у судово-психологічній оцінці виступає як багатовимірний інструмент, що відкриває нові методологічні можливості для кількісної і якісної інтерпретації поведінкових, когнітивних та нейрофізіологічних даних, водночас породжуючи комплексні ризики та вимоги до правового й етичного регулювання. Огляд сучасної літератури демонструє, що найбільш зрілі й досліджені застосування стосуються трьох функціональних областей: прогностики ризиків (зокрема рецидиву, агресії, ризику вчинення насильницьких правопорушень), діагностики психічних розладів на основі аналізу природної мови та нейровізуалізаційних даних, а також детекції симуляції або фальсифікації психопатологічних симптомів. У сфері прогнозування поведінкових ризиків застосування дискримінативних моделей машинного навчання (моделі, що відокремлюють класи, наприклад «рецидивіст/не рецидивіст») показало обнадійливі результати щодо підвищення точності прогнозів у порівнянні з традиційними шкалами; водночас дослідники застерігають про обмежену переносимість моделей між юрисдикціями через різницю вибірок, системи правопорядку та соціокультурних факторів, що робить необхідним локальне калібрування й незалежну валідацію алгоритмів [25].

Підхід до діагностики психічних розладів із використанням методів обробки природної мови та аналізу мовних і мовленнєвих патернів, а також штучних нейронних мереж, що працюють з нейро-

візуалізаційними (нейроіміджинг) даними, дозволяє виявляти ранні маркери депресії, посттравматичних і тривожних розладів або дисоціативних станів із точністю, яка в ряді досліджень конкурує з клінічною оцінкою [25]. Наприклад, моделі, які аналізують мовлення пацієнта, міміку та аудіовимірювання, показали дозвіл діагностувати симптоми депресії з високою чутливістю й специфічністю. У дослідженні, яке використовувало 3D-зображення виразу обличчя та мовні дані, алгоритм досягав чутливості 83,3 % та специфічності 82,6 % щодо виявлення великого депресивного розладу (Major Depressive Disorder) [18, 13]. Крім того, глибинні нейромережі, включаючи рекурентні моделі, аналізують акустичні ознаки (наприклад, мел-кепстральні коефіцієнти), щоб прогнозувати як діагноз, так і тяжкість депресії, даючи змогу неінвазивно підтримувати клінічні рішення [23].

Щодо виявлення симуляції (симуляція – демонстрація або перебільшення симптомів з метою отримання вигоди), останні дослідження і практичні розробки застосовують машинне навчання для підвищення чутливості тестів валідності. Італійські дослідники та інші групи показали, що комбінація часових ознак відповіді (реакційний час), патернів відповідей у стандартизованих інструментах (наприклад, MMPI – MMPII) і поведенкових показників дозволяє будувати коректні класифікатори для виявлення фальсифікацій із високою точністю у експериментальних умовах. Праці Мацця і співавторів показали, що при навантаженні часовим тиском (обмеження часу) алгоритми машинного навчання досягали високих показників розпізнавання симуляції (приблизно 90–95% у контрольованих випробуваннях), тоді як точність знижувалася при природних, менш контрольованих умовах, що підкреслює необхідність контекстуалізації таких методів у судово-психологічній практиці [18]. Систематичні огляди показують, що останнім часом методи машинного навчання (підтримувані векторні машини, лінійні дискримінанти, глибокі нейронні мережі) значно підвищують точність детекції брехні порівняно з традиційною інтерпретацією ЕЕГ-ERP, досягаючи рівня >80% класифікації у різних експериментальних умовах [25].

Існують і конкретні інноваційні рішення, наприклад, поєднання тип-2 нечітких множин із глибокими графовими згортковими мережами для аналізу просторових та часових характеристик ЕЕГ-сигналів, що дає можливість класифікувати брехню та правду з високою точністю

(близько 86%) [22]. Водночас критичні дослідники попереджають, що поєднання поліграфа (детектор брехні) з методами машинного навчання може бути небезпечним, якщо йдеться про застосування в кримінальному праві [17].

В аналізі українських правознавців виділяються питання, що стосуються алгоритмічності упередженості (bias), недостатньої прозорості (проблема «чорної скриньки»), ненадійності навчальних даних, а також етичних і психологічних наслідків для усіх учасників системи [3, с. 316]. Однією з ключових проблем є алгоритмічна упередженість, коли моделі, натреновані на історичних даних, відтворюють або навіть підсилюють існуючі соціальні нерівності. У Сполучених Штатах широке вжиття здобула система COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) – комерційна алгоритмічна платформа, що генерує скорингові індекси ризику рецидиву й небезпеки на основі багатьох змінних. Оцінки COMPAS привернули значну увагу через дослідження журналістів і вчених, які виявили: при загальній прогнозній точності, близькій до 60–65%, алгоритм у деяких вибірках систематично упереджувався щодо чорношкірих підозрюваних, тобто алгоритм систематично класифікує афроамериканських підсудних у групу високого ризику рецидиву з вищою ймовірністю неправильної класифікації (false positive) у порівнянні з білими підсудними, що спричинило запити щодо доступу до вихідних моделей і даних для незалежної верифікації [5].

Також проблемою є прозорість алгоритмів, або її відсутність – моделі часто функціонують як «чорні скриньки», коли механізми прийняття рішення для зовнішніх зацікавлених осіб (експертів, підсудних, адвокатів) непрозорі. Ця непрозорість підриває справедливість і довіру, оскільки суб'єкти, на яких впливають рішення ШІ, не можуть детально перевірити, чому їм присвоєно певний рейтинг ризику. Наприклад, у справі Loomis проти Вісконсину (Wisconsin v. Loomis) було зазначено, що підсудний не отримав повної інформації про внутрішню структуру алгоритму COMPAS, оскільки вона є комерційною та закритою для стороннього аудиту. У наукових колах автори попереджають, що без пояснення роботи моделі (наприклад, через так званий explainable AI або пояснюваний ШІ) складно забезпечити підзвітність та справедливість рішень [25].

Наукові групи університету Оксфорда розробили відкритий інструмент OxRec (Oxford Risk of Recidivism). На відміну від закритих комер-

ційних рішень, ОхРес позиціонується як прозорий інструмент із відомим набором предикторів і публічною методологією для прогнозу насильницького рецидиву, що дозволяє незалежну перевірку і адаптацію до локальних умов [11].

У контексті застосування ІІІ в судово-психологічній оцінці критичною є розробка чітких етичних рамок, які б забезпечували баланс між інноваціями та захистом фундаментальних прав людини. Основні етичні принципи, які частіше за все виділяють у наукових та політичних дискусіях, – це справедливість, прозорість, підзвітність і відповідальність [16].

Правовий і етичний досвід міжнародної практики демонструє, що багато країн і міжнародних організацій намагаються вбудувати ці принципи в регуляторні моделі. Так, згідно з аналізом Белова, Белової та співавторів, у міжнародній практиці вже сформовані різні підходи до регулювання ІІІ в досудових розслідуваннях, які включають етично-правові засади – прозорість алгоритмів, захист приватності, забезпечення справедливого судового процесу та недискримінації [3, с. 317].

У порівняльному дослідженні Фагхірі Каболя підкреслено, що хоча ІІІ може ефективно прогнозувати злочинну поведінку та підтримувати рішення слідчих або суддів, у багатьох країнах – зокрема в Об'єднаних Арабських Еміратах – відсутня легітимна нормативно-правова база для його повноцінного застосування [15].

Однією з найважливіших сучасних регуляторних ініціатив є Рамкова конвенція Ради Європи з питань ІІІ, прав людини, демократії і верховенства права (CETS No. 225), яка стала першим міжнародним юридично обов'язковим договором у цій сфері [9]. Ця конвенція покриває весь життєвий цикл систем штучного інтелекту, накладаючи обов'язки на держави. Положення конвенції включають обов'язок документувати інформацію про системи ІІІ та надавати достатні дані про рішення, щоб постраждалі могли оскаржити рішення або навіть використання системи; це є суттєвим механізмом підзвітності [8].

Разом із цим, конвенція запроваджує рамку для оцінки ризиків (risk-management). У цьому контексті розроблено неформальні методології, як-от HUDERIA (Human Rights, Democracy and the Rule of Law Impact Assessment), яка слугує інструментом для держав і розробників систем ІІІ, щоб систематично оцінювати вплив на права людини, верховенство права та демократичні принципи [8].

На регіональному рівні Європейський Союз також просуває законодавчі ініціативи. Напри-

клад, Акті про штучний інтелект (AI Act), який уперше на міжнародному рівні створює комплексну класифікацію систем ІІІ за рівнем суспільного ризику [11; 21]. У межах цієї класифікації системи, що застосовуються правоохоронними органами та в судовій практиці, зараховані до категорії високого ризику. До них належать інструменти індивідуальної оцінки ризику повторного правопорушення, аналізу емоційного стану особи, оцінювання надійності доказів, а також системи, що здійснюють профілювання особистості на підставі рис характеру чи поведінкових патернів.

Відповідно до регламенту, використання таких систем є можливим лише за умов дотримання низки суворих вимог: забезпечення обов'язкового контролю з боку людини, наявності адекватних заходів безпеки, ведення детальної документації, фіксація дій моделі, а також дотримання принципу прозорості при розгортанні системи в реальному середовищі.

У Акті є заборона так званих «неприйнятних» або «надвисокоризикових» практик штучного інтелекту. AI Act прямо забороняє соціальний скринінг, прогнозування індивідуальних злочинів на основі профілювання, маніпулювання вразливими групами тощо. Порушення регламенту передбачає істотні санкції – від багатомільйонних штрафів до накладення відсоткових стягнень від глобального обороту компаній-розробників або впроваджувачів таких систем [7; 21].

Однак навіть такі передові нормативні рамки зіштовхуються з регуляторними викликами. Ще до введення актів, дослідники вказували на труднощі з поясненням алгоритмів: зокрема, використання post-hoc пояснень може виявитися недостатнім або маніпульованим в адвесаріальному контексті (коли інтереси сторін конфліктують) [6].

Інший виклик – незалежний аудит. Такі організації мають включати фахівців з психології, права, етики, представників громадськості, щоб оцінити не лише технічні, але й соціальні ризики [8].

Етичні рамки мають супроводжуватися освітніми програмами, щоб фахівці розуміли, як працюють алгоритми, які можуть бути їхні обмеження, упередження, джерела помилок. У науковому дискурсі також виокремлюють необхідність мультидисциплінарного комітету, який би постійно переглядав стандарти і надавав практичні рекомендації. У деяких країнах такі консультативні групи вже створено: наприклад, у Канаді існує експертна група зі ІІІ для правоохоронних органів, яка розробляє етичні кодекси з акцентом на справедливість, прозорість та захист прав гро-

мадян. У Сінгапурі також діє етичний кодекс для правоохоронних органів, який вимагає регулярної оцінки ризиків, відкритого діалогу з громадськістю, а також гарантує право на приватність при зборі біометричних даних [3, с.318–319].

Висновки. Проведений аналіз засвідчує, що використання штучного інтелекту у судово-психологічній оцінці є перспективним, але водночас вимагає надзвичайно обережного та регульованого підходу. ШІ демонструє значний потенціал у прогнозуванні ризиків рецидиву, діагностиці психічних розладів, аналізі поведінкових та нейрофізіологічних даних, а також у виявленні симуляції симптомів. У низці досліджень моделі машинного навчання досягають точності, співставної або навіть вищої за традиційні клінічні методи.

Разом із тим, результати роботи підкреслюють, що впровадження ШІ супроводжується суттєвими ризиками: алгоритмічною упередженістю, непрозорістю моделей, ненадійністю або неповнотою навчальних даних, потенційними порушеннями прав людини та етичними дилемами. Приклади з міжнародної практики – зокрема дискусії навколо COMPAS – демонструють, що навіть відносно точні моделі можуть бути соціально несправедливими та непрозорими для учасників процесу.

Юридичний аналіз підтвердив, що у більшості країн, включно з Україною, нормативно-правова база не встигає за розвитком технологій. У від-

повідь на ці виклики міжнародні інституції розробляють регуляторні моделі – такі як Рамкова конвенція Ради Європи щодо ШІ або Акт ЄС про штучний інтелект – що спрямовані на забезпечення прозорості, підзвітності, незалежного аудиту, захисту прав осіб та етичного використання алгоритмічних систем у правозастосуванні.

На підставі проведеного дослідження можна зробити висновок, що застосування ШІ у судово-психологічній оцінці має бути допоміжним, а не визначальним інструментом. Його результати повинні інтерпретуватися лише в поєднанні з професійною експертною оцінкою, із забезпеченням права на оскарження та доступом до інформації про методологію моделей. Ключовим завданням стає розробка локально валідованих алгоритмів, створення прозорих стандартів документування, забезпечення етичного нагляду й формування міждисциплінарних команд, здатних оцінювати як технічні, так і соціально-правові наслідки впровадження ШІ.

Таким чином, інтеграція штучного інтелекту в судово-психологічну практику може сприяти підвищенню точності та ефективності оцінок, однак лише за умов дотримання принципів справедливості, прозорості, підзвітності. Подальший розвиток цієї сфери потребує балансу між інноваціями та захистом фундаментальних прав людини, а також безперервного оновлення регуляторних рамок у відповідь на технологічні зміни.

Список літератури:

1. Авдієва Г. К. Проблеми використання систем штучного інтелекту у правозастосовній діяльності. *Вісник Луганського навчально-наукового інституту імені Е.О. Дідоренка*. (2). 63–80 с. <https://doi.org/10.33766/2524-0323.102.63-80>
2. Баранов О. А. Інтернет речей: регулювання надання послуг роботами зі штучним інтелектом. *Інформація і право*. 4(27). 46–70 с. URL: http://ippi.org.ua/sites/default/files/7_10.pdf
3. Белов Д. М., Белова М. В. Штучний інтелект у судочинстві та судових рішеннях, потенціал та ризики. *Науковий вісник Ужгородського Національного Університету*. 78(2). 315–320 с. <https://doi.org/10.24144/2307-3322.2023.78.2.50>
4. Мічурін Є. Правова природа штучного інтелекту. *Форум Права*. 64(5). 67–75 с.
5. Abd-alrazaq A., Alhuwail D., Schneider J., et al. The performance of artificial intelligence-driven technologies in diagnosing mental disorders: An umbrella review. *npj Digital Medicine*. 5. 87 с. <https://doi.org/10.1038/s41746-022-00631-8>
6. Blott H., Hind E., Brown C., Forrester A. Artificial intelligence in forensic psychiatry: Potential applications and key considerations. *Journal of Forensic and Legal Medicine*, 116 p., 103016. <https://doi.org/10.1016/j.jflm.2025.103016>
7. Bordt S., Finck M., Raidl E., von Luxburg U. Post-hoc explanations fail to achieve their purpose in adversarial contexts. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency (FAccT '22)* (pp. 891–905). Association for Computing Machinery. <https://doi.org/10.1145/3531146.3533153>
8. Council of Europe. Council of Europe adopts first international treaty on artificial intelligence. URL: https://www.coe.int/en/web/portal/full-news/-/asset_publisher/y5xQt7QdunzT/content/id/267650696?_com_liferay_asset_publisher_web_portlet_AssetPublisherPortlet_INSTANCE_y5xQt7QdunzT

9. Council of Europe. Explanatory report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. URL: https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence?utm_source=chatgpt.com
10. Fazel S., Chang Z., Långström N., Fanshawe T., Mallett S. OxRec model for assessing risk of recidivism: Ethics – Authors’ reply. *The Lancet Psychiatry*. 3(9). 809–810 p. [https://doi.org/10.1016/S2215-0366\(16\)30216-4](https://doi.org/10.1016/S2215-0366(16)30216-4)
11. Foo Y. C., Hummel T. Europe sets benchmark for rest of the world with landmark AI laws. *Reuters*. URL: https://www.reuters.com/world/europe/eu-countries-back-landmark-artificial-intelligence-rules-2024-05-21/?utm_source=chatgpt.com
12. Haque A., Guo M., Miner A. S., Fei-Fei L. Measuring depression symptom severity from spoken language and 3D facial expressions. *arXiv preprint arXiv:1811.08592*. <https://doi.org/10.48550/arXiv.1811.08592>
13. How We Analyzed the COMPAS Recidivism Algorithm by Jeff Larson, Surya Mattu, Lauren Kirchner and Julia Angwin. *ProPublica*. URL: https://www.propublica.org/article/how-we-analyzed-the-compas-recidivism-algorithm?utm_source=chatgpt.com
14. Kabol A. The Use Of Artificial Intelligence In The Criminal Justice System (A Comparative Study). *Webology*. 2022. 19. 593 p.
15. Khan A. A., Badshah S., Liang P., Waseem M., Khan B., Ahmad A., ... Akbar M. A. Ethics of AI: A systematic literature review of principles and challenges. In *Proceedings of the 26th International Conference on Evaluation and Assessment in Software Engineering* (pp. 383–392). <https://doi.org/10.48550/arXiv.2109.07906>
16. Kotsoglou K. N., Biedermann A. Polygraph-based deception detection and Machine Learning. Combining the Worst of Both Worlds? *Forensic science international. Synergy*. 9. 100479. <https://doi.org/10.1016/j.fsisy.2024.100479>
17. Malgaroli M., Hull T. D., Zech J. M., Althoff T. Natural language processing for mental health interventions: A systematic review and research framework. *Translational Psychiatry*. 13(1). 309 p. <https://doi.org/10.1038/s41398-023-02592-2>
18. Mazza C., Monaro M., Orrù G., Burla F., Colasanti M., Ferracuti S., Roma P. Introducing machine learning to detect personality faking-good in a male sample. *Frontiers in Psychiatry*. 10. 389. <https://doi.org/10.3389/fpsy.2019.00389>
19. Olawade D. B., Ayoola F. I., Ebo T. O., Asaolu A. J., Egbon E., Clement David-Olawade A. Artificial intelligence in forensic mental health: A review of applications and implications. *Journal of Forensic and Legal Medicine*. 113. 102895. <https://doi.org/10.1016/j.jflm.2025.102895>
20. Piquard A. First measures of European AI Act regulation take effect. *Le Monde*. URL: https://www.lemonde.fr/en/pixels/article/2025/02/02/artificial-intelligence-the-first-measures-of-the-european-ai-act-regulation-take-effect_6737691_13.html?utm_source=chatgpt.com
21. Rahmani M., Mohajelin F., Khaleghi N., Sheykhivand S., Danishvar S. (An Automatic Lie Detection Model Using EEG Signals Based on the Combination of Type 2 Fuzzy Sets and Deep Graph Convolutional Networks. *Sensors*. 24(11). 3598. <https://doi.org/10.3390/s24113598>
22. Rejaibi E., Komaty A., Meriaudeau F., Agrebi S., Othmani A. MFCC-based recurrent neural network for automatic clinical depression recognition and assessment from speech. *Biomedical Signal Processing and Control*. 71. 103107. <https://doi.org/10.48550/arXiv.1909.07208>
23. Taha B. N., Baykara M., Alakus T. B. Neurophysiological Approaches to Lie Detection: A Systematic Review. *Brain sciences*. 15(5). 519 p. <https://doi.org/10.3390/brainsci15050519>
24. Tortora L. Beyond discrimination: Generative AI applications and ethical challenges in forensic psychiatry. *Frontiers in Psychiatry*. 15. 1346059. <https://doi.org/10.3389/fpsy.2024.1346059>
25. Wang C., Han B., Patel B., Rudin C. In pursuit of interpretable, fair and accurate machine learning for criminal recidivism prediction. *Journal of Quantitative Criminology*. 39(2). 519–581 p. <https://doi.org/10.48550/arXiv.2005.04176>

Tashmatov V.A., Heraschenko O.S., Holopoteluk A.O. ARTIFICIAL INTELLIGENCE IN FORENSIC PSYCHOLOGICAL EXAMINATION: OPPORTUNITIES, RISKS AND ETHICAL FRAMEWORKS

The article is devoted to a comprehensive analysis of opportunities, risks and ethical frameworks for the use of artificial intelligence technologies in the field of forensic psychological examination. In the context of the rapid development of machine learning algorithms and the increase in the amount of data available for expert analysis, AI is increasingly viewed as a tool capable of increasing the accuracy of forecasts, standardizing assessment procedures, and minimizing the subjective factor in the work of forensic psychologists. The article outlines the key areas of application of AI, in particular in assessing the risk of recidivism, analyzing behavioral patterns, identifying psychological anomalies, as well as in building complex profiles of a person based on

large arrays of criminological and psychometric data. Examples of algorithmic systems that are already used in different countries of the world to support decision-making in the criminal process are considered.

At the same time, the article emphasizes that the implementation of AI in such a sensitive field as forensic psychological examination is accompanied by a number of significant risks. These include algorithmic biases, interpretation errors, the dependence of results on the quality and representativeness of data, as well as the potential loss of an individual approach to the evaluated person. An important problem is that machine models often reproduce and reinforce social disparities, misinterpret cultural or behavioral characteristics, and can form excessive confidence in automated conclusions. Particular attention is paid to the need to develop clear ethical standards that would regulate the use of AI – transparency of algorithms, responsibility for the consequences of their use, respect for human rights and protection of privacy.

The article emphasizes that the effective integration of AI into forensic psychological practice is possible only under the condition of interdisciplinary interaction of psychologists, lawyers, and specialists in the use of data. This approach will allow technology to be used as a support tool rather than a substitute for an expert, ensuring a balance between innovation and responsibility. It is concluded that AI is potentially able to increase the objectivity and accuracy of expert work, but requires careful control, regulatory regulation and ethical caution when applied in the criminal justice system.

Key words: *artificial intelligence, forensic psychological assessment, criminal justice, recidivism risk, algorithmic bias, AI ethics, expert evaluation.*

Дата надходження статті: 05.11.2025

Дата прийняття статті: 27.11.2025

Опубліковано: 30.12.2025