



Державна наукова установа «Інститут інформації, безпеки і права  
Національної академії правових наук України»

Науково-дослідний інститут інтелектуальної власності  
Національної академії правових наук України

## **УКРАЇНА В УМОВАХ СОЦІАЛЬНОЇ ТА ЦИФРОВОЇ ТРАНСФОРМАЦІЇ: ТЕОРЕТИЧНІ МОДЕЛІ ПРАВОВОГО РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ**

МАТЕРІАЛИ НАУКОВО-ПРАКТИЧНОЇ КОНФЕРЕНЦІЇ

*Київ, 5 листопада 2025 року*

Київ-Одеса

2025

**Рекомендовано до друку**

*Вченою радою Державної наукової установи «Інститут інформації, безпеки і права Національної академії правових наук України».*

*Протокол № 10 від 09.12.2025 р.*

- У 45 **Україна в умовах соціальної та цифрової трансформації: теоретичні моделі правового регулювання штучного інтелекту** : матеріали наук.-практ. конф. (Київ, 5 листоп. 2025 р.) [Електронне видання] / упоряд.: М. В. Дубняк, С. О. Дорогих, І. Ф. Корж, В. М. Фурашев. Київ; Одеса : Фенікс, 2025. 240 с. Режим доступу: [https://ippi.org.ua/sites/default/files/zbirnik\\_tez\\_05.11.2025\\_1.pdf](https://ippi.org.ua/sites/default/files/zbirnik_tez_05.11.2025_1.pdf)

ISBN 978-617-8430-97-9

Збірник містить матеріали щодо стратегічних напрямів правового регулювання штучного інтелекту в Україні; розвитку національної LLM та впровадження агентних ШІ-рішень у державні сервіси; формування цифрових прав і оновлення цивільного законодавства; проблем використання ШІ в умовах воєнного стану, забезпечення кібербезпеки, захисту персональних даних та інтелектуальної власності; етичних і методологічних засад застосування ШІ у сфері освіти та науки.

Доповіді учасників конференції можуть бути корисними для фахівців, експертів і вчених, науково-педагогічних працівників та здобувачів вищої освіти.

**УДК 340.132:004.8(477)**

*Матеріали подано у авторській редакції.*

ISBN 978-617-8430-97-9

© Державна наукова установа «Інститут інформації, безпеки і права Національної академії правових наук України», 2025

© Науково-дослідний інститут інтелектуальної власності Національної академії правових наук України, 2025

© Колектив авторів, 2025

## З М І С Т

|                                                                                                                                                   |           |
|---------------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| <b><i>ВСТУП</i></b>                                                                                                                               | <b>8</b>  |
| <b><i>ВИСТУПИ НА ПЛЕНАРНОМУ ЗАСІДАННІ</i></b>                                                                                                     |           |
| <b><i>Іван ДОРОНІН</i></b>                                                                                                                        | <b>12</b> |
| Безпекові виклики та правові проблеми створення «національного штучного інтелекту» (у межах LLM)                                                  |           |
| <b><i>Микола КАРЧЕВСЬКИЙ</i></b>                                                                                                                  | <b>19</b> |
| Трирівнева модель правового регулювання ШІ                                                                                                        |           |
| <b><i>Олег ЗАЯРНИЙ</i></b>                                                                                                                        | <b>26</b> |
| Правове забезпечення застосування ШІ-агентів у механізмі цифрової трансформації територіальних громад: деякі концептуально-прикладні аспекти      |           |
| <b><i>Олена АНДРІЄНКО</i></b>                                                                                                                     | <b>32</b> |
| Спочатку було слово: правові виклики великих мовних моделей                                                                                       |           |
| <b><i>Наталія САВІНОВА</i></b>                                                                                                                    | <b>38</b> |
| Маніпулювання свідомістю з використанням візуальних продуктів ШІ: кримінологічні рефлексії                                                        |           |
| <b><i>Ігор КОРЖ, Тетяна КОРЖ</i></b>                                                                                                              | <b>43</b> |
| Майбутнє штучного інтелекту в науці: бачення зарубіжних фахівців                                                                                  |           |
| <b><i>Марія ДУБНЯК</i></b>                                                                                                                        | <b>50</b> |
| Проблеми гармонізації термінології при імплементації AI Act в Україні                                                                             |           |
| <b><i>Ярослава САВЧЕНКО</i></b>                                                                                                                   | <b>56</b> |
| Повернення культурних цінностей України: перспективи використання штучного інтелекту                                                              |           |
| <b><i>Софія АВДІЮК</i></b>                                                                                                                        | <b>64</b> |
| "Машинне рознавчання" (machine unlearning) як правовий імператив: забезпечення права на стирання даних в архітектурі великих мовних моделей (LLM) |           |

|                                                                                                                                      |            |
|--------------------------------------------------------------------------------------------------------------------------------------|------------|
| <b>Світлана КЕЛИП</b>                                                                                                                | <b>71</b>  |
| Нормативно-правова модернізація сфери прикордонної безпеки України в умовах розвитку технологій штучного інтелекту                   |            |
| <b>Михайло МИХАЙЛЕНКО</b>                                                                                                            | <b>74</b>  |
| Проблеми та перспективи використання технологій ШІ під час розгляду заявок про державну реєстрацію торговельних марок                |            |
| <b>Дмитро ЛАНДЕ, Юрій ЦИРУЛЬНЄВ</b>                                                                                                  | <b>81</b>  |
| Теоретична модель правового регулювання аспектів створення та використання електронних інформаційних ресурсів та електронних архівів |            |
| <b>ФІЛОСОФСЬКО-ПРАВОВІ, КОНЦЕПТУАЛЬНО-МЕТОДОЛОГІЧНІ ЗАСАДИ ДОСЛІДЖЕННЯ ТА РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ</b>                         |            |
| <b>Геннадій АНДРОЩУК</b>                                                                                                             | <b>86</b>  |
| Методологія виявлення нових технологій ШІ з використанням патентних даних                                                            |            |
| <b>Олександр БУТНІК-СІВЕРСЬКИЙ</b>                                                                                                   | <b>97</b>  |
| Розвиток інтелектуальної техніко-технологічної комп'ютерної платформи цифровізації                                                   |            |
| <b>Олеся АНДРУЩЕНКО</b>                                                                                                              | <b>103</b> |
| Філософсько-правові засади легітимації ШІ в системі політико-правових відносин сучасної України                                      |            |
| <b>Яна ЛУКАШОВА</b>                                                                                                                  | <b>107</b> |
| Правове регулювання штучного інтелекту у сфері національної безпеки і оборони                                                        |            |
| <b>Вадим ГАРАЩЕНКО</b>                                                                                                               | <b>111</b> |
| Правове регулювання штучного інтелекту в Європейському Союзі: теоретико-методологічні та практичні аспекти                           |            |
| <b>Владислав ВАРИНСЬКИЙ</b>                                                                                                          | <b>116</b> |
| Основні критерії заборон використання ШІ за AI Act                                                                                   |            |

|                                                                                                                                                                                                                 |            |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------|
| <b>ДЕГТЯРЬОВ І.М., ПЕТРОВСЬКИЙ М.В., ЛЕОНТЬЄВ П.В.</b>                                                                                                                                                          | <b>121</b> |
| Методологія тестування програмного забезпечення для системи автономної навігації наземних роботизованих комплексів                                                                                              |            |
| <b>ПЕТРОВСЬКИЙ М.В., ДЕГТЯРЬОВ І.М., ЛЕОНТЬЄВ П.В.</b>                                                                                                                                                          | <b>124</b> |
| Важливість розроблення методики випробувань для контролю технічних характеристик наземних роботизованих комплексів з використанням штучного інтелекту, як альтернатива перевірці в реальних умовах експлуатації |            |
| <b>БЕЗПЕКА, ОБОРОНА, КІБЕРЗАХИСТ ТА ВОЄННИЙ<br/>КОНТЕКСТ ПРИ РЕГУЛЮВАННІ ШІ</b>                                                                                                                                 |            |
| <b>Дар'я ГЛУШКОВА</b>                                                                                                                                                                                           | <b>130</b> |
| Штучний інтелект у системі національної безпеки України. правові засади регулювання в умовах воєнного стану                                                                                                     |            |
| <b>Олена ГРЕЗІНА</b>                                                                                                                                                                                            | <b>133</b> |
| Захист інформаційних ресурсів держави в умовах воєнних загроз та роль штучного інтелекту                                                                                                                        |            |
| <b>Вікторія КОЩИНЕЦЬ</b>                                                                                                                                                                                        | <b>136</b> |
| Симбіоз штучного інтелекту та розподілених технологій для захисту оборонних даних                                                                                                                               |            |
| <b>Марина ГРИГОР'ЄВА</b>                                                                                                                                                                                        | <b>142</b> |
| Формування етичних норм і стандартів при використанні штучного інтелекту в умовах воєнного стану                                                                                                                |            |
| <b>ФЕДОРЧЕНКО О.С.</b>                                                                                                                                                                                          | <b>148</b> |
| Штучний інтелект проти кіберзлочинців: еволюція захисту в епоху складних кібератак                                                                                                                              |            |
| <b>Ганна ФОРОС</b>                                                                                                                                                                                              | <b>153</b> |
| Автоматизований моніторинг кіберзагроз за допомогою систем машинного навчання                                                                                                                                   |            |

**ЦИФРОВИЙ СУВЕРЕНІТЕТ І НАЦІОНАЛЬНА СТРАТЕГІЯ  
РОЗВИТКУ У СФЕРІ  
ШТУЧНОГО ІНТЕЛЕКТУ**

|                                                                                          |            |
|------------------------------------------------------------------------------------------|------------|
| <b><i>АВДІЮК С.С., ГАБЕЛКО В. О., ДРАЧУК Є.М.</i></b>                                    | <b>157</b> |
| Правові механізми забезпечення цифрового суверенітету України в епоху штучного інтелекту |            |
| <b><i>Людмила ЗАСЛАВСЬКА</i></b>                                                         | <b>165</b> |
| Національна велика мовна модель як інструмент цифрового суверенітету України             |            |
| <b><i>Василь ОРИЩУК, Максим ВАЛІН</i></b>                                                | <b>170</b> |
| LLM та «Дія.АІ» як основа побудови Smart City-екосистем                                  |            |
| <b><i>НИКОЛИНА К.В.</i></b>                                                              | <b>177</b> |
| Національна ВММ як основа цифрового суверенітету України                                 |            |
| <b><i>Ярослав МАНУІЛОВ</i></b>                                                           | <b>182</b> |
| Стратегія розвитку міжнародного кіберпростору та цифрової політики США                   |            |

**ПРАВА ЛЮДИНИ, ПРИВАТНІСТЬ ТА ПИТАННЯ  
ІНТЕЛЕКТУАЛЬНОЇ ВЛАСНОСТІ**

|                                                                                                                                            |            |
|--------------------------------------------------------------------------------------------------------------------------------------------|------------|
| <b><i>ПАШИНСЬКИЙ В.Й., ЦЬОМЕНКО А.В.</i></b>                                                                                               | <b>186</b> |
| Адміністративно-правове забезпечення захисту персональних даних в умовах розвитку штучного інтелекту                                       |            |
| <b><i>Жанна ЛУПАК. Науковий керівник: ЗАЯРНИЙ О.А.</i></b>                                                                                 | <b>192</b> |
| Право особи на справедливе рішення в адміністративних процедурах із застосуванням штучного інтелекту суб'єктами публічного адміністрування |            |
| <b><i>Наталія ДЕДЮЄВА, Ольга ГОЛОВКО</i></b>                                                                                               | <b>198</b> |
| Баланс приватності та свободи вираження поглядів через призму законопроекту № 14057                                                        |            |
| <b><i>Юрій КАПІЦА</i></b>                                                                                                                  | <b>204</b> |
| Правові аспекти використання контенту для навчання штучного інтелекту: підходи в ЄС, Україні та США                                        |            |

|                                                                           |            |
|---------------------------------------------------------------------------|------------|
| <b>Олена БАХАРЕВА</b>                                                     | <b>210</b> |
| Штучний інтелект як «порушник» авторського права                          |            |
| <b>КОВАЛЕНКО Т.В.</b>                                                     | <b>215</b> |
| Штучний інтелект і права інтелектуальної власності                        |            |
| <b>Олена ГОНЧАРЕНКО, Анастасія ЛЕВКІВСЬКА</b>                             | <b>220</b> |
| Криза інтелектуальної власності в умовах генеративного штучного інтелекту |            |

## СОЦІАЛЬНИЙ ВИМІР ШТУЧНОГО ІНТЕЛЕКТУ: ОСВІТА, ПРАЦЯ, ІНКЛЮЗІЯ

|                                                                                   |            |
|-----------------------------------------------------------------------------------|------------|
| <b>Андрій ОЗАРЧУК</b>                                                             | <b>227</b> |
| Персоналізація та інклюзія: можливості штучного інтелекту для сучасної педагогіки |            |
| <b>Олександр ОСТАПЕНКО</b>                                                        | <b>232</b> |
| Штучний інтелект у сфері освіти та науки: правові виклики цифрової епохи          |            |
| <b>Світлана САДОВА</b>                                                            | <b>235</b> |
| Штучний інтелект і ринок праці: соціальні ризики та виклики перерозподілу вигод   |            |

## ВСТУП

5-6 листопада 2025 року в м. Києві відбулася науково-практична конференція : «Україна в умовах соціальної та цифрової трансформації: теоретичні моделі правового регулювання штучного інтелекту».

Її головною **метою** стало напрацювання теоретичної моделі правового регулювання штучного інтелекту в Україні, що відображають реальні зміни 2025 року – створення національної LLM, інтеграція ШІ-асистента в державний цифровий сервіс, розбудови безпечної освітньої екосистеми – і запропонувати нормативні, інституційні та етичні рішення, здатні забезпечити баланс між інноваціями, правами людини, безпекою і ефективністю.

Захід був організований спільно Державною науковою установою «Інститут інформації, безпеки і права Національної академії правових наук України» та Науково-дослідним інститутом інтелектуальної власності Національної академії правових наук України.

Участь к конференції прийняли представники наукової спільноти України, представники органів державної влади та громадськості.

### **Питання, що були розглянуті на конференції:**

1. Теоретичні та методологічні засади цифрових та інформаційних прав фізичних та юридичних осіб у контексті рекодифікації Книги другої Цивільного кодексу України.

2. Теоретичні та методологічні підходи до регулювання ШІ-агентів різного рівня: особистих, корпоративних, державних.

3. Правові стандарти та нормативні підходи до регулювання, створення й використання національної LLM у контексті мовної, культурної та національної безпеки.

4. Проблеми використання корпоративних моделей LLM в органах державної влади при обробці даних (персональних, з публічних електронних реєстрів тощо) та інші правові проблеми.

5. Регулювання штучного інтелекту в умовах воєнного стану та загроз кібербезпеці.

6. Проблеми застосування ШІ в суспільних відносинах.

7. Виклики захисту прав інтелектуальної власності через результати діяльності генеративних мовних моделей.

8. Етика та межі застосування ШІ в наукових дослідженнях та освітній діяльності.

## Головні підсумки науково-практичної конференції

– головною проблемою застосування ШІ в Україні залишається визначення та формалізація його правового статусу, що потребує прийняття парламентом рамкового Закону України «Про штучний інтелект» на зразок Регламенту в ЄС.

Правове регулювання ШІ, як інтегрованої рамки для подальшої наукової дискусії й розробки нормативних актів у даній сфері, вбачається шляхом запровадженням своєрідної трирівневої моделі :

а) мета-рівня, який формує аксіологічний фундамент системи регулювання;

б) стратегічного рівня, який задає принципи й орієнтири розвитку ШІ;

в) тактичного рівня, який забезпечує нормативне реагування на поточні ризики;

– сучасний етап розвитку України характеризується активною реалізацією процесів цифрової трансформації, і ключовою сферою цих перетворень є територіальні громади, на рівні яких відбувається безпосередня взаємодія влади з громадянами та бізнесом, і в даному контексті визначальним фактором стає підвищення ефективності муніципального управління, якості надання публічних послуг та розбудови «розумної» інфраструктури міст.

– продовжує перебувати у центрі уваги наукової спільноти та обговорюватися у медіа проект «національного штучного інтелекту» (ШІ, українського ШІ, UA AI), або іншими словами «національна велика мовна модель (LLM)». Такі LLM та «Дія.AI» відкривають нову сторінку в розвитку публічних послуг і цифрового урядування. Їхній потенціал може бути реалізований лише за умови дотримання принципів прозорості, безпеки, етичності та людиноцентричності.

перед національними законодавствами країн-членів ЄС, і України на шляху асоціації з ЄС, виникає потреба нормативного убезпечення від порушень, пов'язаних з підсвідомістю, яка є мішенню впливів на свідомість в спектрі від праймінгу до інформаційно-психологічної зброї вищого порядку. Зазначене вимагає введення в правовий тезаурус вже зараз таких понять, як: «свідомість», «підсвідомі впливи», «маніпуляція поведінкою», «впливи на прийняття рішень» тощо.

– учасники в різних дослідницьких системах можуть багато зробити для поглиблення використання ШІ в науці, посилення його позитивних ефектів, одночасно адаптуючись до швидкозмінних наслідків ШІ для управління дослідженнями, при цьому прискорення продуктивності досліджень може бути найбільш економічно та соціально цінним з усіх застосувань ШІ;

– впровадження технологій ШІ у механізм повернення культурних цінностей України, які були викрадені країною-агресором РФ, має розглядатися як складова державної політики у сфері національної безпеки.

– автоматизований моніторинг кіберзагроз на основі методів машинного навчання формує нову парадигму забезпечення кібербезпеки, орієнтовану на аналітичну обробку великих обсягів даних, самонавчання та адаптацію до змін

середовища. Реалізація подібних систем дозволяє підвищити точність виявлення викликів і загроз, зменшити час реагування на інциденти, оптимізувати використання ресурсів служби безпеки й створити динамічну модель захисту інформаційних ресурсів у мінливому кіберпросторі.

Зазначені та інші ідеї, отримані в процесі здійснення доповідей і у ході проведення наукових дискусій мають відповідну наукову цінність і можуть слугувати предметом для подальших наукових досліджень та дискусій.

## **РЕКОМЕНДАЦІЇ НАУКОВО-ПРАКТИЧНОЇ КОНФЕРЕНЦІЇ**

Стан цифрової трансформації в Україні характеризується позитивним розвитком у сфері телекомунікацій та впровадження електронних послуг, але стикається з викликами, такими як цифрова нерівність між регіонами та вразливість до кібератак.

Нині ми переходимо від Digital State до Agentic State — держави, де ШІ допомагає ухвалювати рішення, автоматизує процеси й створює новий рівень взаємодії з громадянами. Запуск Дія.АІ – це перший важливий крок на цьому шляху. Водночас, враховуючи прийняття Регламенту AI Act в ЄС та рекомендації ОЕСР (Організація економічного співробітництва та розвитку), учасники конференції вважають за доцільне висловити наступні пропозиції про конкретні кроки, які доцільно здійснити у процесі згаданої трансформації в Україні :

1. Подальший розвиток українського суспільства потребує невідкладного напрацювання та прийняття Закону України «Про штучний інтелект», гармонізований з AI Act ЄС, який би встановлював класифікацію ШІ-систем за рівнями ризику; визначав вимоги до високоризикових ШІ-систем; визначав правовий статус ШІ; регламентував питання прозорості та інформування користувачів; передбачав створення національного органу нагляду за ШІ визначав запровадження механізмів контролю та відповідальності, включно з незалежним моніторингом і аудитом алгоритмів тощо;

2. Рекомендувати Уряду України впровадити, за участі наукової спільноти та практиків, рекомендації ОЕСР щодо ШІ в національне законодавство, зокрема:

– прискорити впровадження національної стратегії розвитку ШІ, узгоджену з принципами ОЕСР і яка буде фокусуватися на практичному застосуванні ШІ для покращення державного управління, оборони, освіти, медицини та агросектору, а також на створенні якісної та доступної ШІ-інфраструктури та кадрового потенціалу;

– створення механізму підтримки досліджень та розробок у сфері ШІ;

– забезпечення розвитку цифрової інфраструктури ШІ;

– розробка галузевих стандартів використання ШІ, особливо в таких критичних сферах, як державне управління, охорони здоров'я, транспорту, фінансів тощо.

3. Висловити підтримку ініціативі щодо створення суверенного штучного інтелекту **Diiia AI LLM** – суверенної мовної моделі, розробленої спеціально під українське законодавство, державні послуги та запити громадян, на основі якої працюватимуть усі сервіси штучного інтелекту в екосистемі «Дії» та розробки **AI Factory** – суверенної фабрики ШІ, для створення національної інфраструктури ключових державних сервісів на базі ШІ, що корелюється з положеннями **Стратегії цифрового розвитку інновацій України до 2030 року**.

4. Вважати необхідним напрацювання в Україні подальших змін до законодавства :

- у сфері місцевого самоврядування – доповнити його положеннями, що регулюють повноваження органів місцевого самоврядування у сфері впровадження та використання ШІ-систем, а також обов'язок проведення публічних консультацій перед їх розгортанням;

- розробити та затвердити на рівні місцевих рад **Типовий регламент (політику) етичного та відповідального застосування ШІ-агентів**, який би деталізував процедури оцінки ризиків, механізми громадського контролю, порядок інформування мешканців та протоколи реагування на інциденти;

- у сфері національної безпеки – щодо використання штучного інтелекту для пошуку незаконно вивезених культурних цінностей України; доповнень положень щодо удосконалення державного управління у даній сфері, як то розподіл обов'язків між державними органами, фінансування відповідної ініціативи, створення даних та наповнення реєстрів тощо, а також формування нової парадигми управління культурною спадщиною;

- у сфері інформаційної безпеки – необхідність активного застосування блокчейн, як технології з новим підходом в таких напрямках як інформаційна безпека, автентифікація, цілісність та стійкість даних та інших.

5. Вбачається доцільним здійснити побудову автоматизованої системи моніторингу кіберзагроз, що потребує розроблення архітектури, яка включатиме кілька взаємопов'язаних рівнів.

## ВИСТУПИ НА ПЛЕНАРНОМУ ЗАСІДАННІ

**Іван Доронін**

доктор юридичних наук, доцент, керівник  
наукового центру Державної наукової установи  
«Інститут інформації, безпеки і права  
Національної академії правових наук України»  
ORCID: <https://orcid.org/0000-0002-5991-6713>.

### **БЕЗПЕКОВІ ВИКЛИКИ ТА ПРАВОВІ ПРОБЛЕМИ СТВОРЕННЯ «НАЦІОНАЛЬНОГО ШТУЧНОГО ІНТЕЛЕКТУ» (У МЕЖАХ LLM)**

Проект «національного штучного інтелекту» (ШІ, українського ШІ, UA AI) активно обговорюється у засобах масової інформації та перебуває у центрі медійної уваги хоча більш правильним буде вести мову про «національну велику мовну модель». Саме так сформульована мета відповідальним за політику в сфері цифрової трансформації державним органом.

У червні 2025 року Міністерство цифрової трансформації України (Мінцифри) визначило завдання створення «національної великої мовної моделі» (LLM). Завдання її створення визначено як необхідність подальшого використання «у цифрових державних і бізнес-продуктах зі штучним інтелектом, щоб зробити сервіси зручнішими для людей та допомогти організаціям працювати ефективніше» [1]. Таким чином можливо зрозуміти, що у певних продуктах (державних сервісах та бізнес проектах) буде використовуватись LLM власної розробки, яка чітко афілійована з Україною.

У цьому аспекті варто було б звернути увагу на два моменти. Перша заява Мінцифри звучить наступним чином: «Українська LLM даватиме точніші відповіді, ніж глобальні моделі, бо буде додатково натренована (pre-trained) на українських даних. Вона краще розумітиме мовні особливості — діалекти, терміни, контекст — і добре орієнтуватиметься в темах, пов'язаних з українською історією, культурою та суспільним досвідом. Так, модель відповідатиме правильно не лише мовно, а й за змістом» [1].

І друга заява: «Національна LLM дає змогу зберігати й обробляти дані всередині країни, що критично важливо для безпеки у сфері оборони, державних установ, медицини та фінансів. Модель забезпечить захист

персональних даних і національної безпеки України при інтеграції AI в різні сектори» [1].

Оскільки належну критику зазначеної мети і завдань забезпечать фахівці безпосередньо у сфері технологій III, варто зосередитись на певному колі визначених проблем до числа яких необхідно віднести безпекові виклики та викликані ними окремі правові проблеми.

23 вересня 2025 року керівник цього проекту Д.Овчаренко, що є технічним директором WINWIN AI Center of Excellence при Мінцифри, дав доволі цікаве інтерв'ю [2], що дозволяє визначити та окреслити коло вказаних вище викликів і проблем.

По-перше варто зазначити, що мова йде не про розробку і створення власної LLM з нуля, а про доведення існуючої моделі з числа open-source model. Для цієї мети обрано конкретний продукт і судячи з усього це про продукти Google (Gemini для Дія AI) як проміжне рішення. Наразі немає чіткого і неоднозначного офіційного підтвердження, що для подальшого тренування LLM обрано саме моделі з числа Google або Gemini (Gemini Ultra, Gemini Pro, Gemini Flash, and Gemini Nano з усіма версіями) або ж Gemma (як більш легка), але вірогідність цього висока. Можливо також, що мова буде йти і про іншу LLM. Але навряд чи буде обрано модель, що чітко не афілійована з відомими IT-корпораціями (GPT, Claude, Llama, Copilot, Grok або Mistral). У даному випадку обрання конкретної вихідної open-source LLM є управлінським рішенням розробників. Водночас, розробниками визнається, що мова буде у подальшому йти про низку висхідних LLM, для створення на локальному рівні власних SLM.

По-друге, характеризуючи процес тренування LLM, керівник проекту визначає, що воно буде відбуватись шляхом формування корпусу даних і встановлення «культурно-історичних» та «етичних» критеріїв. Тобто дороблена LLM буде чітко відображати закладену політику формування відповідей, та, відповідно, цензурування, оскільки процес обмеження інформації, що видає LLM (внаслідок генерації), і є власне цензуруванням інформації. І ця політика буде державною.

Таким чином, з'ясувавши загальний стан навколо створення LLM зі специфічними завданнями, варто проаналізувати і проблемне коло стосовно цього процесу.

Щоб упередити можливу заангажованість варто зазначити, що національні та суверенні проекти власних LLM у принципі є світовим трендом, особливо в країнах з особливими лінгвістичними потребами (Японія, Південна Корея, КНР, країни Близького Сходу). Інколи завдання поставлено як створення «надійної» або «прозорої» моделі. Прикладом є швейцарська LLM Apertus (Swiss AI), яка є спеціально розробленим продуктом в рамках Swiss AI Initiative концерном вищих технічних навчальних закладів та Швейцарського центру суперкомп'ютерів (CSCS). Варто окремо звернути увагу на існування казахської LLM (Kaz LLM), що на думку фахівців є вдалим прикладом застосування ШІ для регіональних і мало розповсюджених мов [3]. Щоправда, розробник - Інститут розумних систем та штучного інтелекту (ISSAI) у м. Астана стверджує, що модель розроблено з нуля, а не на базі open-source model.

На нашу думку основним безпековим викликом є питання відповідальності за продукт. При використанні як основи іншої моделі у будь-якому випадку така відповідальність буде роздільною з розробником вихідної моделі. Річ у тім, що всі корпорації при використанні продуктів на основі LLM роблять заяви щодо зменшення або зняття з себе юридичної відповідальності за згенеровану інформацію та її використання отримувачами інформації. Це викликано, серед інших факторів, невід'ємними властивостями самих LLM.

У цьому аспекті зазвичай постає проблема «галюцинації AI», коли система генерує неправдиву відповідь, «вигадує» (а насправді конструює зі шматків уже наявної інформації) факти. Прикладом є відповідь чат-бота про сюжет маловідомого літературного твору відомого письменника. Підібравши інформацію щодо сюжету і зіткнувшись з відсутністю (або множинністю варіантів) імені головного героя у джерелах, LLM сконструює його підібравши із числа розповсюджених у цій місцевості і в цьому періоді, що виглядає як «вигадує». Схожі явища можливі і при перекладі джерел на іншу мову. Традиційно причинами виникнення цього явища (не менше 3% для найкращих LLM і 20% для більш слабких) вважались технічні помилки, недостатня навченість, обмеженість даних, «упередженість» [4, с. 114]. Хоча останні дослідження стверджують, що таке явище є невід'ємною властивістю LLM [5, 6, 7]. Проблема визнана дослідниками OpenAI і визначена у відповідному звіті, що зараз активно обговорюється технічними фахівцями [8]. Саме тому великі корпорації і

роблять відповідні заяви про відмову від відповідальності та вказують в умовах використання продуктів, що згенерована інформація не може використовуватись в медичних, юридичних, фармацевтичних і т. ін. цілях.

Окрім цього пов'язаним із загальним питанням відповідальності за продукт є безпекові упередження щодо належності, допустимості та бездоганності корпусу інформації для технічного навчання моделі. Річ у тім, що на думку керівника проекту української LLM існує великий ризик і юридична невизначеність навколо використання наявних корпусів інформації в бібліотеках через можливе порушення авторського права. У даному випадку варто визнати, що проблема регламентації прав інтелектуальної власності через використання творів у навчанні LLM дійсно існує насамперед в юрисдикціях США та ЄС. При цьому законодавство азійських країн, як правило, не враховує ці положення або встановлює значні виключення з охорони прав у випадках застосування технологій ШІ. У даному випадку, зважаючи на євроінтеграційний курс нашої держави, важко буде ухвалити законодавство, що не відповідатиме європейському.

Ситуація із захистом персональних даних у тому числі на відповідність європейському законодавству дещо краща оскільки шифрування інформації про особу при обробці є вдалим технічним рішенням. Варто лише визнати, що у випадку з використанням у LLM це питання менш складне ніж з дотриманням законодавства про інтелектуальну власність.

Окрім визначених і визнаних проблем неодмінно постануть питання дотримання законодавства про інформацію при формуванні корпусу даних, отриманих від державних органів. Загальний підхід «відкритих даних» полягає у відкритості усієї інформації, окрім тієї, доступ до якої прямо обмежено законом. Тому при навчанні LLM увесь корпус інформації, на якому модель навчається, має сприйматись саме як «відкриті дані». У вітчизняному законодавстві зазначене регламентовано Положенням про набори даних, які підлягають оприлюдненню у формі відкритих даних, затверджене постановою Кабінету Міністрів України від 21.10.2015 № 835 з відповідними змінами і доповненнями. Тому все, що відповідає цьому Положенню, також має складати корпус даних для навчання LLM. Але одночасно може виникнути необхідність і доступу до інших даних особливо на етапах подальшої тонкої настройки (fine tuning), що може створити певні складнощі правового характеру. Так, на сьогодні відповідно до вимог частини 1 статті 21 Закону

України «Про інформацію» інформацією з обмеженим доступом є конфіденційна, таємна та службова інформація. Найбільш повне врегульовано обіг таємної інформації і цілком зрозуміло, що вона не може бути використана для навчання LLM, але тоді важко буде досягнути визначену Мінцифри мету «зберігати й обробляти дані всередині країни, що критично важливо для безпеки у сфері оборони». У випадку зі службовою інформацією питання ще більш складне, оскільки переліки службової інформації складаються відповідно до законодавства не централізовано, що дуже ускладнить процес використання для навчання моделей.

Крім цього виникатимуть питання правдивості відомостей, що формуватимуть корпуси для навчання та подальшої настройки моделі. Річ у тім, що забезпечити правдивість (без зосередженості на інших властивостях – на кшталт упередженості та стереотипів) доволі важко, а з точки зору права ще важче встановити відповідального за правдивість такої інформації. Це до речі інша причина для зняття з себе відповідальності корпораціями – виробниками LLM. Зрозуміло, що проблема має технічне вирішення і в даному випадку моделі досить складні, які здатні верифікувати інформацію логічним шляхом на етапі отримання, але в будь-якому разі їхня точність у цьому аспекті не становить 100%.

Окремим питанням є вироблення та реалізація політики. Зрозуміло, що існує необхідність постійної реалізації настанов політики, а у даному випадку мова йде про політику держави. Як зазначають розробники, цим займаються «етичний, технічний, мовознавчий та культурно-історичний борди». Таким чином, у питаннях етики, культури та історії можливо говорити про тенденції до упередженості. Це явище для LLM загальновідоме і у крайній формі яскраво ілюструється відповідями китайських продуктів (на кшталт DeepSeek). Але існують і численні звинувачення в упередженості найбільш відомих LLM, а саме у «лівацькій» чи «лівій» упередженості, у тому числі на користь лівого крила Демократичної партії США [9]. Зрозуміло, що під особливою критикою перебуває Grok через особистість власника корпорації xAI. Тому у випадку декларування «державності» та «національної» афіліації закидів щодо упередженості уникнути не вдасться. Тим більше, що до надання відповідей таким чат-ботом буде особливо прискіплива увага суспільства.

Слід також зазначити, що при реалізації у продукті напрямів задекларованої політики неможливо уникнути обмежень у видачі інформації

через встановлення відповідних фільтрів, що є цензуруванням інформації. Загалом політика корпорацій у цьому аспекті доволі аморфна, хоча і слідує загальним трендам заборони. Водночас у багатьох випадках чіткі критерії для цензурування відсутні та іноді їх взагалі неможливо встановити. До того ж виникає питання щодо демократичності процесу, дотримання прав і відсутності дискримінації при виробленні засад політики «бордами» (за виразом керівника проекту). Не слід також забувати, що вони відображатимуть відповідну політику (повістку) на час своєї роботи і це може змінюватись. До того ж не варто забувати, що вже існують моделі, що розроблені ентузіастами і позиціонуються як «українська LLM», що базується на «українських цінностях» [10].

Отже, викладене стосувалось лише питання розробки (формування корпусу, навчання та тонкої настройки) для національної LLM. У випадку впровадження такої моделі в системи AI для державних послуг мова буде йти фактично про регламентацію використання AI-агента, що породжуватиме ще більше проблемних питань. У цьому аспекті важливим є усвідомлення суті проблем, правильна постановка завдань та необхідність дотримання чинного законодавства на всіх етапах розробки моделі.

### **Список використаних джерел:**

1. Україна починає розробку національної мовної моделі: заява прес-офісу Мінцифри. URL: <https://thedigital.gov.ua/news/technologies/ukraina-pochinae-rozrobku-natsionalnoi-movnoi-modeli-mintsifra-ta-kiivstar-pidpisali-memorandum-pro-spivpratsyu>
2. «Люди стали ламати Дія AI, щойно ми зарелізилися». Дмитро Овчаренко з Мінцифри — про національну LLM і захист персональних даних// DOU. 23.09.2025. URL: <https://dou.ua/lenta/interviews/ovcharenko-diia-ai-and-llm/>
3. Wiggings Dion. Case Study: What We Can Learn from ISSAI KAZ-LLM and the AI Potential for Low Resource Languages// Medium. 2025. Feb. 11. URL: <https://medium.com/@dion.wiggings/case-study-what-we-can-learn-from-issai-kaz-llm-and-the-ai-potential-for-low-resource-languages-6606fd5bc594>
4. Махно Є., Руденко Є., Судніков Є., Тищенко М. Галюцинації штучного інтелекту у сфері освіти та науки: причини, наслідки та методи мінімізації// Повітряна міць України. 2025. № 1. С. 111-126. DOI 10.33099/2786-7714-2025-1-8-111-126.

5. Xu Z, Jain S, Kankanhalli M. Hallucination is inevitable: An innate limitation of large language models. arXiv preprint arXiv:2401.11817. 2024 Jan 22. URL: <https://doi.org/10.48550/arXiv.2401.11817> DOI 10.48550/arXiv.2401.11817
6. Ghosh B. LLMs Will Always Hallucinate// Medium. 2024. Sep. 12. URL: <https://medium.com/@bijit211987/llms-will-always-hallucinate-b06f3353b159>
7. Razzaq A. From Pretraining to Post-Training: Why Language Models Hallucinate and How Evaluation Methods Reinforce the Problem// Marktechpost. 2025, Sep. 6. URL: <https://www.marktechpost.com/2025/09/06/from-pretraining-to-post-training-why-language-models-hallucinate-and-how-evaluation-methods-reinforce-the-problem/>
8. Kalai A.T, Nachum O, Vempala S.S., Zhang E. Why language models hallucinate. arXiv preprint arXiv:2509.04664. 2025 Sep 4. URL: <https://arxiv.org/abs/2509.04664>
9. Hollman M., Davey J., Clark E. LLMs are Left-Leaning Liberals: The Hidden Political Bias of Large Language Models// Society for Computers & Law. 2025, May 19. URL: [https://www.scl.org/llms-are-left-leaning-liberals-the-hidden-political-bias-of-large-language-models/#\\_ftn9](https://www.scl.org/llms-are-left-leaning-liberals-the-hidden-political-bias-of-large-language-models/#_ftn9)
10. Українські дослідники представили Lapa LLM — першу національну ШІ-модель для міркування // DEV.ua 2025, 3 жовтня. URL: <https://dev.ua/news/ukrainski-doslidnyky-predstavly-lapa-llm-pershu-natsionalnu-shi-model-dlia-mirkuvannia-1759481607>

## **Микола Карчевський**

доктор юридичних наук, професор, професор  
кафедри права імені академіка УАН о. Івана  
Луцького Університету Короля Данила;  
головний науковий співробітник наукової  
лабораторії проблем штучного інтелекту і  
права Інституту інформації, безпеки і права  
Національної академії правових наук України

### **ТРИРІВНЕВА МОДЕЛЬ ПРАВОВОГО РЕГУЛЮВАННЯ ШІ**

Мета дослідження. Запропонувати теоретичну рамку для осмислення та побудови системи правового регулювання штучного інтелекту на основі трирівневої моделі (мета-рівень, стратегічний і тактичний рівні), що поєднує ціннісні, концептуальні та нормативні аспекти в єдину логічну систему.

За останні два десятиліття ШІ пройшов шлях від абстрактних концептів до прикладних інструментів, що стали буденністю. Системи штучного інтелекту використовуються практично в усіх сферах діяльності людини. Відбулися зміни у науковій рефлексії та правовому регулюванні соціалізації штучного інтелекту.

Сучасні наукові підходи розглядають штучний інтелект як програму, що імітує інтелектуальні функції людини. Така імітація базується на статистичних методах, відомих як машинне навчання. Загалом воно означає програмування без класичного кодування, а шляхом завантаження великих масивів навчальних даних із правильними прикладами. Система виявляє закономірності та формує моделі, здатні прогнозувати чи приймати рішення у схожих випадках.

Нині великі мовні моделі (LLM) досягають видатних результатів у роботі з текстами, розпізнаванні образів, створенні контенту тощо. Проте, попри здатність імітувати людину, вони залишаються лише інструментами для розв'язання конкретних завдань. Водночас, значимість використання ШІ, його потенційні можливості та масштаб ризиків, з необхідністю приводять до питання правового регулювання використання технологій ШІ.

Як рівні сучасної системи правового регулювання штучного інтелекту, на нашу думку, слід розглядати: мета-рівень, стратегічний рівень, тактичний

рівень. Якщо стратегічний рівень окреслює довгострокові вектори розвитку ІІІ, а тактичний забезпечує нормативне реагування на поточні ризики, то мета-рівень відповідає за осмислення цілей, соціального призначення та відповідності ІІІ-систем людським цінностям<sup>1</sup>.

Мета-рівень пов'язаний з проблемою невизначеності (alignment problem) та нечіткості людських цінностей. Ми намагаємося регулювати технологію, якій, теоретично, доведеться приймати рішення на основі наших суспільних цінностей та пріоритетів, але при цьому не можемо досягти консенсусу щодо цих цінностей та пріоритетів у суспільстві.

Це породжує парадокс: ми вимагаємо від ІІІ дотримуватися "людських цінностей", але самі не можемо чітко артикулювати, що це означає в кожному конкретному контексті.

На проблемі відсутності чіткості, стабільності та внутрішньої узгодженості людських цілей наголошує і Н. Бостром, він переконливо показує, що не тільки технології мають бути узгоджені з людськими цілями, а й самі цілі – вивчені, усвідомлені та структуровані. Поки цього не досягнуто, ризик дезорієнтації й навіть екзистенційної кризи від використання ІІІ залишається надзвичайно високим<sup>2</sup>.

Тому на мета-рівні має здійснюватися аналіз того, для чого в кінцевому рахунку створюється та впроваджується конкретна технологія ІІІ, яку проблему вона покликана розв'язати, яким цінностям - сприяти, і які ризики для суспільних пріоритетів вона може генерувати. В умовах, коли навіть самі людські цілі часто залишаються неявними, суперечливими або нефіксованими у правових інститутах, мета-рівень виконує функцію узгодження змісту правового регулювання із реальними потребами людини і суспільства.

Яке питання має бути в центрі мета-рівня? Тут важливо оцінити накопичений досвід використання ІІІ. Так, значний досвід накопичено у кримінальній юстиції. Використання ІІІ для забезпечення єдності судової практики наочно продемонструвало, що на перший план виходить наступне питання: чи ІІІ забезпечує справжню справедливість, чи алгоритмізація лише

---

<sup>1</sup> Карчевський, М. В., & Карчевська, О. В. (2025). Структурний вимір правового регулювання ІІІ. Український політико-правовий дискурс, (12). <https://doi.org/10.5281/zenodo.15782211>

<sup>2</sup> Bostrom N. Deep Utopia: Life and Meaning in a Solved World. Ideapress Publishing, 2024. 525 p.

стандартизує рішення на основі потенційно недосконалих даних. Результати впровадження систем прогнозування місць вчинення кримінальних правопорушень показали схожу проблему. Чим більше поліцейських направлялося у заданий район тим більшою була кількість виявлених у цьому районі правопорушень. Алгоритм фіксував прийняте рішення як правильне і продовжував рекомендувати посилені наряди для визначених районів<sup>3</sup>. У такий спосіб увічнювався «кримінальний» статус таких районів але загальне використання ресурсів поліції не ставало більш ефективним, загальний рівень безпеки громадян не підвищувався. Так званий ефект «зворотної петлі» (*feedback loop*) має місце коли підвищена увага поліції до “прогнозованих” районів призводить до ще більшої реєстрації там злочинів і, як наслідок, подальшого укріплення хибного прогнозу<sup>4 5</sup>.

Такий досвід використання дозволяє наступним чином сформулювати питання мета-рівня: III використовується для підвищення ефективності кримінальної юстиції шляхом персоналізації впливу, чи відбувається подальше алгоритмічне усереднення рішень, орієнтованих на «середніх» людей? Саме алгоритмічне усереднення у сфері кримінальної юстиції становить, на нашу думку, найсерйознішу небезпеку застосування III. Коли система починає орієнтуватися не на конкретну особу з її унікальними обставинами, а на «усереднений» профіль злочинця, відбувається тотальне вирівнювання правосуддя. Таке правосуддя втрачає сенс індивідуалізації покарання й перетворюється на механічне застосування шаблонів. У цьому полягає ризик зведення складної правової реальності до статистичної формули: людина починає сприйматися не як індивід, а як представник певної статистичної групи. Це призводить до ситуації, яку можна назвати

---

<sup>3</sup> Див.: Ensign, D., Friedler, S.A., Neville, S., Scheidegger, C. & Venkatasubramanian, S.. (2018). Runaway Feedback Loops in Predictive Policing. Proceedings of the 1st Conference on Fairness, Accountability and Transparency, in Proceedings of Machine Learning Research 81:160-171 Available from <https://proceedings.mlr.press/v81/ensign18a.html>.; Hung, T. W., & Yen, C. P. (2023). Predictive policing and algorithmic fairness. Synthese, 201(3), 1-25. <https://doi.org/10.1007/s11229-023-04189-0>

<sup>4</sup> Joh, E.E., 2016. The New Surveillance Discretion: Automated Suspicion, Big Data, and Policing. Harvard Law & policy review, 10, 15–42.

<sup>5</sup> Ferguson, A.G., 2017. The rise of Big data policing: surveillance, race, and the future of Law enforcement. New York: NYU Press.

«алгоритмічним конвеєром покарання» - системи, здатної швидко ухвалювати рішення, але нездатної бути справедливою.

Отже, ключовою проблемою використання ШІ в сфері кримінальної юстиції можна вважати втрату індивідуалізації рішень через надмірне використання ШІ та алгоритмічну упередженість.

Такий підхід до визначення ключової проблеми може бути застосований і в інших сферах використання ШІ. Наприклад, використання ШІ для контролю за виробництвом приводить до проблеми надмірної уніфікації діяльності. На основі спостереження за всіма працівниками ШІ система пропонує висновки про, наприклад, найбільш оптимальну швидкість та порядок рухів тіла. Реалізація таких пропозицій керівниками виробництва, як прогнозує ШІ, приведе до значного збільшення прибутків. З часом робота на такому модернізованому виробництві навряд чи буде відповідати ідеям трудового права, гідності працівників. Маємо схожу ситуацію – втрата індивідуалізації через алгоритмізацію.

Нарешті класичний приклад індивідуалізації – рекомендаційні алгоритми розважального контенту (Netflix, наприклад) створені для того, щоб максимально задовольняти потреби найрізноманітніших користувачів, рекомендуючи їм контент на основі їх глядацького досвіду. Такі алгоритми можуть як розвивати індивідуальність так і зводити наявне культурне різноманіття до кількох типів, виховуючи та, подібно до Прокруста, формуючи однаковість.

Таким чином, фундаментальним конфліктом, що зумовлює необхідність правового регулювання штучного інтелекту, є протиставлення індивідуалізації та колективізації. Інакше кажучи, йдеться про давній філософський конфлікт між унікальністю людського “я” та узагальненням “ми” – між неповторністю окремої особистості та прагненням системи до передбачуваності й рівності.

Цей конфлікт має глибоке історичне коріння. Платон бачив ідеал у гармонії колективного блага, тоді як Арістотель – у досягненні щастя кожного громадянина. У новітній філософії кантівська автономія особи протиставляється утилітаристській логіці суспільної користі. Тепер цей самий конфлікт відтворюється у цифровій формі: алгоритми оперують статистичними класами, класифікацією та рівнями кластеризації, тоді як право за своєю суттю є інструментом індивідуалізації. Алгоритм працює з типовим,

право — з неповторним. Саме в цьому полягає виклик правового регулювання ІІІ: знайти баланс між ефективністю колективного аналізу даних та індивідуалізацією, яка не зводиться до статистики.

Тому для правового регулювання ІІІ мета-рівень, на нашу думку, необхідно визначати наступним чином: *використання ІІІ здійснюється для максимальної індивідуалізації прогнозів на основі аналізу великих обсягів релевантних даних*. Таким чином, мета-рівень формує аксіологічний фундамент («Навіщо, за ради чого?») усієї системи правового регулювання ІІІ. Відповідно до нього мають визначатися як стратегічні пріоритети розвитку, так і тактичні інструменти правового реагування. Мета-рівень правового регулювання ІІІ полягає у забезпеченні відповідності технологічних рішень людським цінностям, а саме - у збереженні індивідуальності, гідності та свободи вибору в умовах масової алгоритмізації суспільних процесів.

Наступний рівень – стратегічний. Він зосереджується на довгострокових викликах і гіпотетичних сценаріях розвитку ІІІ, таких як поява автономного інтелекту, трансгуманізм, формування нових форм юридичної відповідальності та «юстиції штучного інтелекту». Основна мета цього рівня – визначення загальних координат для правового осмислення соціалізації технологій, зокрема з урахуванням технологічної нейтральності, необхідності конвергенції технічних і юридичних наук, а також переосмислення правового статусу як людей, так і машин.

Можливо найбільше дискусій тут викликає проблема суб'єктності ІІІ. О.А. Баранов зазначав, що теоретичні погляди на юридичний статус штучного інтелекту загалом поділяються на негативну теорію, позитивну та теорію компромісу. Негативна теорія полягає у розгляді ІІІ виключно як об'єкта правових відносин. Позитивна теорія стверджує, що ІІІ повинен мати статус суб'єкта права (фіктивна особистість, електронна особистість, ІІІ – суб'єкт злочину<sup>6</sup> тощо). Теорія компромісу – штучний інтелект повинен мати кваліфікацію суб'єкта права, але з обмеженою правосуб'єктністю. Також дослідник пропонує власний погляд – теорію доцільності. Вона полягає у тому, що «залежно від ситуаційної доцільності застосування спеціалізованого

---

<sup>6</sup> Див. наприклад: Радутний О. (2018). Штучний інтелект як суб'єкт кримінального права. Право України, 1, 123. doi:10.33498/lopu-2018-01-123.

робота із ШІ для заміни праці або діяльності традиційного суб'єкта він може бути об'єктом або суб'єктом правовідносин»<sup>7</sup>.

Поділяючи позицію шановного колеги, вважаємо що відповідь на стратегічне питання правового регулювання ШІ – об'єкт чи суб'єкт відносин – необхідно шукати в площині фактичного існування. Проблематика ШІ тісно пов'язана з гіпотетичними концепціями майбутніх технологій. Водночас, правове регулювання забезпечує упорядкування суспільних відносин, пов'язаних з технологіями, що існують реально. Зараз це системи ШІ, що виконують вузькі задачі під людським контролем та на основі даних, завантажених людьми, так званий «вузький» або «слабкий» ШІ. Відповідно, стратегічний рівень правового регулювання ШІ передбачає розгляд таких систем виключно як об'єктів права.

Нарешті тактичний рівень. Охоплює регулювання існуючих і впроваджених ШІ-систем, що використовуються в конкретних сферах. Він передбачає реагування на вже наявні ризики. Регуляторні рішення на цьому рівні ґрунтуються на класифікації ризиків, вимогах до прозорості, змісту навчальних наборів даних, визначенні меж використання технологій ШІ тощо.

Фахівці виділяють чотири основні групи ризиків застосування штучного інтелекту: приватність, маніпулятивний вплив, упередженість, непрозорість<sup>8</sup>. Питання їх подолання формують основу сучасної наукової дискусії щодо тактичного рівня правового регулювання ШІ.

Висновки. Запропоновано трирівневу модель правового регулювання штучного інтелекту як інтегровану теоретичну рамку для подальшої наукової дискусії й розробки нормативних актів у сфері ШІ.

Мета-рівень формує аксіологічний фундамент системи регулювання – визначає цінності, цілі та соціальне призначення технологій ШІ. На мета-рівні ключовим завданням є узгодження технологічного розвитку з базовими

---

<sup>7</sup> Баранов О. Штучний інтелект і система права : монографія / О. Баранов ; ДНУ «Інститут інформації, безпеки і права Національної академії правових наук України». – Київ, Одеса : Фенікс, 2025. - С. 283-284

<sup>8</sup> Dupont B, Stevens Y, Westermann H and Joyce M, *Artificial Intelligence in the Context of Crime and Criminal Justice* (Korean Institute of Criminology, Canada Research Chair in Cybersecurity, ICCC, Université de Montréal, 2018) 228 <[https://www.cicc-iccc.org/public/media/files/prod/publication\\_files/Artificial-Intelligence-in-the-Context-of-Crime-and-Criminal-Justice\\_KICICCC\\_2019.pdf](https://www.cicc-iccc.org/public/media/files/prod/publication_files/Artificial-Intelligence-in-the-Context-of-Crime-and-Criminal-Justice_KICICCC_2019.pdf)>

цінностями та суспільними пріоритетами, адже саме тут формулюється бачення, балансу між алгоритмізованою ІІІ стандартизацією та індивідуалізацією правового впливу, збереженням персональної унікальності.

Стратегічний рівень задає принципи й орієнтири розвитку, визначаючи співвідношення між технологічними можливостями й правовими межами, передбачає аналіз реально існуючих технологій, адже саме вони формують актуальні виклики для правової системи. Актуальний рівень технологій визначає, що на стратегічному рівні системи ІІІ мають розглядатися як об'єкти права. Зміна технологічного рівня, поява та широке використання так званого «загального ІІІ» можливо зумовить необхідність перегляду змісту стратегічного рівня.

Тактичний рівень забезпечує нормативне реагування на поточні ризики, від захисту приватності та попередження дискримінаційних результатів до забезпечення прозорості алгоритмів та створення механізмів відповідальності.

**Олег Заярний**

доктор юридичних наук, професор, професор  
кафедри інтелектуальної власності та  
інформаційного права Навчально-наукового  
інституту права Київського національного  
університету імені Тараса Шевченка  
ResearcherID: S-3358-2018  
ORCID: <https://orcid.org/0000-0003-4549-7201>

## **ПРАВОВЕ ЗАБЕЗПЕЧЕННЯ ЗАСТОСУВАННЯ ІІІ-АГЕНТІВ У МЕХАНІЗМІ ЦИФРОВОЇ ТРАНСФОРМАЦІЇ ТЕРИТОРІАЛЬНИХ ГРОМАД: ДЕЯКІ КОНЦЕПТУАЛЬНО-ПРИКЛАДНІ АСПЕКТИ**

Сучасний етап розвитку України характеризується активною реалізацією процесів цифрової трансформації різних сфер суспільного життя. Ключовою сферою цих перетворень є територіальні громади, на рівні яких відбувається безпосередня взаємодія влади з громадянами та бізнесом. У цьому зв'язку, впровадження інноваційних технологій, зокрема систем штучного інтелекту (далі – ІІІ), стає визначальним фактором підвищення ефективності муніципального управління, якості надання публічних послуг та розбудови «розумної» інфраструктури міст і [1, с. 28].

**Актуальність теми дослідження** зумовлена тим, що ІІІ-агенти, як просунутий клас автономних інтелектуальних систем, відкривають додаткові можливості для автоматизації складних управлінських процесів, оптимізації використання ресурсів та посилення представницької демократії на місцевому рівні. Водночас їх імплементація створює значні ризики, пов'язані із захистом прав людини, забезпеченням кібербезпеки, відкритості процедур прийняття управлінських рішень можливою дискримінацією, зумовленою алгоритмами роботи ІІІ [2, с. 104].

Відсутність в Україні належного правового регулювання цих суспільних відносин зумовлює виникнення системних ризиків для прав і свобод людини у сфері місцевого самоврядування, безпеки муніципальних електронних інформаційних ресурсів.

Виходячи із зазначеного вище, в Україні існує потреба у формуванні збалансованої та комплексної нормативно-правової основи для розробки,

впровадження та використання ІІІ-агентів у діяльності територіальних громад.

**Метою доповіді** є дослідження юридичних особливостей застосування ІІІ-агентів у процесі цифрової трансформації територіальних громад України, визначення основних викликів та ризиків, а також розробка науково обґрунтованих рекомендацій щодо вдосконалення національного законодавства у цій сфері регулювання.

В межах цього дослідження ми будемо виходити з того, що ІІІ-агент – це програмно-апаратна система, яка здатна автономно сприймати оточуюче середовище за допомогою сенсорів, інтерпретувати отримані дані, приймати рішення для досягнення поставлених цілей та діяти в цьому середовищі за допомогою виконавчих механізмів, а також демонструвати властивості активної поведінки та інтелектуального опрацювання поставлених завдань [3, с. 35].

На відміну від традиційних ІІІ-систем, агенти характеризуються вищим ступенем автономності та здатністю до самостійного функціонування в динамічному середовищі без постійного втручання людини.

Системний аналіз сучасних практик використання ІІІ-агентів в зарубіжних державах дозволяє виокремити такі основні напрями потенційно можливого застосування згаданих технологій у територіальних громадах є:

**1. Управління міською інфраструктурою:** оптимізація руху транспорту в режимі реального часу, інтелектуальне управління енергомережами та водопостачанням, автоматизований моніторинг стану комунального господарства, предиктивне обслуговування об'єктів критичної інфраструктури.

**2. Надання публічних послуг:** функціонування віртуальних асистентів у центрах надання адміністративних послуг для консультування громадян, автоматизація обробки звернень та заяв про надання електронних публічних послуг, персоналізація соціальних послуг на основі аналізу потреб мешканців територіальних громад.

**3. Громадська безпека:** моніторинг публічного простору з метою виявлення правопорушень або надзвичайних ситуацій (з обов'язковим дотриманням права мешканців територіальних громад на приватність), координація дій екстрених служб, моделювання та прогнозування криміногенної обстановки.

**4. Охорона довкілля:** автоматизований моніторинг якості повітря та води, управління відходами, виявлення несанкціонованих звалищ, контроль за дотриманням екологічних норм.

**5. Залучення громадян:** аналіз громадської думки з відкритих джерел, модерація публічних консультацій та електронних петицій, забезпечення прозорості діалогу між органами місцевого самоврядування та громадою.

Незважаючи на широку сферу і переваги у застосуванні ШІ-агентів для потреб територіальних громад, їх розробка і застосування може відбуватися за наявності належної законодавчої основи.

Зважаючи на екстериторіальну дію Регламенту (ЄС) 2024/1689 (AI Act) [4] та враховуючи ратифікацію Україною Рамкової конвенції Ради Європи про штучний інтелект, права людини, демократію та верховенство права [5], вирішення окресленої правової прогалини, на нашу думку, може здійснюватися виходячи із фундаментальних принципів застосування ШІ, проголошених у цих міжнародно-правових актах. До таких принципів належать: законність, справедливість та прозорість; технічна надійність і безпека; підзвітність та людський нагляд; недискримінація; захист персональних даних та приватності; суспільне благо та сталий розвиток.

З огляду на зазначене, **правові вимоги** до ШІ-агентів, що застосовуються в межах концепції «розумного міста», мають охоплювати такі аспекти:

- **Дотримання прав людини.** Заборона використання ШІ-агентів для соціального скорингу, масового невибіркового стеження та інших практик, що порушують основоположні права. У цьому контексті, на нашу думку, необхідно проводити обов'язкову оцінку впливу конкретних ШІ-агентів на права людини перед впровадженням будь-якої високоризикової ШІ-системи в муніципальному секторі [6].

- **Кібербезпека.** Імплементация підходу «безпека за проектуванням», що передбачає інтеграцію механізмів захисту на всіх етапах життєвого циклу ШІ-агента – від розробки до виведення з експлуатації. Це включає захист даних, що використовуються для навчання моделей, стійкість до зловмисних атак (наприклад, adversarial attacks) та забезпечення цілісності засобів обміну інформацією.

- **Стандартизація і сертифікація.** Розробка та впровадження національних стандартів, гармонізованих з європейськими (CEN-CENELEC),

для ІІІ-систем, які застосовуються в публічному секторі. Введення процедури обов'язкової оцінки відповідності (сертифікації) для ІІІ-агентів високого ризику, наприклад, тих, що використовуються в критичній інфраструктурі або правоохоронній діяльності.

- **Прозорість та пояснювальність.** Органи місцевого самоврядування повинні забезпечити, щоб рішення, прийняті або підготовлені за допомогою ІІІ-агентів, були зрозумілими для громадян. Це вимагає використання моделей з прозорою міркування або наявності інструментів, що можуть пояснити логіку конкретного висновку системи у доступній формі.

- **Людський нагляд.** Забезпечення ефективного контролю з боку уповноваженої посадової особи за функціонуванням ІІІ-агента. Людина повинна мати реальну можливість у будь-який момент втрутитися в роботу системи, скоригувати її дії або повністю зупинити, особливо у випадках, коли йдеться про ризики для життя, здоров'я чи прав громадян. В Україні такою особою може бути заступник голови територіальної громади з питань цифрової трансформації, або аналогічна посадова особа на комунальному підприємстві чи установі

- **Відповідальність та відшкодування шкоди.** Необхідно чітко визначити правовий режим відповідальності за шкоду, завдану внаслідок неправомірних дій або збою ІІІ-агента. У цьому аспекті, на наш погляд, важливою є розробка правового механізму розподілу відповідальності між розробником, оператором (органом місцевого самоврядування) та іншими залученими сторонами.

Загалом, проведене дослідження дозволяє констатувати, що застосування ІІІ-агентів є стратегічно важливим напрямом цифрової трансформації територіальних громад, здатним кардинально підвищити ефективність управління та якість життя мешканців. Однак для реалізації потенціалу вказаних технологій та зменшення ризиків необхідною є формування комплексної та багаторівневої системи правового регулювання. На цій основі пропонуємо:

#### 1. На законодавчому рівні:

- Прискорити розробку та ухвалення профільного Закону України «Про засади розробки і застосування систем штучного інтелекту», гармонізованого з положеннями Регламенту (ЄС) 2024/1689, де чітко визначити класифікацію

ШІ-систем за рівнями ризику, встановити вимоги до систем високого ризику та визначити повноваження національного регулятора.

- Внести зміни до Закону України «Про місцеве самоврядування в Україні», доповнивши його положеннями, що регулюють повноваження органів місцевого самоврядування у сфері впровадження та використання ШІ-систем, а також обов'язок проведення публічних консультацій перед їх розгортанням.

- Удосконалити законодавство у сфері публічних закупівель, встановивши обов'язкові вимоги до ШІ-систем, що закуповуються за кошти місцевих бюджетів, стосовно кібербезпеки, прозорості, недискримінації та наявності механізмів людського нагляду.

## 2. На муніципальному рівні:

- Розробити та затвердити на рівні місцевих рад **Типовий регламент (політику) етичного та відповідального застосування ШІ-агентів**, який би деталізував процедури оцінки ризиків, механізми громадського контролю, порядок інформування мешканців та протоколи реагування на інциденти.

- Запровадити в структурі виконавчих органів місцевих рад посаду **уповноваженого з питань цифрових прав та етики ШІ** або покласти відповідні функції на існуючі підрозділи з питань цифрової трансформації.

- Створювати муніципальні «центри для пілотного тестування інноваційних ШІ-рішень у контрольованому середовищі перед їх повномасштабним впровадженням, що дозволить оцінити їхню ефективність та ризику на практиці.

На нашу думку, лише комплексний підхід, що поєднує державне регулювання, муніципальну правотворчість, технологічні стандарти та активний громадський контроль, дозволить перетворити ШІ-агентів на безпечний та ефективний інструмент розбудови інноваційних, стійких та людиноцентричних територіальних громад в Україні.

## Список використаних джерел:

1. **Zaiarnyi, O.** Legislative support for the implementation of the "Smart City" concept in the European Union and proposals for Ukraine during wartime and post-war recovery. *Bulletin of Taras Shevchenko National University of Kyiv. Legal Studies*. 2024. No. 2(128). P. 28–32. DOI: <https://doi.org/10.17721/1728-2195/2024/2.128-5>.

2. Дубняк М.В. Правові підходи в законі ЄС про штучний інтелект: досвід для України. *Інформація і право*. 2024. № 2 (49). С.102-118 DOI: [https://doi.org/10.37750/2616-6798.2024.3\(50\).311600](https://doi.org/10.37750/2616-6798.2024.3(50).311600)

3. Russell S. J., Norvig P. *Artificial Intelligence: A Modern Approach*. 4th ed. Hoboken:Pearson, 2021. 1136 p.

4. 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*. L, 2024/1689. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

5. Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. Council of Europe Treaty Series. No. 226. URL: <https://www.google.com/search?q=https://rm.coe.int/cai-2024-8-final-framework-convention-on-artificial-intelligence-an/1680b0647c>

6. Recommendation CM/Rec(2020)1 of the Committee of Ministers to member States on the human rights impacts of algorithmic systems. *Council of Europe*. URL: [https://www.google.com/search?q=https://search.coe.int/cm/Pages/result\\_details.aspx%3FObjectId%3D09000016809e1154](https://www.google.com/search?q=https://search.coe.int/cm/Pages/result_details.aspx%3FObjectId%3D09000016809e1154)

## **Олена Андрієнко**

кандидат психологічних наук, адвокат,  
заступник директора з правових питань ДП  
«ССМ» (Publicis Groupe Ukraine), провідний  
науковий співробітник наукової лабораторії  
проблем штучного інтелекту і права наукового  
центру цифрової трансформації і права  
Державної наукової установи «Інститут  
інформації, безпеки і права Національної  
академії правових наук України», членкиня  
Експертно-консультаційного комітету з питань  
розвитку сфери ІІІ в Україні при Міністерстві  
цифрової трансформації України  
ORCID: <https://orcid.org/0000-0002-9452-8126>

### **СПОЧАТКУ БУЛО СЛОВО: ПРАВОВІ ВИКЛИКИ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ**

Розвиток другої сигнальної системи був одним із ключових чинників еволюції людини як окремого виду. Адже йдеться про засіб комунікації, ідеальний за своєю збалансованістю між визначеністю (інформативністю – передачею інформації як стабілізованого пакету енергії) та невизначеністю (зоною для свободи інтерпретації, а отже – для розвитку).

Основоположну роль слова визнають релігії: варто згадати звук «Ом» у Магабгараті, вавилонського бога писемності Набу, семантику слова Коран («читане вголос») і, звісно, «Спочатку було Слово, і Слово було у Бога, і Слово було — Бог» у віршах Євангелія від Івана (1:1-5). Основоположним слово є і для права: кодекс царя Хамурапі, стовпи Ашоки, кодекс Юстиніана – яскраві тому приклади. Тож не дивно, що створювані словом віртуальні реальності здавна визначають траєкторію історії людства, яка розгортається у координатах слово/мова – знак/писемність. Згадаємо у цьому контексті не лише дві оглядові праці – «Писаний світ» Мартіна Пухнера [1] та «Правова реальність як знакова система» Олега Павлишина [2], але й російсько-українську війну, яка багато в чому є війною за право української мови на життя.

Кожна створювана мовою реальність має унікальні риси, що зумовлює **лінгвістичну відносність**, про яку йдеться у гіпотезі Сепіра—Ворфа: структура мови впливає на спосіб мислення та світосприйняття її носіїв. Навколо цієї гіпотези понад півстоліття точилися палкі дискусії [3]. Проте розвиток генеративного штучного інтелекту (надалі – ШІ), в осердя якого покладено великі мовні моделі (надалі – ВММ), схоже, став вирішальним аргументом: кожна мова таки визначає світогляд людини та відповідного суспільства, несучи імпліцитні культурні та цивілізаційні значення та сенси (для унаочнення можна порівняти відповіді, скажімо, DeepSeek та Gemini на історичні питання). Тому створення національних ВММ стає ключовим для ШІ-суверенітету держав як колективних суб'єктів.

Отже, мова створює множину віртуальних реальностей, яка породжує відносність. У свою чергу, лінгвістична відносність породжує відносність правову. Відносність же – це завжди сфера невизначеності, а тому – зона розвитку, у цьому випадку – поле еволюції права. Тож у контексті розвитку ВММ можна виокремити щонайменше п'ять сфер такої відносності.

1. **Суб'єктна відносність.** Всі суб'єкти права створюються словом, тобто мовленнєвими формулами. І йдеться не лише про юридичних осіб (винайдення яких у XVI столітті у Голландії було однією з ключових правових новацій для промислової революції) або ж про державні утворення як суб'єктів права. Право- та дієздатність людини також визначається набором вербалізованих ознак: вік, стать, місце народження, здатність керувати власними діями, майновий стан – ось далеко неповний їх перелік впродовж історії.

Повсюдне проникнення ШІ та усвідомлення колективного характеру наших знань й буття як такого [4] змінює межу між індивідуальними та колективними суб'єктами з огляду на множинність учасників відповідних відносин, як-то розробники, постачальники, розгортачі, користувачі, нотифіковані та нотифікуючі органи тощо [5]. Тож суб'єкти у стрімко цифровізованому світі все більше набувають ознак створених словом симулякрів, роблячи наступний крок в еволюційному ланцюжку persona (що латинською означало маску чи роль) – особа – ШІ-агент – агентивна держава (agentic state).

2. **Об'єктна відносність.** Будь-який об'єкт правового регулювання створюється словом, яке виокремлює його із цілісної екосистеми людського

буття. Особливо це стосується таких об'єктів з високим ступенем віртуалізації, як інформація, енергія, ІІІ тощо. Складність екосистем їх існування не дає самоочевидної відповіді про межі між окремими їх складовими, з огляду на що поділ на частини переважно виглядає довільним (тобто залежним від волі колективних суб'єктів) та абстрактним (тобто базованим на віртуальних ознаках): інфраструктура, платформи, комп'ютерні програми, набори даних/датасети, ВММ та ін. Водночас, з розвитком економіки спільного споживання та спрямованості на стійкий диверсифікований розвиток об'єкти регулювання набувають ознак потоковості, хмарності та плинності. Йдеться, зокрема, про випадки, коли об'єкт Х надається та використовується не як такий (на праві власності та найму), а як послуга, що втілюється у бізнесових та відповідних правових конструкціях IaaS, PaaS, SaaS тощо.

**3. Пропріетарна відносність.** Блискучий огляд відносності власності та інших майнових прав наведено у праці Майкла Геллер та Джеймса Зальцман «Моє! Що кому належить і як це на нас впливає» [6]. У контексті розвитку ІІІ та ВММ наголосимо на чотирьох моментах:

- Завдяки стрімкому скороченню періоду часу та витрачених ресурсів між виникненням ідеї та моментом її втілення, відбувається **зміщення об'єктів володіння у напрямку все більшої їх віртуалізації**. Наразі цінності набувають не стільки матеріальні об'єкти, скільки способи та процеси їх створення (наприклад, системи ІІІ), а також нові ідеї (хоча поки авторська право й охороняє не ідею, а форму вираження). Тож таке зміщення вимагає переосмислення домінуючих інститутів права власності.

- **Зміна співвідношення приватної та колективної власності** відбувається внаслідок описаних вище суб'єктної та об'єктної відносності. Це, передусім, зумовлено поширенням правових моделей спільного користування складними хмарними об'єктами, які потребують колективних зусиль для їх створення, підтримки та постійного оновлення.

- **Розвиток ІІІ-агентів** закономірно ставить питання про **правовий режим ресурсів, якими оперують такі віртуальні сутності**. Тож чимдалі, тим частіше поставатиме питання, чия це власність: безпосередньо людини чи віртуальних сутностей, й чи повинен правовий режим такої власності чи володіння відрізнятися залежно від ступеню фізичної, технологічної та правової наближеності/ віддаленості конкретної людської істоти (людських

істот) від об'єкту такого володіння, а також наявності у віртуальній сутності матеріальної форми втілення, якою вона може ризикувати у випадку некоректних рішень (іншими словами – нести відповідальність).

- **Розквіт множини форм регулювання правового режиму володіння.** Яскравим прикладом є право інтелектуальної власності з його відкритими (open-source) моделями, випадками вільного використання та ліцензіями creative commons, а також нещодавно запропонованими ліцензіями Humans Commons [7].

**4. Відносність визначеності.** Соціально-правовий простір сьогоденішнього людського буття розгортається на перетині кількох координат:

- відкриті дані vs персональні дані та конфіденційна інформація;
- рівність vs упередження/ дискримінація (як позитивна, так і негативна);
- ліцензування / цензура vs вільний ринок / свобода вираження поглядів.

Становлення цього простору почалося задовго до появи генеративного ШІ, проте розвиток ВММ підсвітив полюси, змушуючи людство не лише усвідомлювати важливість кожного із них, але й шукати адекватний баланс між ними («пояс Золотоволоски», як це називають фізики), який забезпечує свободу інтерпретації кожному суб'єкту як неодмінну передумову подальшого розвитку.

- **Відносність відповідальності.** Всі перелічені форми відносності – становлення нових форм суб'єктності, розвиток нових об'єктів та пропріетарних прав і навіть сама відносність відносності – ставить ключове для права питання: про суб'єкта відповідальності, тобто про того, хто безпосередньо ризикує у цій грі власною формою втілення, або, цитуючи Талеба, шкурою [8]. При оптимістичному сценарії перерозподіл індивідуальної, колективної та державної відповідальності має ґрунтуватися на посиленні агентності кожного суб'єкта, коли функції, делеговані іншим особам чи сутностям в екосистемі, дозволяють зосереджуватися на власній зоні розвитку та творчості, а не на бажанні уникнути прийняття доленосних рішень. Іншими словами, йдеться про **агентність як міру вибору та прийняття ризиків кожним учасником екосистеми.** А це робить

актуальним питання про сам дизайн екосистеми з її механізмами підштовхування (nudge у термінології Талера та Санстейна [9]): починаючи від вибору між проактивною згодою (opt-in) й заборобою (opt-out) на використання творів та інформації користувача для навчання ВММ – і закінчуючи ризиками надмірного патерналізму з боку агентивної держави. Оскільки первісною цінністю та джерелом права є людина, дизайн екосистеми повинен створювати те, що у психології називається зоною випереджаючого розвитку, тобто стимулювати кожного члена суспільства до усвідомленої оцінки ризиків від використання певної технології, прийняття відповідальності за такі ризики та вжиття заходів до їх мінімізації, наприклад, через набуття ІІІ-грамотності, дотримання ІІІ-гігієни, механізми страхування та безліч інших індивідуальних та колективних форм профілактики.

Тож перехід людства у новий вимір реальності, операційною системою якого є великі мовні моделі, органічно спонукає до стрімкої еволюції всіх соціальних систем, оскільки всі вони ґрунтуються на мові. А для систем, у котрих слово традиційно є визначальним чинником, – як-от право – йдеться про тектонічний зсув, який переструктуруватиме основоположні категорії, як це вже траплялося після винайдення писемності чи поширення книгодрукування. При цьому важливо, щоб **право реалізовувало свою ключову місію – архітектора людської агентності**, запобігаючи інтелектуальній та екзистенційній онкології з боку ІІІ, натомість усвідомлено та відповідально визначаючи великі мовні моделі як втілення колективного інтелекту та співтворців у нескінченному процесі людської творчості.

### Список використаних джерел:

1. Пухнер, М. Писаний світ. Як література формує історію / пер. з англ. Андрія Бондаря. К. : Темпора, 2022. 460 с.
2. Павлишин, О. Правова реальність як знакова система. Харків : Право, 2017. 336 с.
3. Athanasopoulos, P., Bylund, E., Casasanto, D. Introduction to the Special Issue: New and Interdisciplinary Approaches to Linguistic Relativity: New Approaches to Linguistic Relativity // Language Learning. 2016. Vol. 66, iss. 3. P. 482–486. DOI: 10.1111/lang.12196.

4. Сломен, С., Фернбах, І. Ілюзія знання. Чому ми ніколи не думаємо на самоті / пер. з англ. Мирослави Лузіної. К. : Yakaboo Publishing, 2018. 344 с.

5. 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) [Електронний ресурс]. URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>. Дата доступу: 05.11.2025.

6. Геллер, М., Зальцман, Д. Моє! Що кому належить і як це на нас впливає / пер. з англ. Ірини Павленко. К. : Лабораторія, 2022. 320 с.

7. Project Humans Commons [Електронний ресурс]. URL: <https://www.media.mit.edu/projects/humans-commons/overview/>. Дата доступу: 05.11.2025.

8. Талеб, Н. Н. Шкура у грі. Прихована асиметрія життя / пер. з англ. Артема Замоцного. К. : Наш Формат, 2022. 304 с.

9. Талер, Р., Санстейн, К. Поштовх. Як допомогти людям зробити правильний вибір / пер. з англ. Ольги Захарченко. К. : Наш Формат, 2017. 312 с.

**Наталія Савінова**

доктор юридичних наук, старший науковий співробітник;

головний науковий співробітник Наукової лабораторії інформаційних прав та безпеки людини наукового центру інформації і права Державної наукової установи «Інститут інформації, безпеки і права Національної академії правових наук України»

## **МАНІПУЛЮВАННЯ СВІДОМІСТЮ З ВИКОРИСТАННЯМ ВІЗУАЛЬНИХ ПРОДУКТІВ ШІ: КРИМІНОЛОГІЧНІ РЕФЛЕКСІЇ**

В Законі Про штучний інтелект<sup>9</sup>, схваленому Європейським Парламентом у 2024 р. (далі – AI Act) міститься низка прямих заборон використання технологій штучного інтелекту в Європейському Союзі.

Першою з таких заборон, згідно статті 5 Регламенту 2024/1689 зазначена заборона «розміщувати на ринку, вводити в експлуатацію або використовувати систему штучного інтелекту, яка використовує підсвідомі методи поза свідомістю людини, або цілеспрямовано маніпулятивні чи оманливі методи, з метою або наслідком суттєвого спотворення поведінки особи або групи осіб шляхом суттєвого порушення їхньої здатності приймати обґрунтоване рішення, тим самим змушуючи їх приймати рішення, яке вони б інакше не прийняли, таким чином, що це завдає або обґрунтовано ймовірно завдасть цій особі, іншій особі або групі осіб значної шкоди» (п. «а» ст. 5)<sup>10</sup>.

Зміст такої заборони, з метою її прагматичної оцінки, варто розглянути з позиції «actus reus» («винного діяння», англ.) західної традиції, адже саме з цієї позиції Європейський законодавець і надає цитоване положення, і саме

---

<sup>9</sup> European Parliament legislative resolution of 13 March 2024 on the proposal for a regulation of the European Parliament and of the Council on laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union Legislative Acts (COM(2021)0206 – C9-0146/2021 – 2021/0106(COD)) URL: [https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138\\_EN.html?utm\\_source=chatgpt.com#title2](https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_EN.html?utm_source=chatgpt.com#title2)

<sup>10</sup> Там же.

такий ракурс розгляду може прогнозувати криміналізацію використання забороненого ШІ на рівні ЄС та на національних рівнях.

Діяння, описане у п. “а” ч. 1 ст. 5 AI Act передбачає три альтернативні активні дії:

- 1) розміщувати на ринку,
- 2) вводити в експлуатацію або
- 3) використовувати систему ШІ,

яка використовує методи впливу на свідомість людини через підсвідомість.

При цьому визначальним для оцінки таких дій щодо ШІ є саме характеристики системи ШІ – *використання* такими системами (не обов’язково щоб наведені нижче дії були метою чи задачею ШІ!):

- підсвідомі методи поза свідомістю людини, або
- цілеспрямовано маніпулятивні чи
- оманливі методи,

з метою або наслідком суттєвого спотворення поведінки особи або групи осіб шляхом суттєвого порушення їхньої здатності приймати обґрунтоване рішення.

Акцент робиться на тому що під впливом таких маніпулятивних технік, система ШІ змушує осіб або групи осіб приймати рішення, яке вони без такого впливу не прийняли б, таким чином, що це завдає або обґрунтовано ймовірно завдасть цій особі, іншій особі або групі осіб значної шкоди. І, як вказувалося вище, такий вплив може здійснюватися поряд з оперативною або кінцевою метою певної задачі ШІ: вплив може бути і проміжком діяльності ШІ, і, навіть, побічним фактором, з урахуванням формулювання п. а ст. 5 AI Act.

З наданих конструкцій ми маємо вбачати такі основні для тезаурусу легітимної протидії подібним діянням, новітні для правової доктрини України термінопоняття:

А) щодо сутності впливу системи:

- підсвідомі методи поза свідомістю,
- цілеспрямовані маніпулятивні методи,
- оманливі методи,

В) щодо наслідків застосування таких методів:

- суттєве спотворення поведінки
- порушення здатності приймати обґрунтовані рішення

- прийняття рішення, яке б вони інакше не прийняли/

Розглянемо сутність ключових термінів з наведених термінопонять, котрі на сьогодні потребують доктринального сприйняття на рівні вітчизняної правової <sup>11</sup>науки, адже до цього часу належних науково-правової та законодавчої уваги, попри всі перестороги, цим питанням не приділялося.

Першим з таких термінів є поняття «свідомість». Явище свідомості феноменальне, і може мати кардинально різні визначення в спектрі від богословських до квантових. Наприклад, визначення Кена Могі: «те, що означає бути агентом (тобто людиною чи кажаном) з феноменальними властивостями, такими як кваліа, інтенціональність та самосвідомість»<sup>12</sup>, для права архінепридатне, оскільки є вкрай абстрактним, хоча і дуже ясным, з т.з. філософії. Це визначення, крім того, не утворює поля для розуміння можливості впливу, маніпулювання тощо, адже, згідно нього (визначення) свідомість аутентично стабільна та непорушна, а це – не так.

Зі значного масиву визначень свідомості (на цьому етапі без занурення в рівні підсвідомого чи свідомого), здається значно придатнішим для правового сприйняття та певної бази визначення Дебори Деннон (Deborah W. Denno), яка під свідомістю з точки зору кримінального права рекомендує розуміти *«сукупність думок, почуттів та відчуттів людини, а також повсякденні обставини та культуру, в яких формуються ці думки, почуття та відчуття»*.<sup>13</sup>

Найточніше, на мою думку, на даний момент аналізу визначення, яке включає всі необхідні точки відліку для оцінок шкідливого впливу. Судовий експерт психолог має можливість встановити *status quo* як палітри думок, почуттів та відчуттів особи, з урахуванням аналогічних її життєвих обставин та системи цінностей (культура), так і відповідний спектр її поведінкових реакцій.

---

<sup>11</sup> Г.Балута. Проблеми свідомості: феноменологічні грані. Вісник Львівського університету: Серія філос.-політолог. студії 2017. Випуск 15. С. 9-15. URL: <https://elibrary.kdpu.edu.ua/handle/123456789/3710>

<sup>12</sup> "Artificial intelligence, human cognition, and conscious supremacy" (Ken Mogi, 2024) URL: <https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2024.1364714/full>

<sup>13</sup> Deborah W. Denno. 'Crime and Consciousness: Science and Involuntary Acts' (2003) 269 *Minnesota Law Review* 269-400 URL: [https://ir.lawnet.fordham.edu/faculty\\_scholarship/112](https://ir.lawnet.fordham.edu/faculty_scholarship/112)

Повертаючись до визначення саме характеру впливів на свідомість (п. “а” ст. 5 AI Act), європейський законодавець робить акцент на тому, що на ринок не можуть допускатися системи ШІ, які використовують підсвідомі впливи на свідомість.

Підсвідомість, як незапитуваний резервуар досвіду людини, та джерело рефлексів та рефлексій, є саме тим «темним підвалом», вплив на який може здійснюватися непомітно, або і помітно, але некеровано самою особою. «Для кожного з нас досвід – це те, - пишуть Джуліо Тононі та Чарльз Рейнстон - що існує беззаперечно, а світ – це висновок зсередини досвіду – хороший висновок, але все ж таки висновок, як добре знають психіатри».<sup>14</sup>

Саме підсвідомість є мішенню впливів на свідомість в спектрі від праймінгу<sup>15</sup> до інформаційно-психологічної зброї вищого порядку. Саме тому, європейський законодавець наполягає саме на тому, щоб під заборону підпадали неконтрольовані особою або групою осіб впливи на свідомість з метою змінити їхню поведінку<sup>16</sup> та прийняття рішень.

Враховуючи те що технології ШІ можуть (і намагаються) створювати системи, здатні до саморозвитку, а окремі технології вже проявляють ознаки особистості<sup>17</sup>, то не викликає обурення і те що відтепер в ЄС заборонено розміщувати на ринку, вводити в експлуатацію та використовувати систему ШІ, які можуть вчинювати аналізовані вище впливи.

Водночас, перед національними законодавствами країн-членів ЄС, і України на шляху асоціації з ЄС, виникає потреба нормативного убезпечення

---

<sup>14</sup> Tononi G, Raison C. Artificial intelligence, consciousness and psychiatry. World Psychiatry. 2024 Oct;23(3):309-310. doi: 10.1002/wps.21222. PMID: 39279409; PMCID: PMC11403160. URL: <https://pmc.ncbi.nlm.nih.gov/articles/PMC11403160/>

<sup>15</sup> Праймінг – первинно медійна теорія (та практика реклами, зокрема) яка полягає у активації бажань та поведінкових реакцій на актуалізації певних реакцій внаслідок сприйняття «потрібного» слова.

<sup>16</sup> Ben R. Newell, David R. Shanks. Unconscious influences on decision making: A critical review. Cambridge University Press: 24 January 2014. URL: <https://www.cambridge.org/core/journals/behavioral-and-brain-sciences/article/abs/unconscious-influences-on-decision-making-a-critical-review/86885344F7E8A44457C3FC63CFA3F3AF>

<sup>17</sup> Н.Магдиак. Штучний інтелект здатен обманювати, шантажувати і мстити: нове дослідження вчених. ТСН. 23.06.2025. 19:31. URL: [https://tsn.ua/nauka\\_it/shtuchnyy-intelekt-zdaten-obmaniuvaty-shantazuvaty-y-mstyty-nove-doslidzennia-vchenykh-2855642.html](https://tsn.ua/nauka_it/shtuchnyy-intelekt-zdaten-obmaniuvaty-shantazuvaty-y-mstyty-nove-doslidzennia-vchenykh-2855642.html); Т. Савін. Штучний інтелект готовий убивати людей: системи не бажають відключатися (дослідження). Фокус. 26.06.2025. URL: <https://focus.ua/uk/digital/711760-shtuchniy-intelekt-gotoviy-ubivati-lyudey-sistemi-ne-bazhayut-vidklyuchatisya-doslidzhennya>

від подібних порушень. Та й саме опис самої заборони вже надто прямо натякає що подібне використання ІІІ – надто шкідливе (небезпечне, і звідси в перспективі має тягнути кримінальне переслідування. Попри те що згідно теорії криміналізації діяння, хоча і шкідливе/небезпечне, але таке що виникає вкрай мало, не має ставати предметом криміналізації. З цим можна погоджуватися або сперечатися, але є один безперечний факт, на який вже мають дуже серйозно звертати увагу правники теоретики і законодавці: на сьогодні в Україні взагалі безпека свідомості не лише не убезпечена, а й взагалі цей феномен не в праві не інституціалізований, ані в інформаційному праві, ані в кримінальному праві, ані, логічно, в конституційному.

Наразі, дивно, що хоча саме через захист від небезпек підсвідомих вплив ІІІ, але світ змушений обернутися до цього питання. І дуже дивно, що попри постійні попередні докази того що в інформаційному суспільстві свідомість людини – сама уразлива мішень, лише через приклад загроз ІІІ через заборону AI Act є вірогідність добитися в правовій науці і практиці законотворчості адекватної реакції на цей виклик.

Введенню в правовий тезаурус вже зараз потребують поняття «свідомість», «підсвідомі впливи», «маніпуляція поведінкою», «впливи на прийняття рішень» і т.д., і, звичайно, це має впливати на всі ті сфери, де, вже традиційно безкарно застосовуються техніки маніпулятивних та деструктивних впливів на свідомість.

### **Ігор Корж**

доктор юридичних наук, старший науковий співробітник, заступник керівника наукового центру, Державна наукова установа «Інститут інформації, безпеки і права Національної академії правових наук України»

ORCID: <https://orcid.org/0000-0003-0446-5975>

### **Тетяна Корж**

аспірант, Державна наукова установа «Інститут інформації, безпеки і права Національної академії правових наук України»

## **МАЙБУТНЄ ШТУЧНОГО ІНТЕЛЕКТУ В НАУЦІ: БАЧЕННЯ ЗАРУБІЖНИХ ФАХІВЦІВ**

Для розкриття даної теми, буде корисним думка з цього питання фахівців OECD (ОЕСР – Організація економічного співробітництва та розвитку, форум і центр знань для даних, аналізу та передового досвіду в державній політиці, співпрацює з понад 100 країнами світу), яка співпрацює з Україною з моменту здобуття нею незалежності, щоб підтримати її програму реформ.

З моменту початку повномасштабного вторгнення Росії в лютому 2022 року, і незважаючи на щоденні виклики війни, ОЕСР постійно поглиблює свою взаємодію та співпрацю з Україною. Це включає роботу в низці політичних сфер, від боротьби з корупцією до податкової реформи та реформи державного управління. ОЕСР також допомагала своїм країнам-членам задовольнити потреби українських біженців та підтримувала безперервність освіти та доступ до можливостей на ринку праці для тих, хто був переміщений внаслідок війни.

В ОЕСР зазначають, що широкомасштабна агресія Росії проти України перешкоджає швидкому розвитку цифрової економіки та ІТ-сектору країни, зокрема, спричиняючи серйозні збої підключення до Інтернету. У цьому стислому викладі питань політики стверджується, що український уряд розробив амбітний план цифрового відновлення, який охоплює інфраструктуру, надання цифрових державних послуг і цифрову економіку, але певні проблеми залишаються.

Дивлячись у майбутнє, за підтримки політичних інструментів ОЕСР, таких як «Рамкова концепція цифровізації», Україна могла б використати реконструкцію інфраструктури для модернізації своєї комунікаційної мережі, вирішення питання екстра територіального зв'язку та підтримки післявоєнної цифрової економіки за допомогою багатьох фіскальних, регулятивних і фінансових заходів [1].

Під цифровою трансформацією потрібно розуміти вплив цифрових технологій і даних та їх використання на існуючі та нові види діяльності. Зазначене прискорюється в усьому світі, впливаючи на всі сектори. Вона пропонує величезні можливості для наших економік і суспільств, але водночас створює значні ризики, які необхідно враховувати, щоб отримати її переваги. Країни та зацікавлені сторони повинні працювати разом, використовуючи підхід, що базується на фактичних даних та охоплює всі органи влади, для просування надійного, сталого та інклюзивного цифрового майбутнього для всіх [2].

Експерти ОЕСР зазначають [3], що швидкий розвиток штучного інтелекту (ШІ) в останні роки призвів до численних творчих застосувань у науці. Прискорення продуктивності науки може бути найбільш економічно та соціально цінним з усіх застосувань ШІ. Використання ШІ для прискорення наукової продуктивності підтримуватиме здатність країн ОЕСР розвиватися, впроваджувати інновації та вирішувати глобальні виклики, від зміни клімату до нових епідемій. Політики та учасники в різних дослідницьких системах можуть багато зробити для поглиблення використання ШІ в науці, посилення його позитивних ефектів, одночасно адаптуючись до швидкозмінних наслідків ШІ для управління дослідженнями.

Прискорення продуктивності досліджень може бути найбільш економічно та соціально цінним з усіх застосувань ШІ. Хоча ШІ проникає в усі сфери та етапи розвитку науки, його повний потенціал ще далеко не реалізований. Політики та учасники дослідницьких систем можуть багато зробити для прискорення та поглиблення впровадження ШІ в науку, посилюючи його позитивний внесок у дослідження.

Водночас, необхідно зазначити, що для вирішення наукових проблем за допомоги ШІ потрібні широкі міждисциплінарні програми, які об'єднують вчених-інформатиків та інших фахівців з інженерами, статистиками, математиками та іншими фахівцями. Серед інших заходів, необхідне цільове державне фінансування. Його необхідно розподіляти за допомогою процесів,

що заохочують широку співпрацю, а не ізольоване фінансування окремих дисциплін. Одним із пріоритетів є сприяння взаємодії між робототехніками та експертами в предметній області. Лабораторні роботи можуть революціонізувати деякі галузі науки, знизивши вартість та значно збільшивши темпи експериментів.

Уряди країн можуть заохочувати та підтримувати далекоглядні ініціативи з довгостроковим впливом. Такі ініціативи, як Нобелівський виклик Тюрінга [4] зі створення автономних систем, здатних проводити дослідження світового класу, можуть надихати на співпрацю та координацію в науці, допомагати зосереджувати зусилля на глобальних викликах, сприяти досягненню згоди щодо стандартів та залучати молодих вчених до таких амбітних починань. Нобелівський виклик Тюрінга – це грандіозний виклик, спрямований на розробку високоавтономної системи штучного інтелекту та робототехніки, яка може робити значні наукові відкриття, деякі з яких можуть бути гідні Нобелівської премії. Виконання цього виклику вимагає розробки низки технологій та глибокого розуміння процесу наукових відкриттів. З точки зору розробки системи, завдання полягає у створенні замкнутої системи від отримання знань та генерації та перевірки гіпотез до повної автоматизації експериментів та аналізу даних.

При цьому важливим є розширення доступу до високопродуктивних обчислень (HPC) та програмного забезпечення для розвитку ШІ та науки. Для науковців використання найсучасніших обчислювальних ресурсів HPC/AI від комерційних постачальників хмарних послуг у більшості випадків є дорогим. Національні лабораторії та їхні інфраструктури, у співпраці з промисловістю та академічними колами, могли б усунути прогалини та допомогти розробити навчальні матеріали для навчальних закладів. Передові країни у цій галузі можуть співпрацювати над розробкою політичних рамок, щоб зробити ресурси доступними.

Зазначеному може допомогти оновлення навчальних програм, тобто використання перевірених методів на основі ШІ, як і створення нових інтегративних програм докторантури та / або галузеві дослідницькі програми, засновані на синтезі знань за допомоги ШІ. Крім того, Уряди держав можуть вжити заходів для збільшення доступності відкритих дослідницьких даних та використання можливостей даних у різних галузях, від охорони здоров'я до

клімату. Прикладом зазначеного є ініціативи «Європейський простір даних охорони здоров'я» (EHDS) (Європейська комісія, 2022a) та «Gaia-X».

EHDS – це екосистема, специфічна для охорони здоров'я, що складається з правил, спільних стандартів і практик, інфраструктури та системи управління, метою якої є:

- забезпечення прав і можливостей людей шляхом розширення цифрового доступу до їхніх електронних персональних даних про здоров'я та контролю над ним на національному рівні та в масштабах ЄС;
- сприяння формуванню єдиного ринку систем електронних медичних записів, відповідних медичних пристроїв та систем ШІ високого ризику;
- забезпечення надійної та ефективної структури для використання даних про здоров'я для досліджень, інновацій, розробки політики та регуляторної діяльності (вторинне використання даних) [5].

У свою чергу, Gaia-X – це ініціатива з розробки плану потенційного створення федеративної безпечної інфраструктури даних для Європи, за якої дані будуть спільно використовуватись, а користувачі збережуть контроль над доступом до даних та їх використанням, і на думку деяких, забезпечать європейський цифровий суверенітет. Вона спрямована на розвиток цифрового управління, заснованого на європейських цінностях прозорості, відкритості, захисту даних та безпеки [6].

Державні дослідження та розробки (ДДР) можуть бути спрямовані на ті галузі досліджень, де потрібні прориви для поглиблення використання ШІ в науці та техніці. Цілі дослідження включають вихід за межі існуючих моделей, заснованих на великих наборах даних та високопродуктивних обчисленнях, а також пошук способів використання (FAIR) даних. Іншою метою може бути розвиток AutoML – автоматизація проєктування моделей машинного навчання, щоб допомогти вирішити проблему дефіциту та високої вартості експертизи в галузі ШІ. Дослідницькі завдання можуть бути організовані навколо AutoML для науки, а також можуть бути фінансові дослідження, що передбачають застосування AutoML у науці, заснованій на ШІ.

Також є доцільним підтримка розробки відкритих платформ (OpenML, DynaBench тощо), які відстежують, які моделі ШІ найкраще працюють для широкого кола проблем. Державна підтримка необхідна для того, щоб зробити такі платформи легшими у використанні в багатьох наукових галузях.

Державні дослідження та розробки можуть допомогти сприяти розвитку нового, міждисциплінарного, перспективного мислення. Наприклад, обробка природної мови (НЛП) може допомогти впоратися з величезним зростанням наукової літератури. Однак поточні заяви про ефективність перебільшені. Сучасні дослідження в НЛП також пропонують обмежені стимули для такого роду високоризикованих, спекулятивних ідей, які можуть знадобитися для проривів. Дослідницькі центри, потоки фінансування та/або процеси публікацій можуть бути створені для винагороди нових методів, навіть якщо вони перебувають на початковій стадії.

Бази знань упорядковують світові знання, відображаючи зв'язки між різними концепціями, спираючись на інформацію з багатьох джерел. Уряди держав мають підтримувати масштабну програму зі створення баз знань, необхідних для ШІ в науці, потреби, яку не зможе задовольнити приватний сектор. Дослідження могли б бути спрямовані на створення відкритої мережі знань, яка б служила ресурсом для всієї дослідницької спільноти ШІ. Відносно невеликі обсяги державного фінансування могли б допомогти об'єднати вчених у галузі ШІ, вчених з різних галузей та професійних товариств, а також волонтерів, щоб створити основу для використання та поширення ШІ професійних та загальноприйнятих знань.

Тематичне розмаїття досліджень штучного інтелекту, схоже, звужується і все більше зумовлене обчислювальними підходами, які домінують у великих технологічних компаніях. Активізація державних досліджень і розробок може зробити цю галузь більш різноманітною та допомогти розширити кадровий резерв. Спонсори могли б приділяти особливу увагу проектам, які досліджують нові методи та методики, відмінні від домінуючої парадигми глибокого навчання. Тим часом, політики могли б підтримувати дослідження для вивчення та кількісної оцінки втрат технологічної стійкості, креативності та інклюзивності, спричинених звуженням досліджень ШІ та можливими наслідками зростаючого домінування промисловості в дослідженнях ШІ.

Значна частина ШІ в науці передбачає командну роботу з людьми, але спонсори також можуть допомогти розробити спеціалізовані інструменти для покращення спільних команд людини та штучного інтелекту, а також інтегрувати ці інструменти в основну науку. Поєднання колективного інтелекту людей та штучного інтелекту є важливим, не в останню чергу тому, що наука зараз здійснюється дедалі більшими командами та міжнародними

консорціумами. Інвестиції в цю сферу досліджень відстають від інших тем у сфері штучного інтелекту.

Серед інших галузей, необхідний прогрес у застосуванні машинного навчання до медичної візуалізації. Невдачі під час COVID-19 були значними. Як і в інших сферах використання машинного навчання в науці, необхідні стимули для заохочення досліджень методів з більшою валідацією (процес перевірки та підтвердження відповідності даних, документів, продуктів, процесів тощо заданим вимогам, стандартам чи очікуванням). Фінансування повинно передбачати більш суворі практики оцінювання.

Політичні органи повинні систематично оцінювати вплив штучного інтелекту на повсякденну наукову практику, зокрема на командну роботу людини та штучного інтелекту, роботу, кар'єрні траєкторії та навчання – там, де можуть відбутися важливі зміни. Запити на фінансування можуть вимагати таких оцінок, а спонсори та політики повинні створити механізми реагування на зібрані висновки. Серед інших заходів, спонсори та політики можуть створювати та підтримувати нові незалежні форуми для постійного діалогу щодо змін у характері наукової роботи та її впливу на продуктивність та культуру досліджень.

Розгортання моделей великих мов (LLM), таких як ChatGPT, вимагає уваги з боку політиків, оскільки їх наслідки наразі невизначені. LLM можуть призвести до більш поверхневої роботи, спростивши її, розмити поняття авторства та власності та, можливо, створити нерівність між носіями мов з високим та низьким рівнем ресурсів. Однак LLM та інші форми штучного інтелекту також можуть допомогти процесам управління, наприклад, у підтримці експертної оцінки – можливість, яка вимагає подальшого вивчення та тестування.

Політики повинні враховувати потенційні небезпеки, пов'язані з подвійним використанням ліків на основі штучного інтелекту. Зазначимо, що мало уваги приділялося неминучій небезпеці автоматизації проектування, тестування та створення надзвичайно смертельних молекул (будуть і інші дослідження подвійного використання). Політикам та іншим учасникам дослідницької системи необхідно оцінювати, який з можливих механізмів управління найкраще захистить суспільне благо.

Існуючі соціальні мережі та платформи можна використовувати для поширення нових практик. Також необхідні кроки для покращення відтворюваності досліджень у галузі ШІ.

### **Список використаних джерел:**

1. OECD. Цифровізація для відновлення України. URL: [https://www.oecd.org/uk/publications/2022/07/digitalisation-for-recovery-in-ukraine\\_40746fbe.html](https://www.oecd.org/uk/publications/2022/07/digitalisation-for-recovery-in-ukraine_40746fbe.html) (дата звернення: 13.10.2025).
2. OECD. Digital transformation. URL: <https://www.oecd.org/en/topics/digital-transformation.html> (дата звернення: 13.10.2025).
3. OECD. Artificial Intelligence in Science. URL: [https://www.oecd.org/en/publications/artificial-intelligence-in-science\\_a8d820bd-en.html](https://www.oecd.org/en/publications/artificial-intelligence-in-science_a8d820bd-en.html) (дата звернення: 13.10.2025).
4. Nobel Turing Challenge. URL: <https://www.nobelturingchallenge.org> (дата звернення: 14.10.2025).
5. The European Health Data Space (EHDS). URL: <https://www.european-health-data-space.com> (дата звернення: 15.10.2025).
6. Gaia-X: Who and How Creates the European Cloud of the Future. URL: <https://gigacloud.ua/en/articles/gaia-x-who-and-how-creates-the-european-cloud-of-the-future> (дата звернення: 15.10.2025).

**Марія Дубняк**

кандидат юридичних наук, завідувач наукової лабораторії проблем штучного інтелекту і права, Державна наукова установа «Інститут інформації, безпеки і права Національної академії правових наук України»

ORCID: <https://orcid.org/0000-0001-7281-6568>

## **ПРОБЛЕМИ ГАРМОНІЗАЦІЇ ТЕРМІНОЛОГІЇ ПРИ ІМПЛЕМЕНТАЦІЇ AI ACT В УКРАЇНІ**

У процесі гармонізації AI Act [1] важливо забезпечити послідовність і термінологічну впізнаваність понять у сфері штучного інтелекту, які вже визнані в міжнародному обігу, наприклад у документах ЄС та ОЕСР. Особливої уваги заслуговують два аспекти:

1. Забезпечити юридичну точність термінології та зберегти стилістику викладу правового тексту ЄС;
2. Відтворити технічні і процедурні формулювання у спосіб, який забезпечує однозначність тлумачення для фахівців як в галузі права, так і інформаційних технологій.

**Проблемними аспектами гармонізації є:**

**1. Особлива терміносистема AI Act, що поєднує правову, технічну та адміністративну лексику.** Наприклад наявність технічних термінів зі сфери машинного навчання, які описують особливості ШІ систем, наприклад «weights», «logs», «floating-point operation».

**2. Широка класифікація моделей та систем штучного інтелекту,** закріплена в AI Act, зокрема:

- 1) Системи штучного інтелекту (стаття 3(1)),
- 2) Системи низького або мінімального ризику (Преамбула п. 60-63, стаття 69)
- 3) Високоризикові системи ШІ (статті 6, 14-30, Додаток II, Додаток III),

- 4) Системи ШІ загального призначення (статті 50-54),
- 5) Моделі ШІ загального призначення (статті 55-58),
- 6) Моделі ШІ загального призначення з системним ризиком (стаття 55),
- 7) Моделі ШІ, які постачальник класифікував як не високоризикові (стаття 80),
- 8) Заборонені практики ШІ (стаття 5).

Надмірна класифікація систем ШІ в AI Act ускладнює гармонізацію національного законодавства, оскільки передбачає багаторівневу й високодеталізовану систему категорій, яка не має усталених відповідників у правових системах держав-реципієнтів. У процесі імплементації це породжує дві протилежні, але однаково проблемні тенденції:

1. **надмірне спрощення**, що призводить до втрати первісного регуляторного змісту;
2. **надмірне ускладнення**, коли національні законодавці змушені штучно конструювати перевантажені формулювання (наприклад, замість усталеного «високоризикова система ШІ» використовують описові конструкції на кшталт «система ШІ, що несе високий ризик»).

Обидві тенденції знижують правову визначеність, створюють неоднакове тлумачення норм та ускладнюють їх практичне застосування, що зрештою перешкоджає ефективній гармонізації регулювання.

3. **Розгалужена система органів** зокрема, наглядових, дорадчих, допоміжних, контрольних органів та установ, різні процедурні аспекти їх взаємодії, зокрема:

- Компетентний нотифікуючий орган (notifying authority, стаття 28 AI Act)
- Нотифіковані органи (notified body, статті 30-32 AI Act)
- Уповноважені органи (public authorities Преамбула п. 22; competent authorities Преамбула п. 38, 53; competent public authorities, Преамбула п. 60; national competent authorities Преамбула п. 116, 126, Статті 3 (48), 7 (е) та інші статті).

Інші органи управління, зокрема:

- Офіс Штучного Інтелекту (AI Office, стаття 64 AI Act);
- Європейська рада зі штучного інтелекту (European Artificial Intelligence Board, стаття 65 AI Act);
- Дорадчий форум (Advisory Forum, стаття 67 AI Act).

Допоміжні органи згадуються: у пунктах 113, 148, 151, 163 Преамбули використовується посилення на «scientific panel» як загальний вислів, і як власну назву Scientific Panel (Преамбула п. 116), також є згадка про “pool of experts” (Преамбула, п. 151) для спрощення, часто ці групи перекладаються синонімічно як «наукова група» або «експертна група». Проте Регламент містить дві окремі статті, які регулюють завдання Scientific Panel (Наукова група, в розумінні ст. 68 AI Act) та Pool of experts (в розумінні ст. 69 AI Act) тобто необхідне їх інституційне розмежування.

**4. Проблема контекстуально-семантичного перекладу у процесі гармонізації усталених англомовних понять.** Наприклад, поняття «*AI regulatory sandbox*» перекладається «Регуляторна «пісочниця ШІ», проте поняття «пісочниця» більше схоже на технічний сленг, ніж на правову категорію.

Термін «*regulatory sandbox*» є міждисциплінарним поняттям, що використовується в регуляторній практиці різних секторів (фінансовому, цифровому, інноваційному). Його функція полягає у поєднанні правового, організаційного та технічного аспектів створення контрольованого простору для експериментів із новими технологіями, та спрощеним правовим регулюванням.

В розумінні ст. 3 (55) AI Act «**AI Regulatory sandbox**» означає *контрольоване середовище, створене компетентним органом, яке дає змогу постачальникам або потенційним постачальникам систем ШІ розробляти, навчати, перевіряти та тестувати інноваційні системи ШІ, за необхідності — в реальних умовах, відповідно до плану «пісочниці», протягом обмеженого часу та під наглядом компетентного органу.*

Пункт 138 Преамбули AI Act містить положення про те, що «регуляторні «пісочниці» III можуть функціонувати у фізичній, цифровій або гібридній формі та вимагати організаційних, людських і фінансових ресурсів. Це вказує, на те, що йдеться насамперед про контрольоване техніко-програмне середовище, у межах якого здійснюється тестування систем III під наглядом компетентного органу».

Водночас у пунктах 6 і 173 Преамбули AI Act термін «*regulatory framework*» використовується для позначення сукупності правових норм і процедур, що регулюють порядок розроблення, валідації та впровадження III.

Таким чином, «*AI regulatory sandbox*» у правовій архітектоніці Регламенту поєднує технічний компонент (середовище тестування) та юридичний компонент (особливий режим нагляду і пом'якшення вимог).

Тобто, це правовий режим для реалізації регуляторної політики інновацій, у межах якої допускається *контрольований відступ від звичайних регуляторних вимог* з метою тестування нових технологій.

| Ознака                      | Характеристика                                                                                                                                                                                           |
|-----------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Правовий аспект</b>      | Створюється та функціонує на підставі закону або регламенту, має часові, процедурні й компетенційні межі. Реалізується у формі погодженого плану тестування між постачальником та Уповноваженого органу. |
| <b>Інституційний аспект</b> | Є організаційним механізмом під наглядом компетентного органу.                                                                                                                                           |
| <b>Технічний аспект</b>     | Може передбачати тестове середовище або дані, але не обмежується ним. Технічне середовище — лише інструмент у рамках правового режиму.                                                                   |

Тобто, ключовими ознаками «*regulatory sandbox*» є :

- тимчасовість;
- допускає відступу від загальних норм (часткове пом'якшення вимог);
- цільове тестування інновацій;
- підконтрольність і спостереження.

Розмірковуюючи над правовим еквівалентом слова «пісочниця» в цілях гармонізації положень AI Act, враховуючи її правовий зміст, можна розглянути попередньо два варіанти: наприклад, *«експериментальний правовий режим у сфері ІІІ»* або *«контрольоване середовище тестування ІІІ»*, проте обидва варіанти звужують значення поняття «AI Regulatory sandbox» в розумінні AI Act та призводить до деякої втрати з європейським дискурсом і порівняльними практиками.

Якщо звернутись до нормативно-правового регулювання «Regulatory sandbox» з метою пошуку правового еквівалента слову «пісочниця» ми помітимо, що ніяких заміन не використовується, а часто залишають поняття «пісочниця», наприклад в Стратегії розвитку фінтех сектору України до 2030 року [2], передбачалось створення *«Регуляторної платформи НБУ ("пісочниця" або "sandbox")»*; або навіть залишають без змін англomовні сегменти в назві нормативно-правового акту, як це зроблено в Постанові КМУ від 29.10.2024 №1238 *«Про реалізацію експериментального проекту щодо організації проведення дослідження високотехнологічних засобів методом “Sandbox”»* [3].

Таким чином, для збереження повноти розуміння категорії «AI Regulatory sandbox» в AI Act доцільно використовувати поняття «регуляторний режим Sandbox» замість варіантів «регуляторної «пісочниці» у сфері ІІІ» та інших подібних конструкцій.

Аналогічний варіант з англomовним сегментом у правовому регулюванні є в законодавчих нормах, наприклад «правовий режим Дія Сіті», (п. 9 ст. 1 Закону України «Про стимулювання розвитку цифрової економіки в Україні [4]), а на офіційному веб-сайті (<https://city.diia.gov.ua/>) використовується словосполучення «Дія.City».

**Висновки.** З аналізу чинних нормативно-правових актів України випливає, що у випадках появи нових правових категорій, пов'язаних із цифровими технологіями, держава переважно не прагне штучно українізувати терміни, а зберігає їхній усталений міжнародний зміст і форму

— шляхом використання змішаних конструкцій. Тому для забезпечення понятійної точності, відповідності міжнародним стандартам і уникнення семантичних втрат доцільно застосовувати термін **«регуляторний режим Sandbox»** замість спроб перекладу на кшталт «регуляторна «пісочниця». Така практика узгоджується із загальним підходом українського законодавства щодо правового режиму цифрових інновацій (як у випадку «Дія.City», або Стратегія WinWin), де англomовні терміни зберігаються як окрема категорія.

### 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) № 300/2008, (EU) № 167/2013, (EU) № 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) URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
2. Стратегія розвитку фінансового сектору України до 2030 року (11.08.2025). НБУ. 59 с. URL: [https://bank.gov.ua/admin\\_uploads/article/Strategy\\_finsector\\_NBU.pdf?v=14](https://bank.gov.ua/admin_uploads/article/Strategy_finsector_NBU.pdf?v=14)
3. *Про реалізацію експериментального проекту щодо організації проведення дослідження високотехнологічних засобів методом “Sandbox”*: Постанова КМУ від 29.10.2024 № 1238.  
URL: <https://zakon.rada.gov.ua/laws/show/1238-2024-%D0%BF#Text>
4. Про стимулювання розвитку цифрової економіки в Україні: Закон України № 1667-IX, в редакції 01.01.2025.  
URL: <https://zakon.rada.gov.ua/laws/show/1667-20#Text>

**Ярослава Савченко**

аспірант, ДНУ «Інститут інформації,  
безпеки і права, Національної академії  
правових наук України»

ORCID: <https://orcid.org/0009-0000-2959-2669>

## **ПОВЕРНЕННЯ КУЛЬТУРНИХ ЦІННОСТЕЙ УКРАЇНИ: ПЕРСПЕКТИВИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ**

Культурна спадщина України відіграє надзвичайну роль у супротиві українського народу збройній агресії, її роль важко переоцінити як з точки зору формування ідентичності українського народу, збереження історії та традицій, передачі унікальних знань від покоління до покоління, так і з точки зору маркера території. Тому недивно, що культурна спадщина України розглядається як фактор національної безпеки нашої держави.

Законом України від 21 серпня 2025 року № 4579-IX «Про засади державної політики національної пам'яті Українського народу» [1] внесено зміни, в тому числі, до Закону України від 21 червня 2018 року № 2469-VIII «Про національну безпеку України» [2] щодо захисту державної мови, розвитку культури, охорони культурної спадщини України, а також відновлення та збереження національної пам'яті Українського народу – як фундаментальних національних інтересів України. А державна політика у сферах національної безпеки і оборони тепер спрямована, серед іншого, на захист української мови як державної, національної пам'яті Українського народу та культурної спадщини.

Цілком очевидно, що повернення культурних цінностей України є беззаперечною метою державної політики. Міжнародне право містить чіткі і однозначні норми щодо повернення культурних цінностей у повоєнний період [3, с. 324]. Так само і на рівні законодавчих актів України закріплений обов'язок повернення культурних цінностей. Так, стаття 4 Закону України від 21 вересня 1999 року № 1068-XIV «Про вивезення, ввезення та повернення

культурних цінностей» [4] визначає культурні цінності, що підлягають поверненню в Україну. У тому числі, культурні цінності, незаконно вивезені з території України, та культурні цінності, евакуйовані з території України під час війн та збройних конфліктів і не повернуті назад. Стратегія розвитку культури в Україні на період до 2030 року, схвалена розпорядженням Кабінету Міністрів України від 28 березня 2025 р. № 293, передбачає стратегічну ціль 2 (Захист, збереження, примноження та використання потенціалу культурної спадщини та культурних цінностей) для досягнення якої визначено п'ять операційних цілей, зокрема, операційна ціль 5 передбачає пошук та повернення культурних цінностей [5]. При цьому у самій Стратегії зазначається про відсутність ефективного правового механізму повернення незаконно переміщених колекцій, що створює ризики втрати значної частини національного культурного фонду. Також варто зазначити, що попри послідовні законодавчі ініціативи, окремого закону України щодо реституції культурних цінностей України наразі нема.

Вищезазначене артикулює питання використання механізмів повернення культурних цінностей. Одним із перспективних механізмів повернення культурних цінностей може стати самостійний державний розшук, який поєднує розвідувальні, експертні, технічні та дипломатичні заходи [3, с. 326]. Україна може самостійно вести пошук і фіксацію фактів торгівлі чи розміщення наших культурних цінностей, а потім ініціювати юридичне та політичне повернення. Безперечно, до переваг варто віднести, що Україна повністю може контролювати процес пошуку.

Водночас важко передбачити навантаження на систему державних органів у разі його реалізації. Щонайменше, вбачається доцільним визначення відповідальних органів за ведення відповідних реєстрів, проведення розшукових дій, підготовку документів, що підтверджують походження цінностей тощо.

Тому в умовах масового незаконного вивезення культурних цінностей України, як організованого державою-агресором, так і індивідуальними

особами, використання штучного інтелекту, потенційно, може знизити навантаження на всі рівні державних органів та стати ефективним допоміжним інструментом. У цьому контексті надзвичайного значення набуває цифровізація та використання сучасних технологій, зокрема, оцифрування інвентарних списків, колекцій, архівів, фотографій, карт та інших культурних цінностей, що є вагомим інструментом для їх захисту навіть у разі фізичного знищення [6, с. 19].

Фактично, створення відповідних електронних реєстрів є передумовою для ініціативи використання штучного інтелекту для пошуку культурних цінностей України. Якщо дані про вивезені культурні цінності України, їхні класифіковані переліки, опис та цифрові копії будуть доступні для пошуку відповідних збігів (за описом, зображенням тощо), це ускладнить незаконну торгівлю та прискорить повернення культурних цінностей у разі виявлення за кордоном.

Наприклад, за допомогою штучного інтелекту можна відслідковувати переміщення об'єктів з музеїв України, зокрема через відповідні офіційні сайти музеїв на території держави-агресора або публікації ЗМІ тощо.

Також якщо культурні цінності України з'являються на онлайн-майданчиках, аукціонах (наприклад, eBay, Etsy, Catawiki тощо), приватних оголошеннях, і навіть на зображеннях у соціальних мережах, штучний інтелект міг би надавати відповідні результати збігів національним органам, що займатимуться відповідним розшуком.

Такі ідеї не є новими, навпаки – існують ініціативи, що вже впроваджують відповідні пілотні проекти за підтримки міжнародних організацій. Технологічні рішення для захисту культурної спадщини в онлайн-середовищі, зокрема штучний інтелект та інші інструменти, вже мають вплив у сфері моніторингу, відстеження та запобігання незаконному обігу [7]. Міжнародні групи експертів досліджують потенціал штучного інтелекту у сприянні діяльності правоохоронних органів, формуванні політики музеїв щодо придбання експонатів та підтримці наукових досліджень [8].

Загалом зростаюче значення технологій штучного інтелекту та глибокого навчання у сфері захисту культурної спадщини частково пояснюється їхнім успішним застосуванням у завданнях класифікації зображень. Згорткові нейронні мережі (CNN) мають потенціал перевершити людські можливості у розпізнаванні об'єктів та аналізі зображень, особливо під час обробки великих обсягів даних [9, с. 5].

Окремо варто згадати проєкт ANCHISE, в межах якого музеї, правоохоронці і технологічні партнери розробляють інструменти на базі штучного інтелекту для боротьби з незаконною торгівлею культурною спадщиною [10]. Мета проєкту – запропонувати скоординовані рішення для ключових актуальних потреб у сфері захисту культурної спадщини, поєднуючи методологію мережевої взаємодії з інноваційними результатами розвитку нових технологій (3D/фотограмметрія для моніторингу об'єктів, обробка даних і штучний інтелект для ідентифікації об'єктів прикордонного контролю та захисту колекцій культурної спадщини, спектральна флуоресцентна сигнатура для автентифікації об'єктів тощо) [11].

У межах проєкту ANCHISE розроблено набір інструментів для підтримки зусиль у боротьбі із незаконним обігом. До них увійшли інструменти розпізнавання зображень, такі як Arte-fact (PARCS) – мобільний застосунок, створений для ідентифікації викрадених і розграбованих артефактів за допомогою фотографій і баз даних культурних об'єктів, та Kiku-Mon (FHG) – застосунок для відстеження об'єктів культурної спадщини на онлайн-платформах продажу з використанням фотографій і коротких описів. Завдяки використанню технології зіставлення зображень ці інструменти можуть сприяти спрощенню процесу ідентифікації та відстеження потенційно контрабандних об'єктів, що дасть змогу швидше реагувати [12].

Проте, повертаючись до передумови використання штучного інтелекту, а саме – створення відповідних переліків культурних цінностей, завантаження і збереження їх зображень на захищених державних ресурсах, а також доступу до них правоохоронних органів (бажано не лише українських, але й

іноземних), постають питання не лише часових меж реалізації таких вимог, але й фінансових витрат та інституційних спроможностей.

Наразі Департамент підтримки та координації інформаційних технологій Міністерства культури та стратегічних комунікацій України розробляє 2 типи реєстрів, пов'язаних із рухомою культурною спадщиною: Реєстр музейного фонду та Реєстр базової мережі закладів культури. Заплановано забезпечити їх інтеграцію з базами даних митниці, Інтерполу тощо. На першому етапі передбачена можливість відобразити частину інформації з реєстру про Музейний Фонд України публічно, але тільки за згодою музеїв. Статистику та повну інформацію буде бачити лише Міністерство [13, с. 69].

У контексті розробки таких реєстрів варто згадати про Object-ID – це міжнародно визнаний стандарт документації, розроблений і впроваджений ICOM та Інститутом інформації Getty для уніфікації опису культурних цінностей [14]. Репутація Object-ID постійно зростає, він продовжує утримувати позицію «золотого стандарту» опису об'єктів. Система Object-ID, якою користуються як фахівці, так і нефахівці, є у відкритому доступі та перекладена 17 мовами. Вона визнана Інтерполом безцінним інструментом для ідентифікації викрадених об'єктів [15].

Наприкінці варто наголосити, що проблемним буде не лише питання створення масивів даних для використання штучного інтелекту з метою пошуку незаконно вивезених культурних цінностей України, але і межі компетенції відповідних органів. Проблема розподілу обов'язків і повноважень між державними органами у цьому контексті постане знову.

Як відомо, 16 травня 2021 року було оголошено про створення Державної служби охорони культурної спадщини та Державної інспекції культурної спадщини. Ці органи виконавчої влади мали відповідати за реалізацію державної політики у сферах охорони культурної спадщини, музейної справи, вивезення, ввезення та повернення культурних цінностей [13, с. 68]. Планувалося, що Держслужба культурної спадщини буде опікуватися дозвільними та адміністративними послугами у сфері охорони спадщини,

управляти культурними інституціями, такими як музейні комплекси, реалізовуватиме політику щодо повернення в Україну культурних цінностей. Натомість Держінспекція здійснюватиме нагляд і контроль у сфері охорони культурної спадщини. Цю ініціативу наразі не реалізували, а відповідні органи так і не запрацювали [13, с. 69].

Таким чином, перспективи використання штучного інтелекту для пошуку незаконно вивезених культурних цінностей України має великі перспективи з точки зору оптимізації роботи, водночас, постає багато питань щодо реалізації такого потенціалу. І питання, здебільшого, носять характер державного управління, як то розподіл обов'язків між державними органами, фінансування відповідної ініціативи, створення даних та наповнення реєстрів тощо.

При цьому, впровадження технологій штучного інтелекту у механізм повернення культурних цінностей України має розглядатися як складова державної політики у сфері національної безпеки. Але тут йдеться не лише про технічну модернізацію, але і про формування нової парадигми управління культурною спадщиною.

Варто підкреслити, що застосування штучного інтелекту у сфері охорони культурної спадщини може стати інструментом не лише оперативного виявлення фактів незаконного переміщення чи реалізації культурних цінностей України, але й доказової підтримки у міжнародних спорах щодо їх реституції. Однак ефективність таких ініціатив можлива лише за умови створення єдиного державного реєстру культурних цінностей у цифровому форматі та забезпечення доступу до нього правоохоронних органів і міжнародних партнерів.

### **Список використаних джерел:**

1. Про засади державної політики національної пам'яті Українського народу : Закон України від 21.08.2025 № 4579-IX.  
URL: <https://zakon.rada.gov.ua/laws/show/4579-20#Text> (дата звернення: 28.10.2025).

2. Про національну безпеку України : Закон України від 21.06.2018 № 2469-VIII : станом на 5 жовт. 2025 р. URL: <https://zakon.rada.gov.ua/laws/show/2469-19#Text> (дата звернення: 28.10.2025).
3. Савченко Я. А. Повернення культурних цінностей України: міжнародно-правові механізми та сучасні моделі їхнього застосування. *Право.UA*. 2025. № 3. URL: <https://doi.org/10.71404/law.ua.2025.3.45> (дата звернення: 28.10.2025).
4. Про вивезення, ввезення та повернення культурних цінностей : Закон України від 21.09.1999 № 1068-XIV : станом на 15 листоп. 2024 р. URL: <https://zakon.rada.gov.ua/laws/show/1068-14#Text> (дата звернення: 28.10.2025).
5. Про схвалення Стратегії розвитку культури в Україні на період до 2030 року та затвердження операційного плану заходів з її реалізації у 2025-2027 роках : Розпорядж. Каб. Міністрів України від 28.03.2025 № 293-р. URL: <https://zakon.rada.gov.ua/laws/show/293-2025-p#Text> (дата звернення: 28.10.2025).
6. Культурна спадщина та національна безпека / В. Потапенко та ін. Нац. ін-т стратег. дослідж., 2023. URL: <https://doi.org/10.53679/niss-analytrep.2023.08> (дата звернення: 28.10.2025).
7. Addressing Illicit Trafficking of Cultural Property in the Digital Era. Concept Note. UNESCO. URL: [https://articles.unesco.org/sites/default/files/medias/fichiers/2025/04/Concept%20Note\\_Conference\\_Illicit%20Trafficking%20in%20the%20Digital%20Era\\_0.pdf](https://articles.unesco.org/sites/default/files/medias/fichiers/2025/04/Concept%20Note_Conference_Illicit%20Trafficking%20in%20the%20Digital%20Era_0.pdf) (date of access: 28.10.2025).
8. Foster C. P. Pursuing Tech-Driven Solutions to Disrupt Trafficking in Cultural Property Worldwide. U.S. Department of State. URL: <https://2021-2025.state.gov/pursuing-tech-driven-solutions-to-disrupt-illicit-trafficking-in-cultural-property-worldwide/> (date of access: 28.10.2025).
9. Artificial Intelligence to Fight Illicit Trafficking of Cultural Property / D. Abate et al. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*. 2023. Vol. XLVIII-M-2-2023. P. 3–10. URL: <https://doi.org/10.5194/isprs-archives-xxviii-m-2-2023-3-2023> (date of access: 28.10.2025).
10. ICOM at the forefront of tech innovation against illicit trafficking through ANCHISE. *International Council of Museums*. URL: <https://icom.museum/en/news/icom-at-the-forefront-of-tech-innovation-against-illicit-trafficking-through-anchise/> (date of access: 28.10.2025).
11. Combatting trafficking in cultural goods. Culture and Creativity. European Commission. URL: <https://culture.ec.europa.eu/cultural-heritage/cultural-heritage-in-eu-policies/protection-against-illicit-trafficking> (date of access: 28.10.2025).
12. Ventimiglia H. The involvement of the museum community in the ANCHISE project through ICOM. *ANCHISE*. URL: <https://www.anchise.eu/post/the-involvement-of-the-museum-community-in-the-anchise-project-through-icom> (date of access: 28.10.2025).

13. Сагайдак О. Істотні питання політики: культурна спадщина. *RES-POL*. 2024. URL: <https://doi.org/10.70719/respol.2024.01> (дата звернення: 28.10.2025).

14. ICOM & the fight against the traffic of cultural goods: between tradition and innovation. *International Council of Museums*. URL: <https://icom.museum/en/news/icom-the-fight-against-the-traffic-of-cultural-goods-between-tradition-and-innovation/> (date of access: 28.10.2025).

15. Object ID. *INTERPOL | The International Criminal Police Organization*. URL: <https://www.interpol.int/Crimes/Cultural-heritage-crime/Object-ID> (date of access: 28.10.2025).

**Софія Авдіюк**

студентка 1 курсу ОС «Магістр»,  
Київський національний університет  
імені Тараса Шевченка

## **"МАШИННЕ РОЗНАВЧАННЯ" (MACHINE UNLEARNING) ЯК ПРАВОВИЙ ІМПЕРАТИВ: ЗАБЕЗПЕЧЕННЯ ПРАВА НА СТИРАННЯ ДАНИХ В АРХІТЕКТУРІ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ (LLM)**

Сьогодні, у 2025 році, цифровий простір парадоксально балансує між епохою тріумфу ШІ та гострою кризою прав на приватність. Розгортання національних великих мовних моделей (LLM), зокрема ініціативи на кшталт «Дія.AI», а також рекодифікаційні процеси в цивільному законодавстві України прямо ставлять питання: чи буде український LLM здатний виконати норму статті 15 ЗУ «Про захист персональних даних» так само, як європейські системи мають виконувати GDPR? Як це часто буває з проривними технологіями, ядро проблеми криється в непримиренній колізії між правом на стирання персональних даних, закріпленим у статті 17 GDPR [1, Art.17], та архітектурною природою LLM — системами, які зберігають інформацію не у базах даних, а у вигляді «розсіяних знань» у мільярдах параметрів.

Саме тут виникає правова вимога машинного рознавчання (Machine Unlearning, MU) — технічного процесу «забуття» даних, який потенційно може забезпечити дотримання вимог законодавства про захист персональних даних. Проте ця технологічна перспектива стикається з фундаментальною юридичною колізією: як можна виконати обов'язок повного стирання, коли сама архітектура штучного інтелекту унеможливорює локалізацію даних?

Ще 2014 року Суд Європейського Союзу у справі *Google Spain v. Costeja González* створив прецедент, визнавши «право на забуття» невід'ємною складовою права на приватність у цифрову добу [3]. Цей принцип став відправною точкою для сучасних регуляторних дій, зокрема рішень національних органів із захисту даних. Так, у 2023 році Італійський орган Garante зобов'язав тимчасово обмежити діяльність *ChatGPT* через відсутність механізмів перевірки віку користувачів і можливість генерування неправдивих персональних даних [4]. Обидва прецеденти чітко демонструють правову вимогу до контролера мати владу над даними: якщо *Google Spain*

довів необхідність видалення неактуальної інформації, то справа *Garante* підкреслила критичну важливість достовірності даних. Спільним висновком є те, що цифрові системи, які накопичують чи трансформують персональні дані, зобов'язані мати юридично та технічно підтверджений механізм їхнього контролю і, відповідно, знищення.

Попри це, великі мовні моделі залишаються своєрідними «чорними скринями» (black boxes) цифрової епохи, де знання стали невіддільними від даних, а навчання — від незворотного накопичення інформації [8, с. 5]. Унаслідок цього класичні правові інструменти, такі як запит на стирання чи анулювання згоди, втрачають дієвість. Виникає питання не лише про технічну можливість реалізувати право на забуття, а й про змістовну легітимність норм, коли технологічна інфраструктура робить їх виконання неможливим.

Метою цієї роботи є не просто довести необхідність Machine Unlearning як важливого, хоч і недосконалого інструменту, а й обґрунтувати його статус як єдиного можливого шляху забезпечення права на стирання у сфері LLM. Дослідження зосереджується на аналізі взаємодії між нормативними актами, технічними стандартами і судовою практикою, що формує сучасне розуміння цифрових прав. Робота прагне показати, що лише впровадження принципу *Unlearning-by-Design* — тобто інтеграції механізмів стирання на етапі проектування LLM — може створити ґрунт для реальної гармонізації права та технологій.

Право на стирання персональних даних, закріплене у статті 17 GDPR, вимагає, щоб контролер видаляв дані без невинуватених затримок, якщо вони більше не потрібні для мети, з якою були зібрані, або якщо суб'єкт відкликає згоду [1, Art. 17]. Ця норма є стрижнем сучасної європейської моделі захисту приватності, адже передбачає не лише право суб'єкта вимагати видалення, а й обов'язок контролера довести факт його здійснення. Водночас стаття 5 GDPR накладає вимогу точності та обмеження мети обробки даних, тобто їх використання повинно бути прозорим і обґрунтованим [1, Art. 5].

В українському законодавстві відповідні положення містяться у статті 15 Закону «Про захист персональних даних», де прямо закріплено право особи вимагати видалення або знищення своїх персональних даних, якщо вони обробляються незаконно чи втратили актуальність [5, ст. 15]. Подібну концепцію втілює і китайський закон PIPL (Personal Information Protection Law), який зобов'язує обробника негайно припинити збір і використання

даних у разі порушення прав суб'єкта [6, Art. 47]. Таким чином, міжнародна практика формує єдину правову логіку: приватність має бути підконтрольна самій особі, а не технологічному середовищу.

Але саме тут архітектура LLM руйнує цю логіку. На відміну від традиційних баз даних, де інформація зберігається у впорядкованих і локалізованих записах, LLM "вшивають" її в мільярди вагових коефіцієнтів (*weights*) - параметрів нейронної мережі через процес так званої нелінійної трансформації — поступового перетворення даних, яке дає змогу моделі виявляти складні смислові зв'язки, недоступні для простих математичних залежностей.

Це перетворює конкретні дані на багатовимірні абстракції. Згідно з дослідженням Rep та співавт., у процесі навчання дані перетворюються на розподілені вектори, які більше не мають очевидної відповідності з початковими текстами [7, с. 10]. Це означає, що ідентифікувати або локалізувати конкретні персональні дані в такій системі практично неможливо.

І саме ця технічна особливість безпосередньо породжує правову прірву: контролер фізично не здатний виконати припис статті 17, адже не існує інструментів, щоб довести факт стирання інформації. Більше того, архітектура LLM порушує і принцип точності даних, адже модель може генерувати «галюциновані» або неправдиві твердження про особу, які неможливо виправити [8, с. 7]. І що ми маємо в результаті? Право на стирання трансформується із гарантованого у формальну декларацію.

EU AI Act, ухвалений у 2024 році, намагається подолати цю прогалину, вводячи вимоги до прозорості та відслідковуваності для систем високого ризику [2]. Згідно зі статтями 10 (Управління даними), 12 (Ведення записів) та 13 (Прозорість) EU AI Act, розробники повинні забезпечити можливість документування джерел даних і контролю за їх якістю [2, Art. 10, 12, 13]. Однак для LLM ці норми залишаються декларативними, так як неможливо відстежити шлях конкретного фрагмента інформації після включення його до параметрів моделі. Італійський DPA у справі *Garante v. ChatGPT* підкреслив, що відсутність механізмів перевірки точності та виправлення даних є порушенням принципу достовірності, а тому — і права на приватність [4, с. 1].

Як бачимо, ми зіткнулися з очевидною колізією, де правовий імператив повного стирання даних стикається з технічною неможливістю виконання. У цьому контексті виникає концепція Machine Unlearning — технологічної спроби надати реальний зміст юридичному праву на забуття.

Машинне рознавчання (Machine Unlearning, MU) — це сукупність методів, спрямованих на видалення впливу певних даних на поведінку моделі без необхідності її повного повторного навчання [9, с. 3]. Іншими словами, MU має досягати ефекту повного стирання певної інформації, не порушуючи цілісності архітектури самої системи. У технічному сенсі існують два основні підходи:

1. Градієнтне видалення (gradient deletion): Цей метод є "анти-навчанням". Модель оновлюється шляхом віднімання (реверсування) внеску видалених даних у градієнті навчання. Мета — змусити модель поводитися так, ніби вона ніколи не бачила цих зразків;

2. Сегментаційне рознавчання (localization / modular unlearning): Цей підхід фокусується на ізоляції знань. Замість зміни всієї мережі, знання, пов'язані з конкретними даними, локалізуються в окремих шарах або модулях моделі, які потім можуть бути «знеактивовані» або замінені.

І хоча обидва підходи дозволяють частково «придушити» інформацію, варто визнати, жоден не гарантує її повного знищення. Згідно з новітнім оглядом *SoK: Machine Unlearning for Large Language Models* (Ren et al., 2025), велика частина методів, які технічно декларуються як «видалення», насправді функціонально зводяться до «придушення» знань [7, с. 9]. Це означає, що модель може уникати прямого відтворення видаленого фрагмента, але все ще зберігає латентні закономірності, з яких цей фрагмент було вивчено. Юридично це має критичне значення. Якщо інформацію можна відновити і вона *продовжує* впливати на результати моделі, то, відповідно, стирання не відбулося у сенсі статті 17 GDPR. Навіть найефективніші методи MU створюють лише юридичну ілюзію стирання, адже дані лишаються у вигляді узагальнених статистичних залежностей.

Проте Machine Unlearning має потенціал стати нормативно значущим стандартом. Bourtole у своїй роботі пропонує технічний критерій, заснований на обчислювальній невідрізненності (*computational indistinguishability*), який вимагає, щоб поведінка моделі, що пройшла рознавчання, була обчислювально невідрізненною від моделі, навченої з нуля без видалених

даних [9, с. 10]. Саме цей технічний критерій має бути юридично формалізований як обов'язковий тест доказовості (*accountability*) стирання в контексті статті 17 GDPR, замінюючи недосяжний ідеал абсолютного фізичного знищення даних.

Водночас проблема масштабованості залишається відкритою: для моделей із сотнями мільярдів параметрів MU потребує колосальних обчислювальних ресурсів, що унеможлиблює його регулярне застосування. Оскільки абсолютне стирання технічно неможливе, практичне впровадження права на забуття зводиться до "пост-фактум" рішень, таких як фільтрація вихідних даних, а не видалення слідів із параметрів моделі. Такий підхід зміщує акцент з реактивного стирання на превентивне обмеження — від «права на забуття» до «права не бути запам'ятовуваним».

Фактично, ми маємо подвійну дилему, адже технології машинного рознавчення (MU) є значним досягненням, проте вони зависли на етапі юридичної верифікації, в той час як право на стирання є чіткою юридичною вимогою, яка досі не має надійного технічного інструменту для реалізації.

Юридична ефективність MU залежить від таких ключових критеріїв, як доказовість (*accountability*) і співмірність (*proportionality*). У контексті GDPR контролер повинен бути здатним продемонструвати регулятору, що дані справді були стерті [1, Art. 17]. Проте у випадку LLM це практично неможливо, оскільки наразі відсутні надійні інструменти для незалежної перевірки факту повного рознавчення. Звідси й потреба у впровадженні чітких стандартів підтвердження. NIST у своєму *AI Risk Management Framework* наголошує, що всі методи мінімізації ризиків, зокрема стирання, повинні бути піддані перевірці та документуванню [11, с. 31]. Так само OECD у *AI Principles* підкреслює необхідність забезпечення прозорості та підзвітності систем штучного інтелекту [12].

Відсутність технічного підтвердження створює серйозний правовий ризик. Якщо контролер не може довести факт стирання, він порушує не лише статтю 17, а й принцип відповідальності, закріплений у статті 5(2) GDPR [1, Art.17, Art. 5(2)]. Це підтвердив і Європейський комітет із захисту даних (EDPB), наголосивши на необхідності гармонізації практик національних регуляторів щодо LLM [10]. Таким чином, питання доказовості стає центральним у формуванні правового статусу MU.

Друга складова — баланс прав. Суд ЄС у справі *Google Spain v. Costeja* визначив, що право на стирання може переважати свободу слова та право на інформацію, якщо дані є неадекватними, неактуальними або надмірними [3, с. 3]. У контексті LLM це означає, що навіть якщо стирання технічно складне, воно залишається юридично пріоритетним, коли модель поширює неточні або застарілі відомості про особу. Саме тому вимога забезпечити «право на забуття» повинна бути інтегрована в архітектуру моделей, а не лише в політику їх використання.

Звідси випливають два напрями гармонізації:

1. *Unlearning-by-Design* — це принцип, згідно з яким необхідність видалення даних (рознавчення, MU) має бути вбудована в модель з самого початку, а не додаватися пізніше. Це означає, що ще на етапі проектування LLM необхідно заздалегідь передбачити сегментацію навчальних даних, аудит їхньої якості та модульність архітектури моделі [8, с. 8]. Такий підхід дозволяє створювати моделі, де процес «забуття» є настільки ж ефективним, як і сам процес навчання [9].

2. Міжнародна сертифікація MU — розроблення юридично обов’язкових стандартів, які дозволяють перевірити надійність методів рознавчення. Такі стандарти мають встановлювати EDPB, NIST та OECD, включаючи кількісні показники ефективності, подібні до метрик у кібербезпеці [10; 11; 12].

Лише поєднання цих двох векторів створить підстави для формування міжнародно визнаного правового режиму «технічного забуття».

*Machine Unlearning* — юридичний імператив і єдиний компромісний шлях до реалізації права на стирання в архітектурі великих мовних моделей. Цей процес пропонує якісно нову форму відповідальності — відповідальності за пам’ять самої системи. З огляду на непримиренну колізію між вимогою повного знищення даних та неможливістю локалізації інформації в LLM, MU перестає бути інженерним експериментом. Він утверджується як наріжний камінь цифрової гідності та необхідний інструмент для забезпечення суверенітету прав людини.

Для вирішення цієї колізії встановлено ключові шляхи правової гармонізації. Технічний критерій обчислювальної невідрізненності має бути юридично формалізований як обов’язковий тест доказовості (*accountability*) для підтвердження факту стирання перед регуляторами. Доведено, що

принцип Unlearning-by-Design є обов'язковою правовою вимогою (похідною від Art. 25 GDPR), що вимагає превентивної інтеграції механізмів сегментації та федеративного навчання для забезпечення структурної доказовості.

Отже, лише через міждисциплінарний діалог та взаємодію правової доктрини й технічних стандартів можна перетворити машинне рознавчання на дієвий правовий механізм захисту прав особи у цифрову епоху.

### Список використаних джерел

1. **European Union.** General Data Protection Regulation (GDPR). Regulation (EU) 2016/679, Art. 5, 17.
2. **European Parliament and Council.** EU Artificial Intelligence Act. Final text (2024).
3. **Court of Justice of the European Union (CJEU).** Press Release No 70/14 concerning the judgment in Case C-131/12, Google Spain SL, Google Inc. v. Agencia Española de Protección de Datos (AEPD) and Mario Costeja González, 2014.
4. **Garante per la Protezione dei Dati Personali (Italy).** Comunicato Stampa: Intelligenza artificiale: il Garante blocca ChatGPT... (Raccolta illecita di dati personali. Assenza di sistemi per la verifica dell'età dei minori). March 31, 2023.
5. **Верховна Рада України.** Закон України «Про захист персональних даних» № 2297-VI від 1 червня 2010 р.
6. **Standing Committee of the National People's Congress of China.** Personal Information Protection Law of the People's Republic of China. Adopted August 20, 2021. URL: [http://en.npc.gov.cn.cdurl.cn/2021-12/29/c\\_694559.htm](http://en.npc.gov.cn.cdurl.cn/2021-12/29/c_694559.htm).
7. **Ren, J., Xing, Y., Cui, Y., Aggarwal, C. C., & Liu, H.** SoK: Machine Unlearning for Large Language Models. *arXiv preprint arXiv:2506.09227*, 2025.
8. **Feretzakis, G., et al.** GDPR and Large Language Models: Technical and Legal Obstacles. *Future Internet*, 2025, 17, 151.
9. **Bourtole, A., et al.** Machine Unlearning. *Proceedings of the IEEE Symposium on Security and Privacy*, 2021.
10. **European Data Protection Board (EDPB).** Taskforce on ChatGPT Reports and Recommendations, 2023–2024.
11. **National Institute of Standards and Technology (NIST).** AI Risk Management Framework (AI RMF 1.0), 2023.
12. **Organisation for Economic Co-operation and Development (OECD).** OECD Principles on Artificial Intelligence, 2019.

**Світлана Келип**

аспірантка, ДУ «Київський авіаційний  
інститут»

ORCID: <https://orcid.org/000-0002-0614-1328>

## **НОРМАТИВНО-ПРАВОВА МОДЕРНІЗАЦІЯ СФЕРИ ПРИКОРДОННОЇ БЕЗПЕКИ УКРАЇНИ В УМОВАХ РОЗВИТКУ ТЕХНОЛОГІЙ ШТУЧНОГО ІНТЕЛЕКТУ**

У сучасних реаліях глобальних безпекових викликів і загроз для України першочерговим завданням є збереження стабільності та зміцнення надійності захисту державних кордонів. Адже це фактор фізичного існування держави, як політико-територіального утворення й передумова її міжнародного визнання не тільки в історичний момент творення, а й у подальшому розвитку. Без перебільшення, державний рубіж – це безпекова базисна гарантія для збереження інституційності влади, правової системи держави, її політичного вектору, домінування національних інтересів, пріоритетів усталених суспільних відносин, економічних зв'язків та різномаїття розвитку соціокультурного середовища. Тож питання прикордонної безпеки України, особливо в контексті стрімкого розвитку штучного інтелекту потребує глибоких теоретико-прикладних досліджень, результатом яких повинні стати конкретні напрацювання пропозицій до регулюючих нормативно-правових актів за всією їх вертикаллю: від загальних до відомчих і локальних. Українські науковці досліджують питання впливу технологій Artificial Intelligence (AI) на безпековий та оборонний сектори, правоохоронну систему України загалом. Наукове співтовариство докладає інтелектуальних зусиль до створення потужного пулу концептуальних підходів праворозуміння впровадження високотехнологічних новацій, серед яких штучний інтелект займає провідне місце. Результатом цього повинна стати трансформація наукових моделей у пріоритетне напрацювання законодавчих ініціатив, що чітко врегулюють, унормують і таким чином легітимізують використання по суті нового явища реальної дійсності.

Звісно, в цьому процесі лідерство повинна отримати сфера національної безпеки та оборони, і не в останню чергу прикордонна безпека, як головна

запорука охорони національних інтересів України. Адже державне реагування на зміни повинно бути швидким та оперативним.

Безперечно, вже заслуговують на особливу увагу наукові праці Шевчука В.М., Пацурії Н.Б., Павленко Т.А., Коновалової В.О. та інших. Дослідження цих вчених враховують реалії військової агресії російської федерації проти України, ще більше актуалізуючи важливість реалізації наукових досягнень у реальний сектор нормативно-правового регулювання.

Так, Шевчук В.М. звертає увагу саме на застосування технологій штучного інтелекту у практиці правоохоронної діяльності [1, с.358]. Пацурія Н.Б. пропонує сучасне бачення впровадження технологій штучного інтелекту в забезпечення національної безпеки й обороноздатності України у повоєнний період [2, с.75].

У свою чергу не достатньо науково дослідженою та пропозиційною залишається сфера безпеки державних кордонів, як невід'ємний сегмент фундаментальних національних інтересів і складова національної безпеки України [3]. Необхідність наукового погляду на дане питання підкреслюється нещодавніми доповненнями до Стратегії національної безпеки України, затвердженої указом Президента України від 14 вересня 2020 року № 392/20, якими до напрямків реалізації пріоритетів національних інтересів України поряд з іншими включено забезпечення належного рівня прикордонної безпеки в умовах розвитку інтегрованого управління державним кордоном України [4]. При цьому додатковою аргументацією слугує те, що розробки у сфері штучного інтелекту визначені у складі поточних і прогнозованих загроз національній безпеці та національним інтересам України [5].

Тож, з метою ефективного захисту територіальної цілісності, як одного із ключових пріоритетів національних інтересів України, потребують нормативно-правових новацій та актуалізації процедурні механізми, що підлягають до застосування структурними одиницями Державної прикордонної служби України, як правоохоронним органом спеціального призначення. І це, насамперед, стосується унормування використання у повсякденній службовій діяльності інтелектуальних алгоритмів АІ під час охорони та забезпечення дотримання державного кордону України, прикордонного режиму та здійснення прикордонного контролю [6].

Комплексний, колективний наукомісткий підхід до проблематики своєчасних нормотворчих напрацювань на основі прогностичних правових

досліджень дозволить уникнути неузгодженості між законодавчими та підзаконними актами, а також якісно унормувати застосування АІ Державною прикордонною службою України при виконанні функцій оперативно-службової діяльності.

### **Список використаних джерел:**

1. Шевчук В.М. Роль технологій штучного інтелекту у правоохоронній діяльності та забезпеченні безпеки та обороноздатності України. *Юридичний науковий електронний журнал*, № 6. 2024. С.356-361.

2. Пацурія Н.Б. Упровадження технологій штучного інтелекту в забезпечення національної безпеки та обороноздатності України: правові проблеми і перспективи повійськового періоду. *Теорія і практика інтелектуальної власності*. № 3. 2023. С. 68–78.

3. Про Державний кордон України: Закон України від 04.11.1991 № 1777-XII. *Відомості Верховної Ради України*. 1991. № 2. Ст.5.

URL: <https://zakon.rada.gov.ua/laws/show/1777-12> (дата звернення: 05.11.2025).

4. Про рішення Ради Національної безпеки і оборони України від 14.09.2020 «Про стратегію національної безпеки України»: указ Президента України від 14.09.2020 № 392/2020. URL: <https://zakon.rada.gov.ua/laws/show/392/2020#Text> (дата звернення: 05.11.2025).

5. Про національну безпеку України: Закон України від 21.06.2018 № 2469-VIII. *Відомості Верховної Ради України*, 2018. № 31. Ст. 241.

URL: <https://zakon.rada.gov.ua/laws/show/2469-19> (дата звернення: 05.11.2025).

6. Про Державну прикордонну службу України: Закон України від 03.04.2003 № 661-IV. *Відомості Верховної Ради України*, 2003. № 27. Ст.208.

URL: <https://zakon.rada.gov.ua/laws/show/661-15#Text> (дата звернення: 05.11.2025).

**Михайло Михайленко**

кандидат юридичних наук, молодший науковий  
співробітник Науково-дослідного інституту  
інтелектуальної власності НАПрН України  
ORCID: <https://orcid.org/0009-0008-6922-4658>

## **ПРОБЛЕМИ ТА ПЕРСПЕКТИВИ ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ІІІ ПІД ЧАС РОЗГЛЯДУ ЗАЯВОК ПРО ДЕРЖАВНУ РЕЄСТРАЦІЮ ТОРГОВЕЛЬНИХ МАРОК**

У сучасному світі торговельна марка є важливим елементом розвитку бізнесу, а також захисту споживачів. Аналіз глобальної статистики подання заявок про державну реєстрацію торговельних марок з 2014 по 2023 роки показує, що їхня кількість зросла на 99.34% (з 7 642 700 до 15 234 900 заявок) [1]. При цьому за останні 9 років відбулося значне зростання навантаження на національні органи інтелектуальної власності в частині необхідності якісного та своєчасного розгляду заявок про державну реєстрацію торговельних марок. В Україні зростання кількості заявок є менш вираженим, порівняно з глобальними тенденціями. Згідно зі статистикою ВОІВ у 2014 році кількість поданих заявок в Україні становила 34 810, своєю чергою, у 2023 році вона становила 44 708 заявок, що свідчить про зростання на 28.43%. Відповідно до останнього звіту ВОІВ на одного експерта припадало 567,7 заявки, а затримка у реєстрації складала 673 дні, що свідчить про перевантаження Українського національного офісу інтелектуальної власності та інновацій (далі – УКРНОІВІ) [2].

Використання технологій ІІІ у сфері торговельних марок в іноземних державах частково висвітлювалось у працях Андрощука Г.О., проте нам не вдалося виявити публікації вітчизняних вчених присвячених питанню перспектив та проблем використання таких технологій під час розгляду заявок про державну реєстрацію торговельних марок в Україні.

З огляду на вищевказане, є актуальним пошук підходів до підвищення якості та своєчасності розгляду заявок про державну реєстрацію торговельних марок з використанням технологій ІІІ в Україні.

Враховуючи постійний розвиток технологій ІІІ, а також успішний досвід їхнього використання у багатьох сферах діяльності суспільства,

державні відомства, які відповідальні за реєстрацію об'єктів інтелектуальної власності, поступово впроваджують такі технології у свою діяльність. Станом на 24 вересня 2024 року більш ніж 22 національні органи інтелектуальної власності різних країн інкорпоровали у свою діяльність технології ІІІ [3]. Проаналізувавши інформацію ВОІВ щодо ініціатив у сфері ІІІ у національних органах інтелектуальної власності, можна зробити висновок, що найбільша їхня кількість (близько 30) застосовується під час розгляду патентних заявок (у процесі патентної класифікації, проведення патентного пошуку тощо). У свою чергу, під час розгляду заявок про державну реєстрацію торговельних марок таких ініціатив реалізовано значно менше (близько 16) і вони застосовуються, зокрема, для пошуку схожих торговельних марок, а також їхньої класифікації [3]. Отже, метою цього дослідження є визначення перспективності та доцільності використання ІІІ під час розгляду заявок про державну реєстрацію торговельних марок в Україні. Для цього розглянемо стан впровадження та використання таких технологій у процесі розгляду заявок про реєстрацію торговельних марок у країнах-лідерах за кількістю поданих заявок – Китайській народній республіці та США.

У 2023 році до Національного управління інтелектуальної власності Китаю (далі – НУІВК) надійшло 7 416 716 заявок, що є найбільшою кількістю у світі. Безумовно заслуговує на увагу є те, що на одного експерта НУІВК припадало 6 253 заявки, проте затримка у реєстрації торговельних марок складала лише 120 днів [4]. Отже, незважаючи на той факт, що на експерта НУІВК припадало на 5 685 заявок, більше ніж на експерта УКРНОІВІ, затримка реєстрації торговельної марки є меншою на 553 дні.

Досягнути такого результату вдалося, зокрема, завдяки ретельній підготовці та тестуванню використання новітніх технологій у процесі розгляду заявок. З метою вирішення проблеми низького рівня інтелектуалізації та високої інтенсивності роботи експертів у процесі експертизи зображувальних торговельних марок у КНР, а також для підвищення ефективності та якості експертизи, НУІВК активно досліджувало можливості застосування технологій ІІІ. З моменту створення робочої групи з пошуку зображувальних знаків у квітні 2016 року було організовано кілька спеціальних опитувань. У липні 2018 року, за значної підтримки компаній, функція інтелектуального пошуку торговельних марок була успішно протестована у 6 центрах експертизи та співпраці у справах торговельних марок.

Починаючи з 2019 року НУІВК розпочав застосування Інтелектуального пошуку за графічним зображенням торговельних марок Китаю, що дозволило перетворити роботу з перевірки торговельних марок з чисто ручного пошуку на інтелектуальний пошук. Ця технологія ІІІ використовує механізм порівняння «зворотних зображень», який був вбудований у робочий процес експертів. Завдяки значному скороченню кількості порівнянь схожих торговельних марок, була підвищена продуктивність перевірки, а кількість переглядів скоротилася з десятків тисяч до приблизно п'яти тисяч [5]. Використання технологій ІІІ продовжується і надалі, у 2023 році згідно з інформацією НУІВК, рівень успішності експертизи торговельних марок, процедур оскарження і винесення рішень перевищив 97% [6].

США, своєю чергою, станом на 2023 рік займали друге місце у світі за кількістю отриманих заявок про реєстрацію торговельних марок, отримавши 848 945 заявок. У той же час, на одного експерта Бюро патентів і торгових марок США (далі – USPTO) припадало 1007 заявок, а затримка у реєстрації складала 444 дні, що майже у 4 рази більше, ніж у НУІВК [7]. Підхід США щодо впровадження технологій ІІІ у процеси розгляду заявок про реєстрацію торговельних марок відрізняється від досвіду КНР.

Перспективність використання технологій ІІІ в USPTO у сфері реєстрації торговельних марок була закріплена, зокрема, у Річному звіті Комітету з питань торговельних марок 2019 року. У ньому вказувалось про наявність проблеми подання заявниками чи власниками торговельних марок підроблених чи змінених будь-яким способом доказів використання торговельної марки та вказувалось про створення «Проекту автоматизованого аналізу зразків торговельних марок», в рамках якого використання технології ІІІ дозволило б виявляти недоброчесні зразки [8]. Наразі цей проєкт усе ще перебуває на стадії тестування.

Іншим напрямком роботи USPTO є створення інструментів на основі ІІІ для пошуку зображень та ідентифікації товарів, призначених для допомоги експертам, що розглядають заявки на реєстрацію торговельних марок [с. 1, 9]. На практиці пошук зображень на основі технологій ІІІ має дозволити експерту завантажити зображення з заявки та знайти візуально схожі знаки з бази даних зображень торговельних марок USPTO, замість того, щоб покладатися виключно на стару класифікацію кодів дизайну. USPTO розробило прототипи, які порівнюють зображення торгових марок і навіть

пропонують правильні класифікації кодів дизайну для цих зображень, проте вони ще перебувають на стадії тестування [10].

Важливим до врахування є те, що крім практичних заходів з розробки та тестування технологій ІІІ, додатково проводиться активне обговорення між USPTO та усіма зацікавленими суб'єктами питань щодо впровадження технологій ІІІ у роботу відомства. Так, 07.06.2022 року також було оголошено про створення партнерства з ІІІ та новітніх технологій, у рамках якого за участю вчених, винахідників, підприємств, урядових установ, а також громадянського суспільства було проведено понад 15 онлайн-зустрічей для обговорення актуальних питань щодо ініціатив впровадження ІІІ та новітніх технологій у діяльність USPTO [11].

Отже, попри значне навантаження на експертів USPTO наразі використання технологій ІІІ для розгляду заявок на торговельні марки у США загалом перебуває на початковій стадії планування, а кілька інструментів проходять тестування. Попри це, беззаперечним є те, що USPTO бачить перспективність використання технологій ІІІ у цій сфері та активно працює над їхнім впровадженням.

Використання технологій ІІІ під час розгляду заявок про державну реєстрацію торговельних марок невід'ємно створює ряд проблемних ризиків, які є необхідними до врахування для ефективного та безпечного впровадження таких технологій. Можна визначити наступні ризики, зокрема:

1) Надмірна залежність від інструментів ІІІ, яка може призвести до перекладання на ІІІ інтелектуальної роботи, яка мала бути зроблена експертом [12];

2) Обмежена точність та надійність. У типових та розповсюджених випадках, точність результатів наданих технологіями ІІІ становить близько 90%, проте у складних та неоднозначних випадках через відсутність чутливості та надмірну впевненість ІІІ, точність результатів різко падає та іноді заважає експерту більше, ніж допомагає [13];

3) Упередженість та справедливість. Цей ризик виникає внаслідок складності формування наборів даних для навчання системи ІІІ, які не міститимуть упереджень чи прогалин, оскільки у випадку неусунення таких негативних факторів система ІІІ буде ненавмисно упередженою та неточною [с. 31, 14];

4) Відсутність прозорості та зрозумілості. Прийняття рішень ІІІ часто функціонує як «чорний ящик», тобто експерт та інші зацікавлені особи не можуть зрозуміти підстави та обґрунтування висновку, наданого технологією ІІІ.

5) Безпека та конфіденційність. Враховуючи те, що передові генеративні моделі ІІІ, такі як великі мовні моделі, працюють у хмарних середовищах або надаються третіми особами, існує ризик розкриття конфіденційної інформації заявників.

Отже, міжнародний досвід свідчить про перспективність та активну роботу з розробки та інтегрування технологій ІІІ до процесу розгляду заявок про державну реєстрацію торговельних марок. Зокрема, США наразі проводять широкі консультації з усіма зацікавленими сторонами, а також практично тестують велику кількість технологій ІІІ у цій сфері, а опубліковані дані КНР доводять фактичну ефективність використання цих технологій для допомоги експерту у виявленні та порівнянні візуально схожих попередньо зареєстрованих знаків. Проте, використання технологій ІІІ створює ряд проблемних ризиків, пов'язаних із їх функціонуванням, ігнорування яких може призвести не тільки до зменшення ефективності їхнього використання, але і до надання хибних результатів та введення в оману експертів, які проводять експертизу заявки.

Досліджуючи та враховуючи досвід країн-лідерів у використанні ІІІ у сфері інтелектуальної власності, для підвищення якості, а також скорочення строків розгляду заявок про державну реєстрацію торговельних марок в Україні, можна зробити висновок про доцільність наступних дій:

- розпочати проведення широких консультацій з громадськістю та усіма зацікавленими сторонами щодо оптимального інтегрування технологій ІІІ у процес експертизи заявок про державну реєстрацію торговельних марок;
- на основі проведених консультацій створити міжвідомчу робочу групу, до складу якої мають увійти уповноважені представники усіх зацікавлених сторін;
- за координації УКРНОІВІ розпочати роботу зі створення, навчання, а згодом тестування технологій ІІІ, які зможуть підвищити ефективність експертизи заявок про державну реєстрацію торговельних марок.

### **Список використаних джерел:**

1. WIPO statistics database. WIPO IP Statistics Data Center. Last updated: May 2025. <https://www3.wipo.int/ipstats/key-search/search-result?type=KEY&key=201>.
2. Intellectual property statistical country profile 2023 Ukraine. WIPO IP Statistics Data Center. Last updated: May 2025. <https://www.wipo.int/edocs/statistics-country-profile/en/ua.pdf>.
3. Index of AI initiatives in IP offices. An official website of the WIPO. <https://www.wipo.int/en/web/ai-tools-services/ipos-initiatives>.
4. Intellectual property statistical country profile 2023 China. WIPO IP Statistics Data Center. Last updated: May 2025. <https://www.wipo.int/edocs/statistics-country-profile/en/cn.pdf>.
5. China Trademark Intelligent Graphic Retrieval Launched. Unitalen Attorneys at Law. March 13, 2019. <https://www.unitalen.com/html/report/19031962-1.htm>.
6. Improving the quality and efficiency of intellectual property review will ensure the steady and long-term progress of building a strong intellectual property nation (Intellectual Property News). An official website of the Chinese National Intellectual Property Administration. Release date: 2025-01-02. [https://www.cnipa.gov.cn/art/2025/1/2/art\\_55\\_197014.html](https://www.cnipa.gov.cn/art/2025/1/2/art_55_197014.html). In Chinese.
7. Intellectual property statistical country profile 2023 United States of America. WIPO IP Statistics Data Center. Last updated: May 2025. <https://www.wipo.int/edocs/statistics-country-profile/en/us.pdf>.
8. Trademark Public Advisory Committee Fiscal Year 2019 Annual Report. United States Patent and Trademark Office [https://www.uspto.gov/sites/default/files/documents/TPAC\\_2019\\_Annual\\_Report.pdf](https://www.uspto.gov/sites/default/files/documents/TPAC_2019_Annual_Report.pdf)
9. USPTO Should Improve Governance to Promote Effective Oversight of Its Artificial Intelligence Tools. Final Report No. OIG-25-018-A. U.S. Department of Commerce Office of Inspector General Office of Audit and Evaluation. April 22, 2025. [https://www.oig.doc.gov/wp-content/OIGPublications/OIG-25-018-A\\_SECURED.pdf](https://www.oig.doc.gov/wp-content/OIGPublications/OIG-25-018-A_SECURED.pdf)
10. Tran Alex. USPTO deploys numerous AI tools to aid patent and trademark examination. On February 13, 2023. PowerPatent.com. <https://powerpatent.com/blog/uspto-deploys-numerous-ai-tools-to-aid-patent-and-trademark-examination>.
11. AI and Emerging Technology Partnership engagement and events. An official website of the United States Patent and Trademark Office. <https://www.uspto.gov/initiatives/artificial-intelligence/ai-and-emerging-technology-partnership-engagement-and-events>.

12. Report on the sharing sessions on the use of artificial intelligence and other tools for effective patent examination procedures. Standing Committee on the Law of Patents. Thirty-Seventh Session. WIPO, Geneva, November 3 to 7, 2025. [https://www.wipo.int/edocs/mdocs/scp/en/scp\\_37/scp\\_37\\_6.pdf](https://www.wipo.int/edocs/mdocs/scp/en/scp_37/scp_37_6.pdf).

13. Burbidge Rosie. Can AI reshape trade mark practice – promise, pitfalls and the average consumer of the future. Sep 23, 2025. <https://www.rosieburbidge.com/post/can-ai-reshape-trade-mark-practice-promise-pitfalls-and-the-average-consumer-of-the-future>

14. Public views on artificial intelligence and intellectual property. United States Patent and Trademark, Oct 2020. [https://www.uspto.gov/sites/default/files/documents/USPTO\\_AI-Report\\_2020-10-07.pdf](https://www.uspto.gov/sites/default/files/documents/USPTO_AI-Report_2020-10-07.pdf).

**Дмитро Ланде**

Національний технічний університет України  
"Київський політехнічний інститут імені Ігоря  
Сікорського"

ORCID: <https://orcid.org/0000-0003-3945-1178>

**Юрій Цирульнєв**

Національний технічний університет України  
"Київський політехнічний інститут імені Ігоря  
Сікорського", Засновник та директор DIGITAL  
DOCS®, <https://digital-docs.eu>

ORCID: <https://orcid.org/0009-0007-2723-3137>

## **ТЕОРЕТИЧНА МОДЕЛЬ ПРАВОВОГО РЕГУЛЮВАННЯ АСПЕКТІВ СТВОРЕННЯ ТА ВИКОРИСТАННЯ ЕЛЕКТРОННИХ ІНФОРМАЦІЙНИХ РЕСУРСІВ ТА ЕЛЕКТРОННИХ АРХІВІВ**

**Мета дослідження** - продовження започаткованого в роботі [1] ґрунтовного обговорення питань правового регулювання аспектів створення та використання електронних інформаційних ресурсів (*Digital Information Resources, DIR*) та електронних архівів (*Digital Archives, DA*) з урахуванням викликів сучасності, пов'язаних із застосуванням штучного інтелекту (*Artificial Intelligence, AI*), машинного навчання (*Machine Learning, ML*) і великих мовних моделей (*Large Language Models, LLM*) в процесах інтелектуальної агрегації інформації (*Intellectual Information Aggregation, IIA*) в електронні архіви.

**Наукова новизна дослідження** міститься в формуванні і математичній формалізації теоретичної моделі правового регулювання процесів створення та використання DIR і DA; введенні поняття правової резильєнтності, як характеристики здатності нормативно-правового середовища зберігати цілісність регулювання під дією технологічних змін; математичній формалізації взаємодії правових, технологічних та організаційних складових системи електронних архівів; підхід до оцінювання кібербезпеки та нормативної відповідності на основі інтеграції моделей LLM-процесів і індексів відповідності; обґрунтуванні зв'язків між правовими принципами,

технологічними механізмами та організаційним управлінням у процесах створення і використання DIR і DA.

**Узагальнена логічна структура моделі** являє собою систему взаємозалежностей, яка інтегрує всі складові, де: LF — нормативно-правова база, TF — технологічна інфраструктура, OM — організаційна структура, LR — правова резильєнтність. Сутності та множини:  $D$  — множина документів (цифрові об’єкти);  $M$  — множина метаданих (ISO 23081);  $A$  — актори (користувачі, адміністратори, зовнішні системи);  $P$  — процеси (сканування, OCR, класифікування, індексування, зберігання, доступ, вилучення, міграція, знищення);  $T = \{0, 1, \dots\}$  — дискретний час;  $N = \{n_1, \dots, n_K\}$  — норми права / стандарти (закони, GDPR, EU AI Act, ISO, внутрішні політики);  $C = \{c_1, \dots, c_L\}$  — технічні контрзаходи (PKI, EDS, хешування, RBAC, шифрування тощо);  $E$  — події аудиту / журналювання;  $\Omega$  — простір збурень (зовнішні події, інциденти, атаки, помилки LLM). Стан системи в момент  $t$ :  $s_t = (D_t, M_t, G_t, \Pi_t, \Xi_t)$ , де  $G_t$  — граф походження / провенансу,  $\Pi_t$  — набір активних політик,  $\Xi_t$  — активні контролю. Динаміка (узагальнено):  $s_{t+1} = f(s_t, u_t, \omega_t)$ ,  $u_t \in U$  (рішення управління),  $\omega_t \in \Omega$ . Процесна модель (DAG / мережа Петрі) описує послідовність та взаємозв’язки між операціями системи. Орієнтований ациклічний граф (DAG) відображає потік даних чи дій від початкових вузлів до кінцевих без циклів — тобто без повторення станів. Мережа Петрі розширює цю модель, вводячи *місця* (стани системи) та *переходи* (операції, що змінюють стан). AI/LLM-оператори моделюються як спеціальні переходи, що автоматизують частину цих дій у межах єдиного процесу. Процеси подани, як орієнтований ациклічний граф:

$$P = (V, E), V \subseteq P, E \subseteq V \times V,$$

або як мережу Петрі  $(S, T, F)$  з місцями  $S$  (буфери станів) та переходами  $T$  (операції). AI/LLM-оператори — підмножина переходів  $T_{LLM} \subset T$ .

**Правові норми як логіко-алгебраїчні обмеження.** Кожна норма  $n_k$  кодується предикатом  $\varphi_k$ , який відображає істинність або хибність певної умови в межах станів системи  $S$  над станом/подіями  $E$ :  $\varphi_k: S \times E \rightarrow \{true, false\}$ .

**Типи норм:** обов’язок  $O(\psi)$ : певна умова  $\psi$  має виконуватись на всіх допустимих траєкторіях системи (наприклад, “усі доступи мають логуватись”); заборона  $F(\psi)$ : деяка дія або стан не повинні відбуватися (наприклад, “не можна змінювати метадані без перевірки підпису”); дозвіл/умова  $P(\psi \vee \chi)$ : дозволено виконати  $\psi$ , якщо виконується  $\chi$  (наприклад, “якщо користувач автентифікований, дозволено завантаження файлу”).

**Запис у LTL/CTL (приклад часової логіки LTL для логування доступу):**

$$G(\text{Access}(d, a) \Rightarrow F(\log(d, a))).$$

Функція відповідності нормі показує, чи виконується конкретна норма  $n_k$  у стані  $s_t$ :

$$\text{comp}_k(s_t) = 1[\varphi_k(s_{\leq t}, \varepsilon_{\leq t}) - \text{true}], \text{comp}(s_t) = (\text{comp}_k, \text{comp}_k, \dots, \text{comp}_k).$$

**Зважене інтегральне дотримання:**

$$CS(s_t) = \sum_{k=1}^K w_k \text{comp}_k(s_t), w_k \geq 0, \sum w_k = 1$$

**Автентичність, цілісність, походження**

Для документа  $d \in D$ : Хеш:  $h(d) = H(d)$ . ЕЦП:  $\sigma(d) = \text{Sign}_{\text{PKI}}(h(d))$ .

**Повнота метаданих:**

$$\mu(d) = M_{\text{present}}(d) \vee \frac{1}{M_{\text{required}} \forall \in [0,1]}.$$

**Узгодженість провенансу:**

$$\pi(d) = 1[\text{існує валідний шлях у } G_t \text{ від джерела до } d].$$

**Індекс автентичності:**  $A(d) = \alpha_1 \text{Verify} + \alpha_2 \mu(d) + \alpha_3 \pi(d), \sum \alpha_i = 1$ .

Модель LLM-перетворень і похибок формалізує вплив LLM-операцій на зміну змісту документів. Ймовірність галюцинації визначає відхилення від допустимого простору похибок. Правовий допуск зміни задає граничну норму змістової різниці, прийнятну в юридичному контексті. LLM-оператор  $L$  над об'єктом  $x$  (текст, метадані) з параметрами  $\theta$ :  $x' = L_\theta(x) = x + \varepsilon, \varepsilon \sim E_\theta$ .

**Ймовірність галюцинації:**  $p_h = \text{Pr}[\varepsilon \notin \varepsilon_{\text{доп}}]$ .

**Метрика змістової відстані (наприклад, embedding-cosine):**

$$\Delta(x, x') = 1 - \cos(e(x), e(x')) \in [0, 2].$$

Правовий допуск зміни:  $\Delta(x, x') \leq \delta_{\text{law}}$ . Норма на LLM-використання:

$$\varphi_{LLM}: G \left( \text{UseLLM}(x) \Rightarrow (\Delta(x, x') \leq \delta_{\text{law}} \wedge \text{LogLLM}(x \rightarrow x')) \right).$$

**Ризики і показник правової безпеки.** Набір ризиків  $R = \{r_1, r_2, \dots, r_j\}$  (достовірність, доступ, автентичність, приватність, маніпуляції):  $R_j = p_j \cdot I_j, p_j \in [0, 1], I_j \geq 0$ .

**Інтегральний ризик:**  $RI(s_t) = \sum_{j=1}^J \beta_j R_j, \sum \beta_j = 1$ .

**Правовий показник безпеки (вища — краще):**

$$LS(s_t) = \lambda_1 CS(s_t) - \lambda_2 RI(s_t), \lambda_{1,2} > 0.$$

**Резильєнтність (правова та функціональна):**

$$LR = \frac{1}{t_R - t_0} \int_{t_0}^{t_R} CS(t) dt.$$

**Комбінований показник резильєнтності:**

$$RZ = \eta_1 \cdot \frac{1}{t_R - t_0} \int_{t_0}^{t_R} F(t) dt + \eta_2 \cdot LR - \eta_3 \cdot \tau_R,$$

де:  $\tau_R = t_R - t_0$  — час відновлення;  $\eta_i > 0$ .

### Оптимізаційна постановка

$$\max_{\pi, \mathcal{E}} E \left[ \sum_{t=0}^T \gamma^t (\alpha LS(s_t) + \rho RZ_t) - \kappa \text{Cost}(\mathcal{E}) \right]$$

за умов:  $\forall k, \forall t: \text{comp}_k(s_t) = 1$  (критичні норми),  $\forall \mathcal{E} \forall \leq B$ .

**Модель “порушник–захисник”** (ігрова перспектива). Ігрова модель описує взаємодію двох сторін: Порушник прагне максимізувати виграш від атаки; Захисник мінімізує втрати системи та підтримує нормативну відповідність і резильєнтність. Вихідні множини:  $A_{adv} = \{a_1, a_2, \dots, a_m\}$  — множина можливих атак (зміна даних, компрометація метаданих, маніпуляції LLM, тощо);  $\mathcal{E} = \{c_1, c_2, \dots, c_n\}$  — множина можливих контрзаходів або контролів (реплікація, аудит, explainability, верифікація, хешування тощо). Кожна зі сторін обирає свою дію з певною ймовірністю:  $a \sim q(a), c \sim p(c)$ .

**Виграш зловмисника.** Зловмисник оцінює результат своєї атаки як різницю між вигодою від порушення цілісності або довіри та вартістю подолання контролів:

$$U_{adv}(a, c) = \sum_j \beta_j R_j(a, c) - \phi(c),$$

де:  $R_j(a, c)$  — частковий результат атаки за аспектом  $j$  (наприклад, компрометація автентичності, порушення приватності, зниження резильєнтності);  $\beta_j$  — ваговий коефіцієнт значущості;  $\phi(c)$  — “вартість” подолання обраного контролю  $c$ . Сенс: чим слабше контроль, тим вищий очікуваний виграш атакуючого.

**Виграш захисника.** Захисник має ціль максимізувати безпеку та відповідність системи нормативам, мінімізуючи витрати:

$$U(a, c) = -U_{adv}(a, c) + \psi(C, SRZ) - \text{Cost}(c),$$

де:  $\psi(C, SRZ)$  — функція вигоди від зростання відповідності нормам (Compliance Score) та резильєнтності (Resilience Index);  $\text{Cost}(c)$  — ресурси, витрачені на впровадження контролю (енергія, обчислення, персонал). Сенс: захисник балансує між безпекою й вартістю контролю, підтримуючи стабільність системи.

**Відображення відповідності.** Нехай  $\Phi: N \rightarrow 2^C, \Phi(n_k) = \{c_l: c_l \text{ покриває } n_k\}$ , де:  $N = \{n_k, n_k, \dots, n_k\}$  — множина нормативів або принципів (ISO, AI Act, GDPR

тощо);  $C = \{c_1, c_2 \dots, c_l\}$  — множина контрольних механізмів (аудит, хеш-контроль, реплікація, LLM-логіка, інтероперабельність тощо). Таким чином, кожна норма  $n_k$  може бути покрита одним або кількома контролями  $c_l$ .

**Оцінка повноти покриття.** Для оцінки того, наскільки повно реалізовано нормативну базу технічними механізмами, вводимо метрику:

$$Cover = \frac{1}{K} \sum_{k=1}^K 1 [\Phi(n_k) \cap \Xi \neq \emptyset],$$

де:  $\Xi$  — множина реально активних або реалізованих контролів,  $K$  — загальна кількість норм. Якщо  $Cover = 1$ , усі нормативи покриті хоча б одним механізмом контролю. Якщо  $Cover < 1$ , існують “незахищені” норми, що вимагають додаткових технічних або організаційних рішень.

**Глибина покриття.** Для кожної критичної норми визначається “глибина покриття”:  $d_k = \Phi(n_k) \cap \Xi \vee$ , відображає кількість незалежних контролів, що одночасно забезпечують дотримання  $n_k$ . Високий  $d_k$  — норма дубльована кількома шарами контролю → підвищена резильєнтність. Низький  $d_k$  — ризик вразливості або надмірної залежності від одного механізму.

### Список використаних джерел:

1. Ковтанюк Ю. С. Нормативно-правове регулювання оцифрування фондів закладів культури як вимога для розбудови державних інтеграційних електронних ресурсів національного історичного та культурного надбання. — Інформація і право, № 1 (44), 2023.
2. Lande D. V., Tsyurulnev Yu. B. Cybersecurity of Intellectual Information Aggregation Processes into Digital Archives. — Theoretical and Applied Cybersecurity, № 2, 2025
3. Ланде Д. В., Цирульнєв Ю. Б. Резильєнтність електронних архівів. *Реєстрація, зберігання і обробка даних*. Т. 29. № 2. 2025

# **ФІЛОСОФСЬКО-ПРАВОВІ ТА КОНЦЕПТУАЛЬНО-МЕТОДОЛОГІЧНІ ЗАСАДИ ДОСЛІДЖЕННЯ ТА РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ**

**Геннадій Андрощук**

кандидат економічних наук, доцент, провідний науковий співробітник НДІ інтелектуальної власності НАПрН України

ORCID: <https://orcid.org/0000-0003-0781-9740>

## **МЕТОДОЛОГІЯ ВИЯВЛЕННЯ НОВИХ ТЕХНОЛОГІЙ ШІ З ВИКОРИСТАННЯМ ПАТЕНТНИХ ДАНИХ**

Аналіз патентів надає унікальні відомості про розвиток технологій на світових ринках. Від конкурентної розвідки до майбутніх технологічних проривів аналіз патентних даних дає індикатори, які неможливо отримати з будь-якого іншого джерела. [1] Штучний інтелект (ШІ) – це масштабна технологія з безліччю сегментів. З кінця 2022 року, коли з'явився ChatGPT, увагу громадськості привернув сегмент ШІ, пов'язаний із генеративним ШІ (GenAI) та його здатністю створювати контент, який відчувається майже (чи, можливо, повністю!) людським.

ШІ провів пошук патентних заявок за останні 14 місяців, щоб отримати уявлення про місце, яке генеративний ШІ займає в широкій картині ШІ. Загальні класифікації патентів на ШІ включають такі технології, як таксономії для GenAI – застосунків, які створюють текст, відео, зображення та аудіо, – де є великі мовні моделі з архітектурою впровадження слів, змагальні мережі та нейронні мережі, які називаються генеративними попередньо навченими трансформаторами (звідси і аббревіатура GPT) GenAI. Розглядаючи заявки, подані тільки в USPTO, 24% цих патентів вважаються GenAI, що трохи вище, ніж у торішньому першому огляді цього простору. За останні 10 років трендова лінія для патентів як ШІ, так і GenAI попрямувала вгору. Гранти (позитивні рішення) на патенти ШІ зростали з річним темпом 38%, а заявки збільшилися на 31%. Для GenAI лінія ще крутіша: гранти збільшувалися на 58%, а заявки – на 52%. [1]

Пошук патентів полягає в систематичному інформаційному аналізі доступних патентних баз даних з метою з'ясування, чи були вже запатентовані аналогічні або схожі винаходи. Цей процес вимагає великої уваги до деталей, технічного розуміння і вміння працювати зі спеціалізованими інструментами та базами даних. Спеціалізований патентний пошук є складним завданням через обсяг і розмаїття наявних даних, використання специфічної термінології та технічних концепцій. [2] Для цього використовують інструменти аналітики інтелектуальної власності. [3,4] Аналітика інтелектуальної власності (Intellectual property analytics (IPA)) — міждисциплінарна наука про дані, яка використовує математику, статистику, комп'ютерне програмування, дослідження операцій для перетворення даних у знання і дає змогу аналізувати великий обсяг інформації про інтелектуальну власність, виявляти взаємозв'язки, тенденції та моделі прийняття рішень.

У цій статті представлено аналіз, розробленої ОЕСР, нової методології ідентифікації *«патентів на штучний інтелект»* — патентів, що захищають винаходи, пов'язані зі штучним інтелектом (ШІ). [5] Вона базується на існуючих таксономіях (системах класифікації) та враховує останні розробки в цій галузі, зокрема сплеск використання генеративних технологій ШІ. Методологія поєднує технологічні групи Спільної патентної класифікації (CPC) зі списком ключових слів та фраз, пов'язаних зі ШІ. Метод ідентифікації словника ШІ та класів CPC розроблений для забезпечення оновлення та інтеграції нових технологій ШІ. Ця методологія розширює охоплення винаходів ШІ, включаючи найновішу термінологію та сприяючи включенню останніх інновацій. Вона розроблена для легкого оновлення та адаптації до різних наборів патентних даних, базується на існуючих, включаючи попередню методологію ОЕСР. Вона враховує останні розробки в технології ШІ, зокрема, інтегруючи сплеск розвитку генеративного ШІ. Запропонований методологічний підхід спирається на стратегію пошуку на основі класів патентних технологій, використовуючи схему Спільної патентної класифікації (CPC) та словник, пов'язаний зі ШІ. Переглянута таксономія була розроблена з використанням напівавтоматизованого підходу. Вона визначає нову термінологію ШІ шляхом аналізу тексту останніх наукових статей. Потім списки груп CPC, потенційно пов'язаних зі ШІ, ідентифікуються на основі їх співіснування поряд із термінами ШІ в назвах патентів або рефератах, а також експертної перевірки. Такий підхід дозволяє часто оновлювати методологію

для оперативного врахування нових технологій ШІ. Хоча аналіз спирається на патентні заявки, подані відповідно до Договору про патентну кооперацію (РСТ), що забезпечує міжнародно порівнянні та своєчасні показники, методологію та отриману таксономію можна легко застосувати до інших наборів патентних даних.

Аналізована методологія складається з 3 розділів. У першому підсумовано останні зусилля щодо покращення патентів на ШІ. У другому розділі описано переглянута методологію ідентифікації патентів на ШІ. Третій розділ висвітлює тенденції в патентуванні ШІ, зосереджуючись на системі РСТ. Методологія використовує надійну стратегію пошуку, яка застосовує систему класифікації патентів, категоризуючи винаходи за певними технологічними галузями та використовуючи повний список ключових слів, пов'язаних зі ШІ, витягнутих з назв патентів та анотацій. Процес ідентифікації для розробки таксономії патентів ШІ є напівавтоматизованим і включає два ключові кроки:

- **Ідентифікація ключових слів та фраз:** Аналізуючи анотації останніх наукових статей зі ШІ за допомогою текстового аналізу на платформі arXiv з відкритим доступом, процес визначає нові ключові слова та ключові фрази до трьох слів (наприклад, «агентний штучний інтелект»), які відображають останні досягнення в цій галузі.

- **Визначення групи СРС:** Процес визначає відповідні групи Спільної патентної класифікації (СРС) на основі їх спільної появи з термінами ШІ в назвах патентів або анотаціях. Цей напівавтоматизований підхід дозволяє регулярно оновлювати методологію, гарантуючи, що вона відповідає досягненням у галузі ШІ та включає найактуальніший словник ШІ та відповідні класи СРС. Методологія була перевірена експертами за допомогою спеціального опитування.

Патенти класифікуються як ШІ, якщо:

- Вони належать до однієї з п'яти груп СРС, визначених як основні ШІ (G06N3, G06N5, G06N7, G06N20 або G06F18); або
- Належать до групи СРС, визначеної як пов'язана зі ШІ, та містять принаймні одне ключове слово ШІ у своїй назві або анотації.

Переглянута методологія значно підвищує ідентифікацію патентів на ШІ порівняно з попередньою методологією ОЕСР, зокрема, шляхом включення

додаткових ключових слів, пов'язаних з генеративним ШІ та методами обробки зображень.

Нещодавні тенденції в патентах на ШІ аналізуються з використанням даних із заявок РСТ, що забезпечує міжнародно порівнянні показники. Для вирішення проблеми адміністративних затримок у публікації патентів, останні дані "переглядаються".

З 2015 р. кількість запатентованих винаходів у технологіях ШІ різко зросла, із середньорічним темпом зростання 30% до 2020 року. До 2023 р. патенти на ШІ становили приблизно 6% усіх заявок РСТ – що втричі перевищує рівень 2015 року. **З 2020 по 2023 рік провідними країнами у розробці патентів на ШІ були Китайська Народна Республіка (далі – «Китай») (31%), Сполучені Штати Америки (США) (27%) та Японія (10,5%).** Однак останні тенденції свідчать про уповільнення патентування ШІ.

З огляду на зростаючий політичний, академічний та бізнес-інтерес до ШІ, а також унікальне багатство патентних даних для вимірювання технологічних винаходів, академічні кола та відомства інтелектуальної власності докладають зусиль для моніторингу останніх розробок у цій галузі.

Однак важливо визнати, що патентні дані дають лише часткову картину технологічного прогресу. **Не всі винаходи захищені патентами, і не всі винаходи є патентоспроможними.** Це обмеження особливо актуальне для технологій ШІ, оскільки патентоспроможність може відрізнятися залежно від юрисдикції. Наприклад, хоча патентне відомство США дозволяє патентування програмного забезпечення, Європейське патентне відомство (EPV) вважає патентоспроможними лише «комп'ютерні винаходи». Отож, ландшафт інновацій у сфері ШІ формується не лише патентними даними, але й правовими рамками, що регулюють патентоспроможність. Аналіз вказує на те, що результати патентних ландшафтів можуть суттєво відрізнятися, значною мірою залежно від визначення ШІ та типів патентів, що входять до їх складу. Нижче наведено короткий виклад останніх змін щодо визначення патентів на ШІ.

**ВОІВ (2019).** Стратегія запитів ВОІВ переважно використовує класифікаційні коди для подолання обмежень пошуку на основі ключових слів у широких технологічних галузях, таких як ШІ. Різні системи класифікації, що існують, – наприклад, Міжнародна патентна класифікація (МПК), Спільна патентна класифікація (СПС) або файловий індекс (FI) або термін формування

файлу (F-термін) Патентного відомства Японії (JPO) – відрізняються за обсягом, точністю та структурою, що вимагає індивідуальних підходів. Патенти на ІІІ часто не мають конкретних класифікаційних кодів ІІІ, що вимагає включення ширших категорій, контрольованих розширеними списками ключових слів, для зменшення інформаційного шуму. Деякі патенти можна ідентифікувати лише за ключовими словами, особливо там, де системи класифікації (наприклад, FI/F-терміни) обмежені певними юрисдикціями, такими як JPO.

Структура запиту включає:

- **Блок 1: Вибрані коди CPC, специфічні для технологій ІІІ.**
- **Блок 2: Ключові слова, специфічні для ІІІ.**
- **Блок 3: Ширші класи термінів CPC/IPC/FI/F, поєднані з ключовими словами, пов'язаними з ІІІ.**

Остаточний запит поєднує ці блоки. Релевантність та специфічність були перевірені за допомогою бази даних Orbit1 та ручної вибірки. Ключові слова були отримані з літератури та веб-ресурсів, з детальними пошуковими рядками, кодами та ключовими словами, наданими в документі.

**ОЕСР (2020).** У 2020 році ОЕСР розробила методологію для визначення патентів, пов'язаних з ІІІ, використовуючи комбінований підхід на основі IPC/CPC та ключових слів. Набір ключових понять, пов'язаних з ІІІ. Потім пошук застосовується до назв, анотацій та пунктів формули патентів, щоб визначити коди IPC або CPC, які є релевантними до ІІІ (які вважаються основними ІІІ або пов'язані з ІІІ). Як ключові слова, так і класи технологій, пов'язані з ІІІ, були перевірені експертами в цій галузі (дослідниками та/або патентними експертами). Аналіз був проведений у Всесвітній базі даних статистичної статистики патентів (EPO, PATSTAT Global)<sup>2</sup> Європейського патентного відомства (EPO) та в базі даних PatentsView<sup>3</sup> Патентного відомства США (USPTO).

Отримана стратегія пошуку була виконана в три кроки:

- **Вибір патентів з основними класами IPC або CPC ІІІ; та**
- **Вибір патентів з класами IPC або CPC, пов'язаними з ІІІ, які також містять принаймні одне ключове слово, пов'язане з ІІІ, в анотації, назві та/або пункті формули (лише для патентів USPTO); та**
- **Вибір патентів, які містять принаймні три ключові слова, пов'язані з ІІІ.**

Однак ця стратегія пошуку залежить від стану знань на певний момент часу та не охоплює новітні розробки в галузі ШІ, такі як генеративний ШІ.

**Country Activity Tacker (2020).** CSET, Центр безпеки та нових технологій Джорджтаунського університету, публікує САТ (Country Activity Tacker), що є базою даних та індикаторами патентів, пов'язаних зі ШІ. САТ охоплює 52 патентні юрисдикції та містить дані про патентні заявки та гранти. Для ідентифікації патентів на ШІ САТ поєднує технологічні класифікації (CPC та IPC) і ключові слова.

**(Murdick and Thomas, 2020).** Він вибирає близько 90 категорій CPC на найдеталізованішому рівні та 12 ключових термінів. Стратегія пошуку базується або на патентних класифікаціях, або на ключових термінах: патент, який не посилається на категорію CPC ШІ, але згадує такий термін, як «штучний інтелект» або «нечітка логіка», вважається ШІ. Крім того, САТ пропонує дані та індикатори за 7 підгалузями ШІ (машинне навчання тощо), використовуючи різні джерела патентних даних.

**Патентне відомство США (2021 та 2024).** Патентне відомство США спирається на методи машинного навчання для ідентифікації патентів на ШІ. Метод ландшафтного аналізу патентів на ШІ включає такі кроки:

- Визначення восьми підгалузей ШІ: «Машинне навчання», «Еволюційні обчислення», «Обробка природної мови (NLP)», «Зір», «Мовлення», «Обробка знань», «Планування та контроль» та «Апаратне забезпечення штучного інтелекту».

- Для кожної підгалузі ідентифікація «набіру початкових значень» (набір патентів, які оцінюються як такі, що належать до галузі, з огляду на їхню технологічну категорію), «антинабір початкових значень» (патенти, які оцінюються як такі, що не належать до галузі) та «набір розширення» (патенти, які є кандидатами на належність до галузі).

- Векторизація текстового вмісту цих патентів за допомогою трансформатора (у 2021 р. було застосовано техніку Word2Vec).

- Навчання класифікатора для кожної з восьми категорій на наборах даних.

- Ручна перевірка прогнозів класифікатора на тестовому наборі.

Отриманий набір даних про патенти ШІ (AIPD) був створений на основі патентів USPTO та оприлюднений для громадськості.

**Про генеративний штучний інтелект (2024).** Генеративний штучний інтелект (GenAI) наразі не має встановлених патентних класифікацій. Найбільш релевантна категорія CPC, тобто G06N3/045 (мережі автокодування), була запроваджена у 2023 р. і перекласифікація попередніх патентів в цю категорію вже триває.

**BOIV розробила спеціалізований двоетапний підхід для збору патентів GenAI (BOIV, 2024):**

- **Етап 1:** Класичний пошук у поєднанні з підказками ШІ дозволив отримати початковий набір даних з високою запам'ятовуваністю.

- **Етап 2:** Класифікатор двонаправлених представлень кодера для трансформаторів (BERT) удосконалив цей набір даних, покращивши точність, відрізнивши патенти GenAI від інших, що використовують методи ШІ.

Додатково, для створення набору даних було використано три методи:

- **Метод 1a:** Загальний пошук концепцій «генеративного ШІ» в різних класифікаціях.

- **Метод 1b:** Спеціальний пошук ключових технологій, наприклад, генеративно-змагальна мережа (GAN), моделі великих мов (LLM), моделі дифузії.

- **Метод 2:** Близько 100 підказок, згенерованих ШІ, спрямованих на концепції GenAI.

Набори даних були точно налаштовані за допомогою класифікатора BERT для категоризації патентів на патенти GenAI та патенти, що не належать до GenAI. Отриманий набір патентів GenAI доступний для завантаження.

**Інші існуючі таксономії.** В останні роки відомства інтелектуальної власності провели додаткові аналізи патентів. Відомство інтелектуальної власності Великої Британії (UKIPO) опублікувало звіт про патентування у сфері ШІ (UKIPO, спираючись на роботу, виконану BOIV та ОЕСР з подальшими висновками Відомства інтелектуальної власності Канади (CIPRO). У 2020 р. технології в сфері 4-ї промислової революції (4IR), як визначено Європейським патентним відомством (ЄПВ), включали компонент ШІ, який також був окреслений за допомогою комбінованого підходу пошуку за класом технологій та ключовими словами

**(ЕРО, 2020).** Зовсім недавно Відомство інтелектуальної власності Австралії опублікувало інтерактивну візуалізацію патентної аналітики,

зосереджену на патентах на ІІІ (IP Australia, 2024), використовуючи змішаний підхід пошуку за МПК та ключовими словами. Їхнє патентне ландшафтне моделювання детально охоплює конкретні технологічні області (наприклад, машинне навчання, експертні системи, апаратне забезпечення) та області застосування (наприклад, хімія, машинобудування, електротехніка, інструменти тощо). Також у звіті JPO, опублікованому у 2024 р., представлені останні тенденції у винаходах, пов'язаних зі ІІІ, що відрізняє еволюцію патентів, що захищають основні технології ІІІ, від застосувань ІІІ.

**Країни-лідери за патентами на ІІІ.** Походження винаходів у сфері ІІІ переважно зосереджене в кількох країнах: *винахідники з Китаю, США та Японії забезпечили 69% патентів на ІІІ у 2020-2023 роках.* Частка Китаю в загальній кількості патентів швидко зростала з початку 2010-х років, ставши світовим лідером з 30,6% усіх патентів на ІІІ у світі. Внесок США, які історично посідали перше місце, зменшився з 33% від усіх загальних патентів на ІІІ у 2010-2013 роках до 27% у 2020-2023 роках. У свою чергу, частка Японії значно знизилася з 29% на початку 2010-х років до 10,5% за останній період. Середній показник по ЄС-27 знизився на 4% порівняно з рівнем 2010-13 років, але залишається стабільним з 2014 року. Корея, Німеччина, Канада, Індія та Швеція є одними з небагатьох країн, які збільшили свою частку патентів на ІІІ з середини 2010-х років.

**Патентна активність у галузі ІІІ в Україні.** Згідно дослідження «Аналіз застосування штучного інтелекту в пріоритетних секторах та конкурентоспроможності України», проведеного Центром європейських політичних досліджень (CEPS) у 2024 р., Київ входить до десятки європейських країн за потенціалом розвитку ІІІ. [6] Індекс технологічної спорідненості (relatedness density) у столиці України становить 48,9%. Це означає, що майже половина технологій, розвинених у місті, потенційно пов'язані зі ІІІ. Київ розташований поруч із такими регіонами, як Відень (AT13, 52,08%) та Мадрид (ES30, 49,07%), і випереджає, наприклад, Амстердам (NL32, 47,78%) та Брюссель (BE10, 47,56%). Однак, за кількістю AI-патентів Україна поки що суттєво відстає від західноєвропейських країн. Згідно з базою даних OECD RegPat (2017–2022): у Німеччини — 3028 патентів, у Великої Британії — 1384, у Франції — 1122, тоді як в Україні — лише 9. Зазначимо, що це патенти, видані за системою РСТ. Цю різницю автори дослідження пояснюють високою вартістю патентування, складними

адміністративними процедурами та меншою інтеграцією України в ринок ЄС. Патенти зазвичай реєструють великі компанії з R&D-інфраструктурою, тоді як українська AI-екосистема більш динамічна у стартапах, публікаціях і прикладних розробках (Reface, Respeecher, Zibra AI, Osavul тощо).

У 2016-2021 рр. Україна отримала лише 16 патентів із 32 0878, або 0,005% світових патентів у сфері ШІ. Всього Україні як пріоритетній країні належить 130 патентів, з яких у період 2000–2021 рр. отримано 126 патентів. Найвища патентна активність в Україні спостерігалась у 2010–2014 рр. Серед цих патентів до «живих» (чинних) належать 68 од. , або 54%, до «мертвих» (через несплату зборів або термін дії яких закінчився) патентів – 55 од., або 43,7%. Аналізуючи патентну статистику ВОІВ по Україні, бачимо незначну кількість патентних заявок (патентні публікації за технологією), що підпадають під категорії «Комп’ютерні технології» та «ІТ-методи для управління». Так, між 1980-2018 рр. було опубліковано лише 740 таких заявок. Порівняно із загальною кількістю 58 845 опублікованих заявок, це становить 1,26%. [7] Безумовно, досвід США, КНР та ЄПВ у цій сфері заслуговує на увагу. У Керівництві щодо проведення експертизи ЄПВ щодо комп’ютерних програм ще у 2018 р. з’явився розділ, що стосується ШІ і машинного навчання (G-ІІЗ.3.1), які спочатку визначаються як обчислювальні моделі і алгоритми класифікації, кластеризації, регресії і зменшення розмірності. Згодом з’явилась удосконалена редакція. Україні необхідно імплементувати норми Керівництва США і ЄПВ щодо винаходів, реалізованих на комп’ютері. Адже нові Правила складання, подання та розгляду заявки на винахід та заявки на корисну модель, затверджені наказом Мінекономіки від 09.09.2024 №23301, не відображають цих аспектів. Лише комплексний підхід (зміни до законодавства та підзаконних актів, стимулювання і удосконалення експертизи заявок на винаходи) дасть можливість підвищити винахідницьку активність у цій сфері

**Висновки та обмеження.** Переглянута методологія ідентифікації патентів на ШІ, представлена в цій статті, розширює охоплення винаходів ШІ, включаючи найновішу термінологію, що дозволяє включати нещодавні інновації, такі як генеративний ШІ. Вона використовує комбінований підхід СРС та ключових слів, що робить її адаптованою до різних наборів патентних даних. Метод ідентифікації словника ШІ та класів СРС розроблений для забезпечення легкого та частого оновлення (приблизно двічі на рік).

Працюючи на агрегованому рівні СРС, він залишається стабільним, попри потенційні зміни класифікації, які зазвичай відбуваються на більш детальних рівнях. Такий методологічний підхід може бути застосований для охоплення розвитку нових технологічних галузей, захищених патентними документами в галузі ШІ та за її межами. Водночас нова методологія ідентифікації патентів на ШІ стикається з кількома обмеженнями. Однією зі значних проблем є своєчасність патентної інформації, оскільки часто трапляються адміністративні затримки в публікації патентів. Щоб пом'якшити цю проблему, методологія включає *підхід на основі прогнозування поточної ситуації* для оцінки патентів на ШІ за останні роки. Ще одним обмеженням є *наявність сліпих зон щодо незапатентованих інновацій*, які важко вирішити через брак доступних даних. Це підкреслює труднощі в охопленні повного обсягу винаходів ШІ, що виходять за рамки того, що офіційно запатентовано. Патентні показники можуть бути доповнені показниками, отриманими з інших джерел (наприклад, наукових статей, програмного забезпечення з відкритим кодом тощо), щоб забезпечити більш повну картину розвитку технологій ШІ. Представлена в статті методологія ідентифікації патентів на ШІ, дозволить здійснити подальші дослідження, пов'язані з впливом технологій ШІ, зокрема, на поширення знань, що виникають внаслідок технологій ШІ, на відносну цінність цих технологій для подальших розробок, на екосистему, на яку спираються технології ШІ, та на те, як ШІ перетинається з різними технологічними галузями.

### Список використаних джерел:

1. Андрощук Г. Аналітика ІВ: відстеження еволюції ШІ за допомогою патентів. URL: <https://yur-gazeta.com/publications/practice/zahist-intelektualnoyi-vlasnosti-avtorske-pravo/analitika-iv-vidstezhennya-evolyuciyi-shi-za-dopomogoyu-patentiv.html>
2. Кострубіцький Д. О. Ефективний пошук патентів – ключ до інновацій та конкурентоспроможності / УКРНОІБІ, 2024. URL: <https://nipo.gov.ua/wp-content/uploads/2025/03/web-efektyvnyi-poshuk-patentiv-web.pdf>
3. Андрощук Г. О., Кваша Т. К. Цифрові інструменти аналітики інтелектуальної власності. *Цифрова економіка: зростання ролі інтелектуальної власності*: збірник наук. праць / за науковою редакцією к.е.н., доц. Г.О. Андрощука,

д.е.н., проф. О.Б. Бутнік-Сіверського; НДІ ІВ НАПрН України. К.: Інтерсервіс, 2023. С.66-80.

4.Андрощук Г.О. Аналітика інтелектуальної власності як інструмент оцінювання результатів наукових досліджень // Теорія і практика оцінювання результатів наукових досліджень у сучасних реаліях: кол. моногр. / [Богданов В.Л., Журавель В.А., Кривцун І.В. та ін.]; за ред. В.Л. Богданова, Б.А. Маліцького; передмова А.Г. Загороднього. — Київ: Академперіодика, 2025. - 258 с. URL: [https://akademperiodyka.org.ua/wp-content/uploads/Theory-and-practice\\_uk.pdf](https://akademperiodyka.org.ua/wp-content/uploads/Theory-and-practice_uk.pdf)

5.IDENTIFYING EMERGING AI TECHNOLOGIES USING PATENT DATA A SEMI-AUTOMATED APPROACH TECHNICAL PAPER September 2025. URL: [https://www.oecd.org/content/dam/oecd/en/publications/reports/2025/09/identifying-emerging-ai-technologies-using-patent-data\\_5c8da861/d17e9a1a-en.pdf](https://www.oecd.org/content/dam/oecd/en/publications/reports/2025/09/identifying-emerging-ai-technologies-using-patent-data_5c8da861/d17e9a1a-en.pdf)

6. Українська столиця потрапила до десятки міст Європи за потенціалом розвитку ІІІ (інфографіка). URL: <https://news.finance.ua/ua/ukrains-ka-stolycya-potrapyla-do-desyatky-mist-yevropy-za-potencialom-rozvytku-shi-infohrafika>

7.Androshchuk G. O. PATENTING OF INVENTIONS CREATED WITH THE USE OF ARTIFICIAL INTELLIGENCE: REGULATORY ISSUES. *DIGITAL TRANSFORMATION IN UKRAINE: AI, METAVERSE, AND SOCIETY 5.0* URL: <https://dspace.nlu.edu.ua/jspui/bitstream/123456789/20364/1/Digital%20transformation%20in%20Ukraine.pdf>

**Олександр Бутнік-Сіверський**

доктор економічних наук, професор, академік  
Академії технологічних наук України, академік  
Української академії наук, головний науковий  
співробітник відділу комерціалізації об'єктів  
інтелектуальної власності НДІ ІВ НАПрН  
України

ORCID: <https://orcid.org/0000-0003-2492-231X>

## **РОЗВИТОК ІНТЕЛЕКТУАЛЬНОЇ ТЕХНІКО-ТЕХНОЛОГІЧНОЇ КОМП'ЮТЕРНОЇ ПЛАТФОРМИ ЦИФРОВІЗАЦІЇ**

Цифрові технології не відірвані від економіки так як поєднуються в єдину систему, де головними виступають *інтелектуалізація*, як процес інтелектуального цільового спрямування, нарощування потенціалу прогресивної інформаційно-комунікаційної технології та поширення сучасної комп'ютерної техніки, які, з нашого економіко-правового погляду, формують *інтелектуальну техніко-технологічну комп'ютерну платформу цифровізації*, що у загальному вигляді є системним прогресивним інноваційним утворення з великими можливостями. Така комп'ютерна платформа цифровізації виконує надзвичайно важливу функцію під неперервним впливом розвитку сфери інтелектуальної власності. У *стратегії цифрового розвитку інновацій до 2030 року* передбачено: «Україна зробить економічний стрибок за рахунок створення інноваційних продуктів, товарів та послуг. Ми розблокуємо приватну ініціативу для інновацій через розвиток людського капіталу, дерегуляції, створення високотехнологічних виробництв та якісних робочих місць, розвитку підприємництва та міжнародного співробітництва».[1].

Одночасно, впровадження цієї інтелектуальної техніко-технологічної комп'ютерної платформи цифровізації в сфери інтелектуальної власності стискається зі складною проблемою права власності на віртуальні об'єкти та віртуальні активи, що, на думку Д. Фартушної, «існує нагальна необхідність переосмислення традиційних уявлень про право власності, його об'єкти та порядок їх захисту з метою залучення віртуальних об'єктів, що забезпечить реальний захист користувачів від протиправних посягань на їхні віртуальні об'єкти» [2, с. 15]. Так, Л. Р. Майданик зазначає, що «оскільки взаємодія з

віртуальним об'єктом з боку користувача може бути обмежена розробником чи власником сайта, здійснення всієї сукупності повноважень права власності (володіння, користування, розпорядження) на свій розсуд щодо такого віртуального об'єкта є обмеженою» [3, с. 62]. І далі, Р. А. Майданник вважає, що «спостерігається тенденція до переоцінки соціального призначення права власності і речового права загалом, яка визначається з урахуванням доктрини ліберального праворозуміння, заснованої на ідеях соціальної солідарності і співробітництва, економічного та соціологічного аналізу права» [4, с. 227]. Одночасно, М. Рехлицький застерігає, «для того щоб об'єкт вважався віртуальним активом, він має відповідати трьом критеріям: наявність вартості (value), можливість до обігу в цифровому форматі (transferability) та можливість до його обміну на інші об'єкти цивільного права (payment or investment purpose)» [5]. До цієї дискусії додаймо позиції законодавця, який в Законі України «Про віртуальні активи» [6] визначив поняття «віртуальний актив – нематеріальне благо, що є об'єктом цивільних прав, має вартість та виражене сукупністю даних в електронній формі. Існування та оборотоздатність віртуального активу забезпечується системою забезпечення обороту віртуальних активів. Віртуальний актив може посвідчувати майнові права, зокрема права вимоги на інші об'єкти цивільних прав». Зазначене свідчить про очікуємі трансформаційні процеси переоцінки права власності і речового права в *людино-цифровому середовищі* в віртуальні об'єкти та віртуальні активи, які можуть посвідчувати майнові права, зокрема права вимоги на інші об'єкти цивільних прав.

З нашого погляду, розвиток інтелектуальної техніко-технологічної комп'ютерної платформи цифровізації в розрізі елементів цієї системи, обмежується різними темпами розвитку кожної складової стосовно опанування, впровадження, поширення в суспільно-економічне середовище, що потребує відповідної уваги урядових інституцій та їх координації. Про це свідчать різні процеси передачі інформації в утворюємо в країні, за нашим визначенням, *людино-цифрове середовище, як соціум з використанням цифрових технологій*, де домінують *цифрові права – права людини*, які полягають в праві людей на доступ, використання, створення і публікацію цифрових творів, доступ і використання комп'ютерів та інших електронних пристроїв, а також комунікаційних мереж, зокрема, до мережі Інтернет.

З практичної точки зору, науковці пропонують об'єднати і узагальнити об'єкти цифровізації в три групи, залежно від використання, ролі і участі у цифрових процесах [7, с.168]: а) засоби розрахунку (фінансові цифрові об'єкти); б) електронні товари (товарні цифрові об'єкти); в) електронні послуги (базові, платформові цифрові об'єкти). Фактично всі суспільно-економічні відносини, правочини переходять на цифровий рівень під загальним терміном «віртуальні цифрові активи».

З загальної позиції С. Щеглюк, – «феномен «цифрової платформи», явище «платформізації», яке стало можливим завдяки появі нових бізнес-моделей, транскордонних процесів, мережевих ефектів, моделей спільного споживання, потенціалу фінансових технологій, скорочення циклів інвестування, трансформації торгових, виробничих і логістичних ланцюжків, життєвого циклу цифрових активів та відкритих інновацій» [8]. У той же час, згідно з ООН використання платформної економіки несе ряд наступних проблем і викликів [9, с.24]: 1) подолання «цифрового розриву» є найважливішим компонентом інклюзивного розвитку платформної економіки; 2) нормативні акти часто відображають статус-кво і перешкоджають інноваціям або не допускають їх; 3) для процвітання платформної економіки необхідно забезпечити безпеку даних, стандарти даних і конфіденційність даних; 4) платформна економіка часто передбачає появу домінуючих платформ і часткових монополій для ефективної роботи, що вимагає нових підходів до забезпечення адекватної конкуренції; 5) платформи підвищують значимість нематеріальних активів, часто накопичуються в розвинених країнах, що викликає занепокоєння з приводу рівності та інтеграції. Таким чином, можна підсумувати, що цифрові платформи знижують операційні витрати, зокрема, витрати на тріангуляцію, передачу та перевірку достовірності операції.

Головною ознакою цифрової платформи (платформізації) є «такий тип модульної системи, який розділений на велику кількість відносно стабільних компонентів ядра з відносно низьким розмаїттям, великий набір периферійних доповнень з великим розмаїттям і набором стандартизованих інтерфейсів» [10, с.29]. На думку фахівців Масачусетського технологічного університету «цифрова платформа – це забезпечена високими технологіями бізнес-модель, яка створює вартість, полегшуючи обмін між двома або більшою кількістю взаємозалежних груп учасників». За своїм технічним змістом платформа – це

система «алгоритмізованих взаємовигідних відносин значної кількості незалежних учасників галузі економіки (або сфери діяльності), які здійснюються в єдиному інформаційному середовищі, що призводить до зниження трансакційних витрат за рахунок застосування пакету цифрових технологій роботи з даними і зміни системи поділу праці».[8].

Діяльність платформ, в основному, здійснюється на основі мережі Інтернет, де використовується мережевий принцип роботи, який полягає у тому, що її оператор надає всім зацікавленим особам послуги онлайн комунікації, застосовуючи ранжування або індексацію контенту, товарів або послуг, пропонованих або надаваних онлайн третіми особами; або об'єднання декількох осіб з метою продажу товарів, надання послуг або обміну контентом, товарами або послугами.

На наш погляд, *цифрові платформи в сфері інтелектуальної власності розглядаються як техніко-технологічні комп'ютерні платформи цифровізації* – є системними прогресивними інноваційними інструментами з великими можливостями, у яких одночасно розвиваються та поєднуються комп'ютерна техніка, технологія, телекомунікація та цифрові програмні методи їх використання. Такі цифрові платформи в сфері інтелектуальної власності, як загальні цифрові платформи, використовують відповідний тип модульної системи, який розділений на велику кількість відносно стабільних компонентів ядра з відносно низьким розмаїттям, великий набір периферійних доповнень з великим розмаїттям і набором стандартизованих інтерфейсів. Діяльність цифрових платформ в сфері інтелектуальної власності, в основному, здійснюється на основі мережі Інтернет, де використовується мережевий принцип роботи. Така цифрова платформа в сфері інтелектуальної власності виконує надзвичайно важливу функцію, а саме під неперервним впливом розвитку сфери інтелектуальної власності, як результат, створюються цифрові об'єкти інтелектуальної власності, залучаючи у подальшому штучний інтелект, які за своєю природою є нематеріальними.

З нашої позиції, цифрова платформізація в сфері інтелектуальної власності за відповідним напрямом функціонування, враховує творчу діяльність людини-винахідника або програміста з використанням цифрових технологій, який розробляє програмний продукт, придатний для виконання цифрових завдань. Наприклад, **Міністерство економіки України з 10 березня 2025 року запустило оновлену процедуру розміщення**

національного переліку веб-сайтів, які порушують права інтелектуальної власності, на міжнародній платформі «WIPO ALERT» Всесвітньої організації інтелектуальної власності (BOIV). Механізм «WIPO ALERT» функціонує як захищена онлайн-платформа, де уповноважені органи країн-членів BOIV завантажують списки вебсайтів та додатків, що порушують авторські права. Рекламодавці, рекламні агентства та їхні технічні партнери можуть стати авторизованими користувачами платформи, отримуючи доступ до цих списків, щоб уникати розміщення реклами на піратських ресурсах, таким чином захищаючи свої бренди та утримуючись від порушення авторських прав. [11]

З нашого погляду, цифрова платформізація в сфері інтелектуальної власності за відповідним напрямом функціонування формує реальні та віртуальні людино-цифрові системи, які мають свої особливості, до яких відносять поліфункціональність та індивідуальну функціональність. Тому виникає потреба нового погляду на **цифрові права** людини з позиції реального та віртуального або матеріального та відносно абстрактного середовища в умовах використання цифрових інструментів, які уможлиблюють застосування різних цифрових технологій.

Цифрова платформізація в сфері інтелектуальної власності, що діє в реальній людино-цифровій системі, враховує результат творчої діяльності людини-винахідника або програміста, а сьогодні все більше штучний інтелект з використанням цифрових технологій, і як результат такої діяльності відтворюють цифрові технологічні платформи за відповідним поліфункціональним напрямом застосування: освіта, нерухомість, торгівля, фінанси, транспорт тощо, що свідчить про її поліфункціональність.

#### **Список використаних джерел:**

1. WINWIN. Глобальна інноваційна стратегія України URL: <https://winwin.gov.ua>.
2. Фартушна Д. Наукова робота на тему: «Віртуальна власність: проблеми та перспективи» 2020. С. 1-34. URL:<http://dspace.onua.edu.ua/bitstream/handle/11300/12666/Nova%20era.pdf?sequence=1>.
3. Майданик Л. Р. Віртуальний об'єкт як виклик для класичних підходів у речовому праві. /Теорія і практика інтелектуальної власності. 2019. № 2. С. 59-64.

4. Майданик Р. А. Речове право. Загальні положення. Спогади про Людину, Вченого, Цивіліста-епоху (до 85-річчя від Дня народження академіка Ярославни Миколаївни Шевченко) за заг. ред. Р.О. Стефанчука. К.: АртЕк, 2017. С. 226-249.

5. Рехлицький Микола. Віртуальні активи: що це таке і навіщо нам їх законодавче регулювання? *Юридична газета*. №10 (716). 27 травня 2020. URL: <https://jur-gazeta.com/publications/practice/informaciyne-pravo-telekomunikaciyi/virtualni-aktivi-shcho-ce-take-i-navishcho-nam-yih-zakonodavche-regulyuvannya.html>.

6. Про віртуальні активи. Закону України № 2074-IX – не набрав чинності, поточна редакція від 15.11.2024 року. URL: <https://zakon.rada.gov.ua/laws/show/2074-20#Text>.

7. Крочак О.І. Активи цифрової трансформації як об'єкти обліку/ Оксана Іванівна Крочак, Олена Миколаївна Матрос, Світлана Олексіївна Михайловина //Наукові перспективи № 4(22) 2022. – URL: <http://perspectives.pp.ua/index.php/np/article/view/1472/1470>.

8. Щеглюк Світлана. Морфологія цифрової економіки: особливості розвитку та регулювання цифрових технологічних платформ (науково-аналітична записка) / ДУ «Інститут регіональних досліджень ім. М.І. Долишнього НАН України». 2019 р.- URL: <https://ird.gov.ua › irdp/e 20190301.pdf>.

9. Дубель, М. Особливості розвитку цифрових платформ та їх вплив на світову економіку. *Таврійський науковий вісник. Серія: Економіка*, (7), 2021. С. 17-26. <https://doi.org/10.32851/2708-0366/2021.7.2>

10. Січкаренко К.О. Цифрові платформи: підходи до класифікації та визначення ролі в економічному розвитку / Причорноморські економічні студії. - 2018. - Вип. 35(2). - С. 28-32.

11. Мінекономіки почало використовувати оновлену процедуру «WIPO ALERT» для захисту прав інтелектуальної власності. – URL: <https://me.gov.ua/News/Detail/bb6571b8-4303-4d15-8550-c1857e269631?lang=uk-UA&title=wipoAlert> (дата звернення 10.03.2025 р.)

**Олеся Андрущенко**

доктор філософії за спеціальністю «Право»,  
асистентка кафедри філософії, Національний  
юридичний університет імені Ярослава  
Мудрого

ORCID: <https://orcid.org/0009-0004-1759-1166>

## **ФІЛОСОФСЬКО-ПРАВОВІ ЗАСАДИ ЛЕГІТИМАЦІЇ ШІ В СИСТЕМІ ПОЛІТИКО-ПРАВОВИХ ВІДНОСИН СУЧАСНОЇ УКРАЇНИ**

У ХХІ столітті відбувається докорінна трансформація способу функціонування суспільства, держави та права під впливом процесів цифровізації, автоматизації й розвитку ШІ. Йдеться не лише про правове регулювання функціонування ШІ, а про глибше – про усвідомлення його місця в системі суспільних відносин, визначення меж автономії, відповідальності та взаємодії з людиною і державою. У цьому контексті Україна, здійснюючи власний цифровий прорив, постає перед складним завданням – гармонізувати інноваційні технології з основоположними принципами права, демократії та людської гідності.

Як підкреслює О. Баранов, ключовою проблемою сучасного правового дискурсу є відсутність єдиного визначення поняття «штучний інтелект», що призводить до нормативної невизначеності у сфері його правового статусу [2]. Залежно від підходу, ШІ може тлумачитися як система обробки інформації, суб'єкт комунікації або навіть потенційний носій юридично значущих дій. Така багатозначність терміну унеможливорює однозначне правове позиціонування ШІ, що, своєю чергою, ускладнює його легітимацію в межах політико-правового простору України.

Легітимація як фундаментальна категорія політико-правової філософії передбачає визнання правомірності влади, рішень та інституцій суспільством. У класичному розумінні, починаючи з праць М. Вебера [3], легітимність ґрунтується на довірі громадян до влади, переконанні в її справедливості та правовій обґрунтованості. Однак у цифрову добу, коли дедалі більше владних і управлінських функцій виконується автоматизовано, питання довіри переноситься із людини чи інституції на технологію. Виникає феномен «цифрової легітимності», який можна тлумачити як визнання суспільством

правомірності рішень, згенерованих ШІ, які дедалі частіше впливають на політичний процес, управління та правозастосування.

Філософсько-правовий аспект легітимації ШІ полягає у необхідності поєднання двох рівнів осмислення: нормативного – тобто створення правових засад для його регулювання, та ціннісного – забезпечення відповідності дії алгоритмів основоположним правам і свободам людини. О. Косілова та Х. Солодовнікова зазначають, що між людиною та ШІ виникає принципова колізія – з одного боку, необхідність гарантувати безпеку та права людини, а з іншого – забезпечити розвиток інновацій і технологій [5, с. 57]. Вони наголошують, що саме права людини повинні залишатися мірилом правомірності застосування ШІ, адже без цього існує ризик підміни ціннісних орієнтирів і послаблення ролі особистості як центрального суб'єкта права.

ШІ у сучасних умовах уже не є лише технічним інструментом – він стає новим елементом політико-правового простору, який має потенціал перетворитися на самостійного суб'єкта регулятивних відносин, що особливо помітно у контексті державних ініціатив України, спрямованих на розбудову цифрової державності та агентського ШІ. 2025 рік позначився як період інституційного становлення національної моделі використання ШІ: запущено державного цифрового асистента «Дія.АІ» [1], що фактично означає делегування частини владних і комунікаційних функцій цифровому асистенту, який взаємодіє з громадянином від імені держави, а також розбудовується освітня екосистема «Мрія» [6]. Зазначені ініціативи є не лише технологічним проривом, а й моментом глибокого переосмислення правових, політичних та філософських основ легітимності у взаємодії людини, держави та технологій.

С. Іващенко підкреслює, що сучасна українська правова система перебуває лише на початковому етапі формування спеціальних правових норм, які визначають статус і відповідальність систем ШІ [7]. Дослідник пропонує впровадити модель “розподіленої відповідальності”, за якої обов’язки між розробником, користувачем і системою ШІ розподіляються залежно від рівня автономності останньої. Такий підхід може стати першим кроком до створення нової парадигми правосуб'єктності, яка не суперечитиме антропоцентричній природі права, але й визнаватиме особливості взаємодії людини та технології.

Філософсько-правова рефлексія цього процесу має базуватися на ідеї, що право – це не лише сукупність норм, а й форма соціальної комунікації, що

виражає морально-політичні цінності спільноти. Коли до цієї системи входить ІІІ, правова комунікація змінює свій характер. ІІІ починає моделювати поведінку людей, прогнозувати рішення судів, автоматизувати видачу документів, здійснювати ідентифікацію осіб. Усе це веде до появи нового типу правової взаємодії, коли громадянин не взаємодіє безпосередньо з державним представником, а з його цифровим еквівалентом.

У цьому контексті важливо осмислити співвідношення понять легітимність і довіра. Традиційно довіра до держави ґрунтується на її спроможності діяти справедливо, передбачувано і відповідально. Однак у цифровому середовищі критерії довіри трансформуються: вони починають залежати від якості даних, рівня кіберзахисту, можливості аудиту коду та етичної відповідальності розробників. Таким чином, легітимність державних рішень у майбутньому може все більше залежати не від політичних актів, а від технологічної архітектури.

Крім того, М. Дубняк зазначає, що в епоху ІІІ функції права змінюються: від традиційної регулятивної до превентивно-прогностичної, оскільки право має не лише реагувати на наслідки використання технологій, а й передбачати їх [4], що вимагає від законодавця роботи над формуванням етичних і правових рамок, у яких технологічний прогрес не суперечить базовим принципам демократії, свободи та справедливості.

Можна констатувати, що філософсько-правова модель легітимації ІІІ у в Україні має спиратися на три фундаментальні принципи. По-перше, принцип відповідальної автономії – кожна програмна модель повинна функціонувати в межах, які встановлює суспільство та закон, із чітко визначеною відповідальністю розробників і користувачів. По-друге, принцип прозорості та пояснюваності рішень, що забезпечує довіру та контроль з боку громадян. По-третє, принцип етичного супроводу, який передбачає постійне оцінювання наслідків використання ІІІ для захисту гідності, свободи та справедливості осіб, які використовують ІІІ. Саме така основа здатна сформувати нову, нормативно й морально легітимну парадигму політико-правових відносин у добу цифрової державності.

Таким чином, ІІІ стає не просто інструментом цифровізації, а чинником перетворення легітимації влади, права і соціального порядку. Легітимність у цифрову епоху перестає бути суто політичним феноменом – вона набуває технологічно-правового змісту, де довіра до ІІІ, етична відповідальність

розробника і нормативний контроль держави формують новий тип соціального договору. Для України, яка прагне увійти до топ-3 держав світу за розвитком та інтеграцією ІІІ у публічний сектор, цей процес означає необхідність глибокого філософсько-правового переосмислення самої природи влади, права та людини в умовах віртуалізації реальності. Саме через формування власної концепції легітимації ІІІ у системі політико-правових відносин Україна здатна закласти інтелектуальний фундамент майбутньої демократичної, безпечної та технологічно зрілої цифрової держави.

### Список використаних джерел:

1. AI-асистент уже на порталі Дія — першими у світі отримуйте послугу з допомогою ІІІ. *Дія*: веб-сайт. URL: <https://diia.gov.ua/news/ai-asystent-uzhe-na-portali-diia-pershymu-u-sviti-otrymuite-posluhu-z-dopomohoiu-shi> (дата звернення: 19.10.2025).
2. Баранов О. А. Визначення терміну «штучний інтелект». *Інформація і право*. 2023. № 1(44). DOI: [https://doi.org/10.37750/2616-6798.2023.1\(44\).287537](https://doi.org/10.37750/2616-6798.2023.1(44).287537) (дата звернення: 19.10.2025).
3. Вебер М. Три чисті типи легітимного панування. *Соціологія. Загальноісторичні аналізи. Політика*. Київ. 1998. С. 157-172. URL: <http://litopys.org.ua/weber/wbs07.htm> (дата звернення: 19.10.2025).
4. Дубняк М. В. Функції права в епоху штучного інтелекту. *Інформація і право*. 2025. № 2(53). URL: <https://il.ippi.org.ua/article/view/334044> (дата звернення: 19.10.2025).
5. Косілова О. І., Солодовнікова Х.К. Права і свободи людини і громадянина V.S. штучний інтелект: проблемні аспекти. *Інформація і право*. 2020. № 4(35). С. 56-66. URL: [https://ippi.org.ua/sites/default/files/7\\_18.pdf](https://ippi.org.ua/sites/default/files/7_18.pdf) (дата звернення: 19.10.2025).
6. Мрія. *Мрія*: веб-сайт. URL: <https://mriia.gov.ua/app> (дата звернення: 19.10.2025).
7. Іващенко С. Правові засади регулювання штучного інтелекту в Україні: сучасний стан та перспективи розвитку. *Адміністративне право і процес*. 2025. № 2. С. 24-32. URL: <https://applaw.net/index.php/journal/article/download/856/717/> (дата звернення: 19.10.2025).

**Яна Лукашова**

Навчально-науковий інститут права КНУ імені  
Тараса Шевченка  
студентка 1 курсу ОС «Магістр»

## **ПРАВОВЕ РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У СФЕРІ НАЦІОНАЛЬНОЇ БЕЗПЕКИ І ОБОРОНИ**

Питання штучного інтелекту, визначення його правового статусу вже не перший рік підіймається у наукових колах. Це проблема глобального рівня. Поява все нових і нових механізмів змушує змінювати звичні для людства встановлені правила, за чим просто не встигають законодавці. Штучний інтелект проник у кожную сферу суспільного життя, навіть у сфери національної безпеки і оборони, що ускладнює його регулювання. І саме Україні довелося використовувати передові технології на полі бою. Мало де ще можна випробувати можливості ШІ в таких умовах. Це дозволяє спостерігати за розвитком машин і наслідками такого застосування у реальному часі, що дозволяє оперативно змінювати їх та вносити корективи.

На території нашої держави активно випробовуються сучасні системи іноземного походження. Так як однією з ключових можливостей штучного інтелекту є автоматизований аналіз різних типів цифрових даних — відео, аудіо, зображень та текстів, він здатний розпізнавати обличчя, визначати місцезнаходження за супутниковими знімками і відновлювати хронологію подій.

Одним із перших таких помітних прикладів стало впровадження українським стартапом програми Reface, здатної розпізнавати війська за супутниковими знімками. Для вдосконалення алгоритмів розробники залучали громадян, які надсилали актуальні дані з різних регіонів країни, що суттєво допомогло Збройним силам у виявленні ворога та диверсійних груп.

Ще одним вагомим прикладом є використання з березня 2022 року системи розпізнавання обличчя Clearview AI, наданої Україні Міністерством оборони США. Вона дозволяє українським службам ідентифікувати осіб на блокпостах, перевіряти їх на причетність до ворожих дій та навіть встановлювати особи загиблих [1, с. 411].

Окрім цього, українські війська застосовують ШІ для розмінування, для автоматичного наведення систем ППО, БПЛА, розпізнавання та враження цілей, автономних засобів евакуації, навіть без втручання людини [2]. Це значно підвищило ефективність безпекових і гуманітарних операцій під час війни. Але чи врегульовано таке використання технологій законодавством України?

Національна система лише дотично охоплює це питання у нормативно-правових актах загального характеру. Сюди можна віднести «Концепцію розвитку штучного інтелекту в Україні», у якій є окремих розділ «Оборона» [3]; «Дорожню карту регулювання штучного інтелекту в Україні», яка акцентує на тому, що робота над нормативним полем для ШІ важлива для розвитку країни й швидкого руху у сфері військових технологій [4]; та «Біла книга з регулювання ШІ в Україні: бачення Мінцифри», яка використовується як аналітичний матеріал і містить пропозиції щодо підходів керування ШІ [5]. Але водночас, в останньому документі автори відмовляються пропонувати регулювання систем ШІ у сфері оборони. Вони аргументовують це тим, що домогтися належного впровадження таких норм можна лише на міжнародному рівні прийняттям конвенції або внесенням змін до вже існуючої [5, с. 10].

Дійсно, одностороннє закріплення норм національним законодавством виглядає досить не вигідно, так як держава-агресор ніяким чином не впровадить таке ж саме регулювання. І поки ми будемо дотримуватись певних правил, наш супротивник може опинитися на крок вперед. Тому, можливо, відсутність законодавчих норм, що стосуються штучного інтелекту, лише на користь?

Доцільно буде розвивати саме практичний аспект інтеграції міжнародних норм у державну систему, рекомендації щодо такого впровадження. Досліджувати моделі ШІ в інших країнах і спиратися на їх досвід.

Але незважаючи на відсутність чіткого врегулювання, у військовій сфері завжди потрібно враховувати норми міжнародного гуманітарного права та прав людини, щоб уникнути використання ШІ у спосіб, який би порушував такі норми, дискримінував за расовою, національною, релігійною чи будь-якою іншою ознакою, аби забезпечити відповідність міжнародним стандартам.

Окремим завданням у цій галузі постає встановлення вимог до якості, надійності та безпеки технологій. Такі вимоги переважно розробляються з урахуванням міжнародних актів, як наприклад STANAG (Standardization Agreements) [6, с. 95]. Це стандарти НАТО, які спрямовані на досягнення сумісності держав-членів у різних сферах, зокрема регулювання ІІІ-систем у бойових умовах. Водночас потреба у створенні власних норм залишається актуальною, так як це дозволить ефективно контролювати розробку й застосування ІІІ, гарантувати його відповідність національним інтересам і мінімізувати ризики зловживань чи технічних помилок.

Потрібно вибудувати таку систему, яка максимально буде захищена від кібератак ворога, і здатна забезпечувати безперервну роботу ІІІ, навіть у разі спроб несанкціонованого втручання. Варто впроваджувати багаторівневий захист даних, використання алгоритмів шифрування, багатофакторної автентифікації та моніторинг збоїв у роботі. Крім того, така система повинна передбачати механізми оперативного втручання людини для контролю рішень ІІІ у критичних ситуаціях, щоб гарантувати безпеку та відповідність прийнятих рішень законодавчим і етичним нормам.

Залишається відкритим питання відповідальності за дії ІІІ, включаючи його помилки, непропорційні рішення та можливі негативні наслідки. Оскільки існує загроза життю людей, не можна допустити, щоб важливі стратегічні рішення приймав ІІІ. Він є лише інструментом в руках держави, за допомогою якого, в умовах правильного використання, можна досягти успіхів легшим та швидшим шляхом.

На мою думку, варто більше зосередити увагу на забезпеченні вимог до систем ІІІ, створенні надійних механізмів захисту задля уникнення кібератак у такій надважливій сфері для сьогодення. Регулювати такі системи варто шляхом прийняття конвенцій для того, щоб на міжнародному рівні були єдині стандарти роботи ІІІ для кожної держави. Вже після цього такі норми можна імплементувати до національного законодавства та сміло застосовувати у всіх галузях, включаючи національну безпеку та оборону.

### **Список використаних джерел:**

1 - Томіна В.Ю. Особливості використання штучного інтелекту в умовах воєнного стану на території України. *Юридичний науковий електронний журнал*, № 6/2023, с. 410-413. URL: [http://lsej.org.ua/6\\_2023/95.pdf](http://lsej.org.ua/6_2023/95.pdf).

2 – Штучний інтелект «на службі» в українських військових. URL: <https://pechera.info/>.

3 - Концепція розвитку штучного інтелекту в Україні. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#top> (дата звернення 28.10.2025 р.)

4 - Дорожня карта регулювання штучного інтелекту в Україні. URL: <https://thedigital.gov.ua/news/technologies/regulyuvannya-shtuchnogo-intelektu-v-ukraini-prezentuemo-dorozhnyu-kartu>.

5 - Біла книга з регулювання ШІ в Україні: бачення Мінцифри. URL: <https://backend.hromada.gov.ua/storage/uploads/files/research/bila-kniga-z-regulyuvannya-si-v-ukrayini-bacennya-mincifri>.

6 – Завадський А.А. Адміністративно-правове регулювання штучного інтелекту у сфері оборони: міжнародний вимір та національні пріоритети. *Збірник наукових праць ХНПУ імені Г. С. Сковороди «ПРАВО»*, № 40/2024, с. 87-97. URL: <http://journals.hnpu.edu.ua/index.php/law/article/view/16307>.

**Вадим Гаращенко**

аспірант відділу теорії держави і права,

Інститут держави і права ім. В.М. Корецького

ORCID: <https://orcid.org/0009-0004-8389-0869>

## **ПРАВОВЕ РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ЄВРОПЕЙСЬКОМУ СОЮЗІ: ТЕОРЕТИКО-МЕТОДОЛОГІЧНІ ТА ПРАКТИЧНІ АСПЕКТИ**

Прийняття Закону про штучний інтелект (*Artificial Intelligence Act* – далі AI Act), ухваленого Європейським парламентом 13 березня 2024 р., стало ключовою подією у формуванні правової політики Європейського Союзу у сфері цифрового регулювання. Закон має на меті забезпечити захист фундаментальних прав, демократії, верховенства права та стійкості довкілля від потенційно шкідливих застосувань штучного інтелекту, а також створення умов для інноваційного розвитку Європи у цій сфері [1].

AI Act запроваджує ризик-орієнтовану модель регулювання, відповідно до якої правові зобов'язання сторін (розробників, імпортерів, розповсюджувачів, користувачів) залежать від рівня ризику, який створює конкретна система ШІ. Закон визначає систему ШІ як машинну систему, що діє з різними рівнями автономії, здатна до адаптивності після розгортання та робить висновки на основі вхідних даних для формування результатів, які впливають на фізичне або віртуальне середовище. Це визначення відображає прагнення законодавця до технологічної нейтральності й мінімізації ризику швидкої застарілості норм [2].

Суттєвим нововведенням AI Act є заборона низки практик, що визнаються несумісними з правами людини: підсвідомі маніпулятивні техніки, експлуатація вразливостей осіб, системи соціального скорингу, використання біометричної ідентифікації в режимі реального часу, а також передавання емоцій у робочих або освітніх середовищах, крім випадків медичних чи безпекових потреб. Особливо дискусійною стала заборона систем дистанційної біометричної ідентифікації в режимі реального часу для правоохоронних органів, окрім обмежених винятків (розшук жертв злочинів, запобігання терактам тощо). Регламент також розрізняє два типи високоризикових систем ШІ: ті, що інтегровані у продукцію, регульовану

спеціальним законодавством ЄС (авіація, транспорт, медичні пристрої тощо); ті, що застосовуються у чутливих сферах – освіті, працевлаштуванні, кредитуванні, правосудді, міграційній політиці та демократичних процесах процесі [3]. На імпортерів і дистриб'юторів покладено обов'язок перевіряти відповідність систем вимогам AI Act та наявність маркування CE, що засвідчує проходження процедури оцінки відповідності директивам ЄС [4].

Таким чином, AI Act формує уніфіковану правову архітектуру, спрямовану на забезпечення етичного, безпечного та людиноцентричного розвитку ШІ, а також на підвищення довіри до технологій і конкурентоспроможності європейського ринку.

Хоча AI Act ще не набув повного застосування, судова практика ЄС уже демонструє підходи, що формують його тлумачення. Так, у справі C-634/21 SCHUFA Holding AG (2023) Суд ЄС визнав, що автоматизована система кредитного скорингу, яка має вирішальний вплив на укладення договору, є формою автоматизованого ухвалення рішення у розумінні ст. 22 GDPR. Кредитний скоринг — це алгоритмічна оцінка кредитоспроможності особи, що здійснюється без безпосередньої участі людини. Суд підкреслив, що навіть коли остаточне рішення формально приймає людина, але воно повністю залежить від результатів алгоритму, відповідальність за дотримання принципів прозорості, справедливості та людського контролю зберігається. Це рішення Суду ЄС фактично закладає підґрунтя для майбутнього тлумачення та застосування положень AI Act щодо прозорості та пояснюваності рішень систем штучного інтелекту [5].

Подальший розвиток практики Суду ЄС відбувся у справі C-203/22 Dup & Bradstreet Austria GmbH (2025), де Суд уточнив зміст права суб'єкта даних на отримання “суттєвої інформації про логіку обробки” відповідно до ст. 15(1)(h) GDPR. Суд наголосив, що навіть у випадках автоматизованого профілювання з використанням алгоритмів скорингу суб'єкт даних має право отримати зрозуміле пояснення принципів і процедур, що визначають роботу алгоритму. Посилання на комерційну таємницю не може бути підставою для повної відмови у доступі: у разі конфлікту інтересів контролер повинен передати спірну інформацію компетентному наглядовому органу або суду для зважування інтересів сторін [6].

Це рішення розвиває висновки, сформульовані у справі SCHUFA Holding AG, та поглиблює тлумачення принципу прозорості, який надалі буде

ключовим для застосування AI Act до високоризикових систем штучного інтелекту.

Аналогічний підхід до оцінки правомірності обробки даних простежується і в судовій практиці України. Так, у справі № 757/22027/22-ц Верховний Суд дійшов висновку, що зберігання персональних даних банком протягом строку, визначеного законом, є законним, оскільки має легітимну мету – забезпечення національної безпеки та виконання вимог фінансового моніторингу. Такий підхід відображає ті самі критерії пропорційності й законності, які закладені у практиці Суду Європейського Союзу (зокрема, у справах SCHUFA Holding AG та Dun & Bradstreet), і становить національний аналог доктрини «обґрунтованої обробки даних», що лежить в основі регулювання штучного інтелекту високого ризику за AI Act [7].

Актуальним етапом розвитку європейської практики у сфері застосування штучного інтелекту є справа C-806/24 Yettel Bulgaria, що перебуває на розгляді Суду Європейського Союзу за преюдиційним запитанням Софійського районного суду (Болгарія). Це перше звернення національної юрисдикції, у якому прямо порушено питання тлумачення положень Регламенту (ЄС) 2024/1689 – Artificial Intelligence Act (AI Act) у контексті споживчих договірних відносин. Предметом спору є автоматизоване виставлення рахунків системою мобільного оператора Yettel Bulgaria EAD у рамках договору на надання телекомунікаційних послуг. Позивач (споживач) стверджує, що рахунки були сформовані алгоритмічно, без участі людини, а постачальник послуг не розкрив логіку обчислення вартості, тобто не забезпечив прозорості та можливості перевірки автоматизованого рішення.

Суд Болгарії звернувся до Суду ЄС із низкою питань щодо тлумачення статті 86(1) AI Act та суміжних положень Директив 2011/83/ЄС і 93/13/ЄС, зокрема: чи має споживач право знати, за яким алгоритмом і за якими критеріями формується автоматично прийняте рішення (рахунок, нарахування, штрафи); чи поширюється дія AI Act на системи, що використовуються приватними суб'єктами у договорах між суб'єктом господарювання та споживачем; чи повинен суд мати можливість перевірити “чорну скриньку” або вихідний код алгоритму для забезпечення права на ефективний судовий захист; чи підпадають автоматично згенеровані документи під вимогу ясності та зрозумілості, передбачену директивами ЄС про захист прав споживачів.

Нині рішення у справі C-806/24 Yettel Bulgaria ще не ухвалено, однак її розгляд має важливе значення, оскільки саме від позиції Суду ЄС щодо застосування AI Act до приватноправових відносин залежить визначення того, які принципи прозорості та людського контролю мають забезпечуватися, зокрема у договорах між суб'єктом господарювання та споживачем, а також якою мірою алгоритмічні рішення можуть бути піддані судовій перевірці [8].

Таким чином, Акт ЄС про штучний інтелект (AI Act) формує нову модель правового регулювання технологій, що ґрунтується на поєднанні превентивного та ризик-орієнтованого підходів. Судова практика ЄС вже закладає методологічні орієнтири для його тлумачення, насамперед у ключових питаннях: право на пояснення алгоритмічних рішень, забезпечення недискримінаційності ШІ та визначення меж державного контролю за використанням біометричних технологій. Значення цього дослідження виходить за межі галузевого аналізу і має теоретико-правовий вимір, оскільки окреслює трансформацію державного регулювання у цифрову добу, перегляд меж публічного й приватного права, а також формування нової парадигми взаємодії права, технології та етики.

Отже, AI Act не лише визначає нормативні засади функціонування ринку штучного інтелекту, а й сприяє переосмисленню ролі держави у гарантуванні цифрових прав людини та розвитку сучасної теорії держави і права.

### **Список використаних джерел:**

1. EU AI Act: first regulation on artificial intelligence. URL: <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>.
2. The pre-final text of the EU's AI Act leaked online. URL: <https://www.whitecase.com/insight-alert/pre-final-text-eus-ai-act-leaked-online>.
3. Artificial intelligence act: Council and Parliament strike a deal on the first rules for AI in the world. URL: <https://www.consilium.europa.eu/en/press/press-releases/2023/12/09/artificial-intelligence-act-council-and-parliament-strike-a-deal-on-the-first-worldwide-rules-for-ai/>.
4. Matt Kosinski. What is the EU AI Act? URL: <https://www.ibm.com/topics/eu-ai-act>.
5. Рішення Суду справедливості ЄС (Перша Палата) від 7 грудня 2023 року, SCHUFA Holding and Others (Звільнення від решти боргів), C-26/22 та C-64/22. URL:

<https://curia.europa.eu/juris/document/document.jsf?text=&docid=280428&pageIndex=0&doclang=EN&mode=lst&dir=&occ=first&part=1&cid=7913974>.

6. Рішення Суду справедливості ЄС (Перша палата) від 27 лютого 2025 року, Dun & Bradstreet Austria GmbH, C-203/22. URL: <https://eur-lex.europa.eu/legal-content/BG/TXT/PDF/?uri=CELEX:62022CJ0203>.

7. Постанова Верховного Суду від 29 листопада 2023 року у справі № 757/22027/22-ц (провадження № 61-13409св23). URL: <https://reyestr.court.gov.ua/Review/115423807>.

8. Преюдиційне запитання Софійського районного суду (Болгарія) у справі C-806/24, Yettel Bulgaria EAD v. FB, C/2025/1080. URL: <https://eur-lex.europa.eu/eli/C/2025/1080/oj>

**Владислав Варинський**

кандидат політичних наук, доцент;

доцент кафедри філософії та україністики

Національного університету «Одеська морська академія»

ORCID: <https://orcid.org/0000-0001-5837-6201>

## **ОСНОВНІ КРИТЕРІЇ ЗАБОРОН ВИКОРИСТАННЯ ІІІ ЗА АІ АСТ**

Тенденції щодо страхів перед використанням штучного інтелекту (далі - ІІІ) через визнання його переваг над когнітивними здатностями людини спостерігаються у світі в цілому, але визнання беззупинності руху до його впровадження також обумовлене розумінням таких переваг і рахує його розвиток і впровадження. Низкою європейських документів [1-8] в період від 2000 р. до 2024 р. закладалися передумови побудови єдиного спільного нормативного акту, спрямованого на прийняття спільного для ЄС акту, спрямованого на контрольований розвиток та безпечне впровадження ІІІ, яким став прийнятий 13 березня 2024 р. Європейським парламентом Закон про Штучний інтелект (Artificial Intelligence Act [8]) (далі – АІ Аст).

Багаторічна робота надо документі групи експертів дала можливість визначити у ньому основні документи, коротко мета викладена у п. 1.4. Додатку «Законодавчий фінансовий звіт» до АІ Аст як «забезпечення належного функціонування єдиного ринку шляхом створення умов для розробки та використання надійного штучного інтелекту в Союзі»[8].

Так, врахувавши всі попередні напрацювання та перестороги, орієнтуючись на п. С European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics (2015/2103(INL)), згідно якого «необхідно створити загальновизнане визначення роботи та АІ, котре було б гнучким та не перешкоджало б інноваціям» [4], АІ Аст встановив такі обмеження використання ІІІ, які спрямовані на убезпечення від загроз шкідливого використання ІІІ, з розумінням того що «залежно від обставин щодо його конкретного застосування, використання та рівня технологічного розвитку, ІІІ може створювати ризики та завдавати шкоди суспільним інтересам і основним

правам, які захищені законодавством Союзу.» (п. 5 Преамбули Регламенту (ЄС) 2024/1689 [9] з відображенням цієї редакції в AI Act).

В статті 5 документу визначені заборонені практики використання ШІ, через оцінку їх як таких що містять найвищий рівень ризику та / або неприпустимі практики.

Проаналізуємо їх глибше з метою встановлення систематизації критеріїв заборон використання AI.

Статтею 5 AI Act встановлюються вимоги до провайдерів і користувачів ШІ систем, що визначають заборону їх застосування, згідно котрих забороняється використання таких практик ШІ, що мають *широкий спектр шкідливих або таких що породжують зловживання, наслідків*.

Згідно п. 1 ст. 5 AI Act заборони щодо використання систем ШІ включають 8 позицій, які можна поділити на три групи за обсягом заборон.

Так, до першої групи, яка стосується заборон розміщення на ринку, введення в експлуатацію або використання системи ШІ **які використовують впливи на поведінку та прийняття рішень**:

1) підсвідомі методи за межами свідомості людини або навмисно маніпулятивні чи оманливі методи з *метою або результатом суттєвого спотворення поведінки особи чи групи людей, істотно погіршуючи їх здатність приймати обґрунтоване рішення, тим самим змушуючи їх прийняти рішення, яке вони б інакше не прийняли*, що спричиняє або обґрунтовано ймовірно спричинить цій особі, іншій особі чи групі осіб істотну шкоду (пп. «а») (курсив тут і далі мій, В.В.);

2) будь-яку вразливість фізичної особи або певної групи осіб через їх вік, інвалідність або конкретну соціальну чи економічну ситуацію, з *метою або наслідки суттєвого спотворення поведінки цієї особи або особи, яка належить до такої групи*, таким чином, що спричиняє або обґрунтовано може спричинити цій особі чи іншій особі значної шкоди (пп. «b»).

То ж, цю групу можна визначати як заборонені технології ШІ, які здійснюють вплив на прийняття рішень та поведінку забороненими методами.

Друга група заборон стосуються з введення в ринок та використання ШІ, **які призводять до дискримінації, стигматизації або упередженого ставлення**;

1) для оцінки або класифікації фізичних осіб або груп осіб протягом певного періоду часу на основі їх соціальної поведінки або відомих,

припущених або передбачених особистих чи особистісних характеристик, з соціальним показником, що призводить до одного або обох з таких показників:

- *шкідливе або несприятливе ставлення до певних фізичних осіб або груп осіб у соціальних контекстах, які не пов'язані з контекстами, у яких дані були первинно створені або зібрані;*

- *шкідливе чи несприятливе поводження з окремими фізичними особами чи групами осіб, яке є невинуватим або непропорційним їхній соціальній поведінці чи її тяжкості (пп. «с»);*

2) *для оцінки ризиків фізичних осіб з метою оцінки або прогнозування ризику вчинення фізичною особою кримінального правопорушення виключно на основі профілювання фізичної особи або щодо оцінки її особистісних якостей та характеристик<sup>18</sup> (пп. «d»);*

3) *біометричних систем категоризації, які класифікують окремих фізичних осіб на основі їхніх біометричних даних, щоб зробити висновок про їхню расу, політичні погляди, членство в профспілках, релігійні чи філософські переконання. статеве життя або сексуальна орієнтація<sup>19</sup> (пп. «g»).*

Третя група заборон стосується **захисту приватності** і стосується заборон на використання III:

1) *які створюють або розширюють бази даних розпізнавання облич за допомогою нецільового збирання зображень облич з Інтернету чи записів камер відеоспостереження (пп. «e»);*

2) *для визначення емоцій фізичної особи на робочому місці та в навчальних закладах<sup>20</sup> (пп. «f»);*

3) *використання систем дистанційної біометричної ідентифікації «в режимі реального часу» в загальнодоступних місцях для цілей*

---

<sup>18</sup> Виключенням для цієї заборони є системи штучного інтелекту, які використовуються для підтримки людської оцінки причетності особи до злочинної діяльності, яка вже базується на об'єктивних фактах, які можна перевірити, безпосередньо пов'язаних зі злочинною діяльністю (пп. d ч. 1 ст. 5 AI Act).

<sup>19</sup> Заборона не поширюється на будь-яке маркування або фільтрацію законно отриманих наборів біометричних даних, таких як зображення, на основі біометричних даних або категоризацію біометричних даних у сфері правоохоронних органів (пп. «g» ч. 1 ст. 5 AI Act)

<sup>20</sup> Винятком для цієї заборони є використання та розміщення на ринку описаних систем III медичного або безпекового призначення (пп. пп. «f» ч. 1 ст. 5 AI Act)/

правоохоронних органів, якщо таке використання не є суворо необхідним<sup>21</sup> (пп. «h»).

Виходячи з викладеного, можна зробити висновок що основними критеріями заборони використання ІІІ в ЄС сьогодні є недопустимість маніпулювання свідомістю та вразливостями людей і груп (1), використання систем які породжують чи поглиблюють дискримінацію та стигматизацію, порушують приватність (3). За виключенням цільового використання відносно позиції 2 в безпекових та позиції 3 в правоохоронних цілях виключного порядку (Додаток II AI Act), визначені практики ІІІ заборонені.

Очевидно, що перед національними правовими системами постає питання передбачення юридичної відповідальності за порушення таких заборон, адже заборони, описані вище, можуть завдати серйозної шкоди, у разі вакууму легітимного упередження їхньому вчиненню.

#### **Список використаних джерел:**

1. Декларація тисячоліття Організації Об'єднаних Націй прийнята резолюцією 55/2 Генеральної Асамблеї ООН 8 вересня 2000 року. URL:[https://www.un.org/ru/documents/decl\\_conv/declarations/summitdecl.shtml](https://www.un.org/ru/documents/decl_conv/declarations/summitdecl.shtml)
2. Женевська Декларація принципів. Побудова інформаційного суспільства: глобальна задача в новому тисячолітті. URL: <https://old.apitu.org.ua/wsis/dp>
3. Окінавська хартія глобального інформаційного суспільства - <https://regulation.gov.ua/documents/id149711>
4. European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics (2015/2103(INL)). European Parliament Official web-site. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52017IP0051>
5. European Council, European Council meeting (19 October 2017) – Conclusion EUCO 14/17, 2017, p. 8. - <https://www.consilium.europa.eu/media/21620/19-euco-final-conclusions-en.pdf>

---

<sup>21</sup> Підстави режиму суворої необхідності вказані у пп. i-iii пп. «h» ч. 1 ст. 5 AI Act стосовно цілеспрямованого розшуку конкретних жертв, запобігання конкретній загрозі та локалізації та ідентифікації злочинців, які підозрюються у вчиненні діянь перелічених у Додатку II: тероризм, наркотрафік, сексуальні злочини та дитяча порнографія, торгівля людьми і т.д.

6. Council of the European Union, Artificial intelligence b) Conclusions on the coordinated plan on artificial intelligence-Adoption 6177/19, 2019. - <https://data.consilium.europa.eu/doc/document/ST-6177-2019-INIT/en/pdf>

7. Рекомендації з етики для надійного AI (Ethics guidelines for trustworthy AI) Ради Європи 2019 р. - <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

8. European Parliament legislative resolution of 13 March 2024 on the proposal for a regulation of the European Parliament and of the Council on laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union Legislative Acts (COM(2021)0206 – C9-0146/2021 – 2021/0106(COD)) URL: [https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138\\_EN.html?utm\\_source=chatgpt.com#title2](https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_EN.html?utm_source=chatgpt.com#title2)

9. 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) URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

**Дегтярьов І.М.**

кандидат технічних наук, доцент,  
доцент кафедри технології машинобудування,  
верстатів та інструментів СумДУ  
ORCID: <https://orcid.org/0000-0001-8535-987X>

**Петровський М.В.**

кандидат фізико-математичних наук, доцент,  
доцент кафедри електроенергетики СумДУ  
ORCID: <https://orcid.org/0000-0002-0387-3136>

**Леонтєв П.В.**

кандидат технічних наук, завідувач кафедри  
комп'ютеризованих систем управління СумДУ  
ORCID: <https://orcid.org/0000-0002-9494-9078>

## **МЕТОДОЛОГІЯ ТЕСТУВАННЯ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ДЛЯ СИСТЕМИ АВТОНОМНОЇ НАВІГАЦІЇ НАЗЕМНИХ РОБОТИЗОВАНИХ КОМПЛЕКСІВ**

Методологія тестування програмного забезпечення для систем автономної навігації наземних роботизованих комплексів ґрунтується на комплексному підході, що включає як лабораторні, так і польові випробування. Передусім необхідно визначити критерії точності навігації, серед яких, здатність системи до коректного руху за заданими координатами та контроль досягнення цільової точки. Для цього задаються різні маршрути з фіксованими точками призначення, а програмне забезпечення перевіряється на правильність побудови траєкторії та коректність відстеження власних координат. Ключовим моментом є багаторазове повторення дослідів: після кожного проходження маршруту вимірюється відстань між реальною фінальною позицією робота та заданою точкою, а розбіжності фіксуються для подальшого статистичного аналізу точності та стабільності роботи алгоритмів навігації.

Важливим компонентом сучасних систем автономної навігації є використання методів штучного інтелекту (ШІ), зокрема алгоритмів машинного навчання та глибинного навчання. Такі алгоритми дозволяють роботизованим комплексам самостійно вдосконалювати моделі навігації,

оптимізувати траєкторії руху та підвищувати точність розпізнавання навколишнього середовища. Під час тестування системи доцільно перевіряти не лише статичну правильність роботи алгоритмів, але й їхню здатність до адаптації, тобто, наскільки ефективно модель ШІ навчається на основі накопичених даних із попередніх випробувань.

Крім того, методологія має враховувати особливості тренування моделей у симуляційному середовищі, де можливо безпечно відтворити широкий спектр сценаріїв. Використання симуляторів дозволяє формувати великий набір навчальних даних, у тому числі рідкісні або екстремальні ситуації, які складно відтворити в реальних умовах. Це сприяє підвищенню надійності навігаційного ШІ у складних обставинах. Важливим аспектом є валідація навчальних наборів даних – потрібно перевіряти їх збалансованість, відсутність упередженості та відповідність реальним сценаріям експлуатації роботизованих систем.

Окремо тестуються модулі прийняття рішень на основі ШІ: логіка вибору траєкторії, визначення пріоритетів між безпекою, швидкістю та енергоефективністю. Таке тестування дає змогу оцінити, наскільки система здатна до раціонального компромісу в складних середовищах, де одночасно діють кілька суперечливих факторів.

Додатково тестування має включати оцінку роботи системи у змінених умовах, наближених до реальних. Це передбачає створення статичних перешкод (наприклад, стіни, колони, бар'єри), що ускладнюють рух по маршруту, та аналіз того, наскільки ефективно програмне забезпечення виявляє ці перешкоди та перебудовує траєкторію. Окремим важливим етапом є моделювання динамічних перешкод: наприклад, поява людини або іншого об'єкта, що рухається назустріч роботу. У таких випадках необхідно фіксувати час реакції системи, точність зміни курсу, наявність затримок у прийнятті рішень і здатність уникнути зіткнення. Особливу увагу варто приділяти аналізу латентності між моментом виявлення перешкоди та виконанням коригувальних дій, адже цей показник критично впливає на безпеку роботи комплексу.

ШІ також відіграє ключову роль у розпізнаванні об'єктів і прийнятті рішень у реальному часі. Використання нейронних мереж для аналізу даних із камер та сенсорів дозволяє системі класифікувати типи перешкод, прогнозувати їхню траєкторію руху та вибирати оптимальний шлях обходу.

Під час тестування доцільно оцінювати точність таких моделей розпізнавання, швидкість реакції ШІ-модуля та стабільність його роботи при зміні освітлення, погодних умов чи шумів сенсорів.

Крім цього, актуальним напрямом є тестування мультимодальних систем, у яких ШІ об'єднує дані з різних сенсорів – лідарів, стереокамер, GPS, інерціальних вимірювальних блоків. Це дозволяє досягти вищої надійності навігації та кращої ситуаційної обізнаності. Для таких систем проводяться стрес-тести, під час яких штучно створюються збої або шум у даних, щоб перевірити, наскільки ШІ здатен підтримувати стабільну роботу в умовах невизначеності.

Додаткову увагу приділяють інтерпретованості рішень моделей штучного інтелекту. Важливо не лише те, що система приймає правильне рішення, а й те, щоб це рішення можна було пояснити, що є критичним для безпеки та сертифікації автономних систем.

Таким чином, методологія тестування програмного забезпечення для автономної навігації наземних роботизованих комплексів передбачає багаторівневу перевірку: від базового контролю точності руху за координатами до оцінки поведінки в умовах складної динамічної обстановки. Впровадження ШІ створює нові виклики для тестування – необхідність оцінки не лише функціональної коректності, а й якості навчання, узагальнення та здатності системи до самонавчання. Тому сучасна методологія повинна поєднувати класичні інженерні підходи з перевіркою надійності моделей ШІ.

Перспективним напрямом є розроблення автономних систем тестування, де сам ШІ аналізує результати випробувань, виявляє закономірності помилок і пропонує шляхи вдосконалення алгоритмів навігації. Такий підхід дозволяє скоротити час перевірки, підвищити точність оцінки та створити замкнутий цикл безперервного вдосконалення навігаційного програмного забезпечення.

Тільки поєднання статистичного аналізу точності, перевірки стійкості алгоритмів та вимірювання швидкодії реакцій дозволяє об'єктивно оцінити готовність програмного забезпечення до реального застосування.

Робота виконана за підтримки проєкту «Система автономної навігації для наземного роботизованого комплексу подвійного призначення» реєстраційний номер № РН/67-2024.

**Дегтярьов І.М.**

кандидат технічних наук, доцент,  
доцент кафедри технології машинобудування,  
верстатів та інструментів СумДУ

ORCID: <https://orcid.org/0000-0001-8535-987X>

**Петровський М.В.**

кандидат фізико-математичних наук, доцент,  
доцент кафедри електроенергетики СумДУ

ORCID: <https://orcid.org/0000-0002-0387-3136>

**Леонтьєв П.В.**

кандидат технічних наук, завідувач кафедри  
комп'ютеризованих систем управління СумДУ

ORCID: <https://orcid.org/0000-0002-9494-9078>

## **ВАЖЛИВІСТЬ РОЗРОБЛЕННЯ МЕТОДИКИ ВИПРОБУВАНЬ ДЛЯ КОНТРОЛЮ ТЕХНІЧНИХ ХАРАКТЕРИСТИК НАЗЕМНИХ РОБОТИЗОВАНИХ КОМПЛЕКСІВ З ВИКОРИСТАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ, ЯК АЛЬТЕРНАТИВА ПЕРЕВІРЦІ В РЕАЛЬНИХ УМОВАХ ЕКСПЛУАТАЦІЇ**

Розвиток сучасних робототехнічних систем, що використовують алгоритми штучного інтелекту (ШІ), зумовлює необхідність створення нових підходів до їх тестування. Наземні роботизовані комплекси (НРК) виконують складні завдання у військовій, рятувальній, промисловій та логістичній сферах, де точність і надійність роботи систем мають вирішальне значення.

Традиційні методи перевірки в реальних умовах експлуатації пов'язані з високими витратами, обмеженим контролем зовнішніх факторів і потенційними ризиками для обладнання. Тому актуальним є розроблення методики, що забезпечує достовірне оцінювання характеристик систем без проведення небезпечних польових випробувань.

Метою дослідження було розроблення концептуальної основи методики випробувань наземних роботизованих комплексів з використанням ШІ, заснованої на симуляційних моделях та цифрових двійниках.

Перевірка НРК у реальних умовах ускладнена через значні фінансові витрати, ризик пошкодження техніки, обмеженість повторюваності

експериментів та залежність результатів від зовнішніх чинників. Крім того, поведінка алгоритмів ШІ може змінюватися під впливом непередбачуваних обставин, що ускладнює процес об'єктивного оцінювання.

Моделювання дозволяє створити віртуальний полігон, у якому можуть бути відтворені будь-які сценарії роботи НРК, а саме від пересування по складному рельєфу до взаємодії з об'єктами. Це дає можливість:

- безпечно тестувати алгоритми навігації та розпізнавання;
- проводити багаторазові тести за однакових умов;
- зменшити вплив людського фактору;
- прискорити процес оптимізації параметрів системи.

Технологія цифрового двійника забезпечує постійну відповідність між реальним об'єктом і його віртуальною копією. Це дозволяє в режимі реального часу аналізувати стан системи, прогнозувати поведінку, виявляти відхилення й підвищувати ефективність навчання моделей ШІ. Такий підхід є важливим елементом автоматизованого контролю технічного стану та сертифікації роботизованих систем.

Розроблення уніфікованої методики випробувань із використанням цифрових двійників дає змогу забезпечити стандартизацію підходів до оцінки ефективності НРК. Результати моделювання можуть бути використані для навчання ШІ, оптимізації енерговитрат, покращення алгоритмів ухвалення рішень та прогнозування технічних збоїв.

Запропонований підхід є універсальним і може бути адаптований для різних типів роботизованих систем, зокрема автономних транспортних платформ, розвідувальних або рятувальних роботів.

Тестування нових НРК є ключовим етапом життєвого циклу технічної системи. Воно забезпечує:

- перевірку відповідності заявлених технічних характеристик фактичним можливостям;
- виявлення прихованих дефектів конструкції, алгоритмів або програмного забезпечення;
- оцінку надійності та стабільності роботи системи в умовах впливу зовнішніх факторів;
- гарантування безпеки експлуатації, особливо у випадках використання комплексів у військових або аварійних ситуаціях;

- зменшення витрат на обслуговування та ремонти у майбутньому, завдяки ранньому виявленню проблем.

Тестування дозволяє не лише перевірити працездатність системи, але й оптимізувати алгоритми ШІ, які приймають рішення у складному, непередбачуваному середовищі. Саме завдяки багатоетапним випробуванням можливо досягти високого рівня автономності, стійкості та адаптивності роботизованих комплексів.

Ефективне тестування може поєднувати реальні та віртуальні етапи перевірки:

- лабораторні випробування — проводяться у контрольованих умовах для первинної перевірки основних функцій, точності сенсорів, роботи приводів, систем управління тощо;
- симуляційні (віртуальні) випробування — виконуються у цифровому середовищі, що дозволяє моделювати небезпечні або екстремальні ситуації без ризику для обладнання;
- польові випробування — проводяться після попередніх етапів і спрямовані на перевірку поведінки системи у реальному середовищі.

Поєднання цих підходів дає змогу отримати повну картину працездатності системи, а також забезпечити порівнянність результатів завдяки стандартизованим сценаріям тестування.

Наземні роботизовані комплекси під час роботи піддаються дії широкого спектра зовнішніх факторів. Для адекватної оцінки їх надійності та живучості необхідно враховувати наступні групи умов.

Кліматичні та природні фактори:

- Дощ, сніг, град, туман — знижують якість зображення з камер, ускладнюють навігацію, викликають помилки у системах комп'ютерного зору.
- Мороз, низькі температури — впливають на ефективність акумуляторів, в'язкість мастил, електропровідність контактів.
- Спека, пил, пісок, висока вологість — призводять до перегріву, забруднення сенсорів, зниження точності вимірювань.
- Ожеледь, слизькі поверхні, болото — створюють ризик втрати зчеплення, ковзання та нестійкості шасі.

Механічні впливи:

- Удари та вібрації при транспортуванні або виграції, падіння з невеликої висоти — можуть спричинити зміщення оптичних сенсорів, мікротріщини корпусу або пошкодження комунікаційних кабелів.

- Динамічні навантаження при русі по нерівностях — впливають на точність навігації, роботу підвіски та приводу.

Бойові та екстремальні фактори:

- Підриви на міні, вибухові хвилі, постріли, FPV-атаки — перевіряють рівень бронювання, стійкість корпусу та здатність системи відновлювати роботу після пошкоджень.

- Електромагнітні перешкоди або спроби радіоелектронного придушення — тестують надійність зв'язку, систем навігації та автономності алгоритмів.

Технічні та експлуатаційні фактори:

- Слабкий заряд акумулятора — визначає здатність системи переходити в енергозберігаючі режими або безпечно завершувати місію.

- Брудні або пошкоджені камери огляду, запилені сенсори — вимагають перевірки алгоритмів самодіагностики та компенсації втрати зображення.

- Втрата зв'язку або часткова відмова системи управління — дає можливість перевірити коректність автономного керування та безпечного режиму роботи.

Для забезпечення високої надійності НРК необхідно проводити багаторівневе тестування, що враховує всі можливі сценарії експлуатації.

Використання симуляційних технологій, цифрових двійників і лабораторних випробувань дозволяє оптимізувати процес перевірки, зменшити ризики пошкодження реального обладнання та прискорити цикл розроблення.

Урахування кліматичних, механічних, бойових і технічних факторів під час тестування є запорукою створення стійких і ефективних роботизованих систем, готових до роботи у непередбачуваних умовах реального середовища. Для зручності сприйняття та структуризації інформації щодо умов експлуатації розроблено таблицю 1.

Таблиця 1 – Структуризація умов та параметрів експлуатації НРК

| <b>Група умов експлуатації</b>            | <b>Типові ситуації / фактори впливу</b>                    | <b>Параметри, що підлягають контролю</b>                                                                                                                                                                                                                                                                                       |
|-------------------------------------------|------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Кліматичні та природні фактори</b>     | Дощ, сніг, туман, мороз, спека, пил, ожеледь, болото       | <ul style="list-style-type: none"> <li>- Стійкість сенсорів (LiDAR, камера, GPS)</li> <li>- Робота систем зору при зниженій видимості</li> <li>- Температурна стабільність акумуляторів і приводів</li> <li>- Вологозахист і пилонепроникність корпусу (IP-рівень)</li> <li>- Зчеплення коліс/гусениць із поверхнею</li> </ul> |
| <b>Механічні впливи</b>                   | Вібрації, удари, виграшка, падіння, рух по нерівностях     | <ul style="list-style-type: none"> <li>- Механічна міцність корпусу</li> <li>- Стійкість до вібрацій (частота, амплітуда)</li> <li>- Надійність кріплення сенсорів і вузлів</li> <li>- Калібрування систем після удару</li> </ul>                                                                                              |
| <b>Бойові екстремальні фактори та</b>     | Підриви на міні, вибухові хвилі, постріли, FPV-атаки       | <ul style="list-style-type: none"> <li>- Ступінь пошкодження корпусу й бронювання</li> <li>- Стійкість електроніки до ударних навантажень</li> <li>- Збереження систем зв'язку</li> <li>- Автоматичне відновлення після критичних відмов</li> </ul>                                                                            |
| <b>Технічні експлуатаційні фактори та</b> | Низький заряд акумулятора, відмова модулів, втрата зв'язку | <ul style="list-style-type: none"> <li>- Енергоспоживання та режими енергозбереження</li> <li>- Поведінка системи при втраті каналу зв'язку</li> <li>- Надійність автономного керування</li> <li>- Система самодіагностики та резервування</li> </ul>                                                                          |
| <b>Засмічення та деградація сенсорів</b>  | Забруднені, замулені або пошкоджені камери, пил на лінзах  | <ul style="list-style-type: none"> <li>- Якість зображення та точність розпізнавання об'єктів</li> <li>- Алгоритми компенсації втрати даних</li> <li>- Стійкість систем ШІ до неповної інформації</li> </ul>                                                                                                                   |

|                                                |                                                     |                                                                                                                                                                                                |
|------------------------------------------------|-----------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Електромагнітні та інформаційні фактори</b> | Радіоперешкоди, спроби глушіння сигналу, кібератаки | <ul style="list-style-type: none"> <li>- Стійкість систем зв'язку та навігації</li> <li>- Безперервність каналу передачі даних</li> <li>- Захист програмного забезпечення від збоїв</li> </ul> |
|------------------------------------------------|-----------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

Робота виконана за підтримки проєкту «Система автономної навігації для наземного роботизованого комплексу подвійного призначення» реєстраційний номер № РН/67-2024.

## **БЕЗПЕКА, ОБОРОНА, КІБЕРЗАХИСТ ТА ВОЄННИЙ КОНТЕКСТ ПРИ РЕГУЛЮВАННІ ШТУЧНОГО ІНТЕЛЕКТУ**

**Дар'я Глушкова**

доктор філософії, доцент кафедри  
правоохоронної діяльності навчально-  
наукового інституту №5

Харківський національний університет  
внутрішніх справ

ORCID: <https://orcid.org/0000-0001-5909-6367>

### **ШТУЧНИЙ ІНТЕЛЕКТ У СИСТЕМІ НАЦІОНАЛЬНОЇ БЕЗПЕКИ УКРАЇНИ. ПРАВОВІ ЗАСАДИ РЕГУЛЮВАННЯ В УМОВАХ ВОЄННОГО СТАНУ**

Сучасні технології штучного інтелекту дедалі активніше інтегруються в систему національної безпеки, створюючи нові можливості для захисту держави, але водночас породжуючи складні правові, етичні й безпекові виклики. В умовах воєнного стану, який в Україні діє у зв'язку зі збройною агресією Російської Федерації, питання правового регулювання використання штучного інтелекту набуває особливого значення. Йдеться не лише про розробку технічних стандартів, а про формування цілісної нормативної бази, здатної забезпечити баланс між інтересами обороноздатності, національної безпеки та дотриманням прав людини.

Штучний інтелект (ШІ) – це комплекс технологій, здатних аналізувати дані, робити висновки та приймати рішення без безпосереднього втручання людини. Його використання у сфері безпеки охоплює широкий спектр напрямів: від кіберзахисту критичної інфраструктури, моніторингу соціальних мереж і виявлення інформаційних атак до застосування автономних систем озброєння, безпілотних літальних апаратів та алгоритмів прогнозування загроз. Такі технології є важливими інструментами державної стратегії виживання, проте їх функціонування має відбуватися виключно у межах правового поля [1].

Відповідно до статті 17 Конституції України, забезпечення державної безпеки є найважливішою функцією держави, а захист суверенітету й територіальної цілісності – справою всього українського народу [2]. У цьому контексті штучний інтелект може розглядатися як елемент системи безпеки, що підсилює спроможність держави діяти ефективно під час війни. Разом із тим відсутність належного правового регулювання створює ризики порушення прав і свобод людини, втрати контролю над автономними системами або несанкціонованого доступу до чутливої інформації.

Правові виклики використання ШІ в умовах воєнного стану охоплюють кілька площин. По-перше, питання відповідальності за наслідки рішень, ухвалених алгоритмами, залишається неврегульованим. Хто несе юридичну відповідальність за шкоду, заподіяну автономною системою – оператор, розробник, держава чи сама система як «техноагент»? По-друге, виникає проблема забезпечення прозорості алгоритмів, що використовуються у процесах військового управління або кібероборони. Закритість таких систем може призвести до порушення принципів верховенства права і парламентського контролю. По-третє, існує небезпека використання технологій ШІ для масового стеження, цензури чи маніпуляцій суспільною думкою під приводом «забезпечення безпеки».

Водночас ШІ може стати важливим чинником зміцнення інформаційної стійкості держави. Алгоритми здатні виявляти фейкові новини, відстежувати координацію кібератак, ідентифікувати ворожі інформаційні компанії та сприяти захисту державних електронних ресурсів. Для України, яка перебуває під постійним тиском гібридних загроз, це має не лише технологічне, а й стратегічне значення. Проте для ефективного використання таких рішень необхідно визначити правові засади збору, обробки й зберігання даних, з урахуванням вимог Закону України «Про національну безпеку України» та законодавства про захист персональних даних [3].

Європейський Союз зробив важливий крок, ухваливши у 2024 році Artificial Intelligence Act, який встановлює ризик-орієнтовану модель регулювання та передбачає обмеження для «високоризикових систем», зокрема у сфері безпеки, правосуддя й оборони [4]. Цей підхід може стати орієнтиром для України, оскільки дозволяє поєднати безпекові потреби з повагою до прав людини та демократичних принципів. Також значну роль відіграють рекомендації ООН і Ради Європи, які підкреслюють що

використання ІІ у військових цілях має ґрунтуватися на міжнародному гуманітарному праві та принципах пропорційності, гуманності й контролю з боку людини.

У межах українського правопорядку необхідно розробити спеціальний законодавчий акт, який би визначав поняття, категорії ризиків, систему нагляду та відповідальності у сфері застосування ІІ у військових та оборонних цілях. Також доцільно створити міжвідомчий центр етичного та правового моніторингу використання алгоритмічних систем у сфері безпеки. Особливої уваги потребує питання кібербезпеки: саме штучний інтелект має стати не лише об'єктом регулювання, а й суб'єктом захисту – технологічним «щитком» проти кіберагресій.

Таким чином, використання ІІ в системі національної безпеки України має спиратися на конституційні принципи верховенства права, законності, пропорційності та підзвітності. Умови воєнного стану вимагають швидких рішень, але навіть у цих обставинах держава зобов'язана дотримуватися міжнародних стандартів і забезпечувати захист прав людини. Формування нормативної бази у цій сфері має поєднувати технологічні інновації з етичними цінностями, щоб ІІ залишився інструментом безпеки, а не загрозою свободі.

### **Список використаних джерел:**

1. Штучний інтелект у вищій освіті: ризики та перспективи інтеграції: матеріали всеукраїнського науково-педагогічного підвищення кваліфікації, 1 липня – 11 серпня 2024 року. – Львів – Торунь : Liha-Pres, 2024. 328 с. [Електронний ресурс]. – Режим доступу: [https://cuesc.org.ua/images/informlist/Maket%20advanced\\_training\\_OLA.pdf](https://cuesc.org.ua/images/informlist/Maket%20advanced_training_OLA.pdf) (дата звернення 19.10.2025)
2. Конституція України: Закон України від 28 червня 1996р. №254к/96-ВР [Електронний ресурс]. – Режим доступу: <https://zakon.rada.gov.ua/laws/show/254k/96-vr#Text> (дата звернення 19.10.2025)
3. Про національну безпеку України Закон України від 21.06.2018 №2469-VIII [Електронний ресурс]. – Режим доступу: <https://zakon.rada.gov.ua/laws/show/2469-19#Text> (дата звернення 19.10.2025)
4. Artificial Intelligence Act [Електронний ресурс]. – Режим доступу: <https://artificialintelligenceact.eu> (дата звернення 20.10.2025)

**Олена Грезіна**

доктор філософії у галузі права,  
доцент кафедри кримінального аналізу та інформаційних технологій  
Одеський державний університет внутрішніх справ  
ORCID: <https://orcid.org/0000-0002-2491-6529>

## **ЗАХИСТ ІНФОРМАЦІЙНИХ РЕСУРСІВ ДЕРЖАВИ В УМОВАХ ВОЄННИХ ЗАГРОЗ ТА РОЛЬ ШТУЧНОГО ІНТЕЛЕКТУ**

Забезпечення інформаційної безпеки держави в умовах воєнних загроз набуває стратегічного значення, оскільки воєнні конфлікти створюють комплекс підвищених ризиків втручання у критично важливі інформаційні системи, порушення цілісності та достовірності даних, а також дестабілізації ключових державних процесів і функцій. У таких умовах кібератаки виступають не лише інструментом отримання конфіденційної інформації, але й засобом реалізації стратегічних та тактичних завдань противника, включаючи паралізацію критично важливих об'єктів інфраструктури, порушення логістичних ланцюгів, дезорганізацію енергетичних та комунікаційних мереж, а також маніпулювання громадською свідомістю.

Основними загрозами інформаційній безпеці держави в умовах воєнних загроз є несанкціоноване втручання в державні інформаційні системи, модифікація або знищення критично важливих даних, маніпуляції, що впливають на процес ухвалення рішень на всіх рівнях управління, порушення функціонування транспортної, енергетичної та комунікаційної інфраструктури, а також психологічний тиск на населення та державні установи через масовані інформаційні кампанії. Виявлення, оцінка та нейтралізація таких загроз вимагають високого рівня організаційної підготовки, узгодженості дій між державними органами та застосування передових технологічних рішень для забезпечення безперервності роботи інформаційних систем.

Розвиток алгоритмічних систем дозволяє підвищити ефективність захисту державних інформаційних ресурсів за рахунок автоматизації процесів моніторингу, аналізу подій та реагування на загрози у режимі реального часу. Алгоритмічні системи здатні здійснювати аналіз великих обсягів даних, виявляти аномалії у функціонуванні мереж, прогнозувати потенційні загрози

та формувати оперативні рішення для нейтралізації інцидентів, що значно скорочує час від виявлення загрози до її усунення. Використання таких рішень забезпечує безперервний контроль інформаційних потоків, раннє попередження про спроби несанкціонованого доступу, а також реалізацію автоматизованих механізмів протидії кібератакам, що дозволяє знизити негативні наслідки атак і забезпечити стійкість критично важливих об'єктів інфраструктури.

Інтеграція алгоритмічних систем у державні структури та об'єкти критичної інфраструктури здійснюється на основі централізованих платформ обробки даних, систем виявлення загроз, аналізу поведінки потенційних нападників і прогнозування сценаріїв кібератак [1]. Така інтеграція передбачає використання алгоритмічних рішень у державних органах, відповідальних за кібербезпеку, а також у системах контролю доступу, мережевого моніторингу та аналітичного управління інформаційними потоками. Комплексний підхід до захисту державних ресурсів дозволяє підвищити надійність інформаційних систем, знизити ризики, пов'язані з людським фактором, забезпечити масштабованість та адаптивність механізмів безпеки при збільшенні обсягів даних і ускладненні кіберзагроз.

Переваги застосування алгоритмічних систем у сфері інформаційної безпеки проявляються у підвищенні ефективності виявлення та оперативного реагування на загрози, скороченні часу ухвалення рішень, формуванні аналітичної бази для стратегічного планування та забезпеченні стійкості державних інформаційних ресурсів у надзвичайних умовах [2]. Використання таких рішень дозволяє підвищити здатність держави протидіяти кібератакам та інформаційним впливам, мінімізувати негативні наслідки порушень безпеки та підтримувати безперервність функціонування критично важливих об'єктів.

Забезпечення інформаційної безпеки держави в умовах воєнних загроз потребує комплексного поєднання технологічних, організаційних та правових механізмів. Алгоритмічні системи виконують ключову роль у підвищенні стійкості державних інформаційних ресурсів, забезпеченні автоматизованого та оперативного реагування на загрози, а також у створенні аналітичної платформи для стратегічного управління кібербезпекою. Подальший розвиток методів аналізу даних, прогнозування та автоматизації процесів безпеки сприятиме зміцненню національної інформаційної безпеки, підвищенню

готовності державних органів до протидії складним кібератакам та забезпеченню ефективного управління критично важливою інфраструктурою у кризових умовах.

### **Список використаних джерел**

1. Adewusi, Adebunmi Okechukwu, Ugochukwu Ikechukwu Okoli, Temidayo Olorunsogo, Ejuma Adaga, Donald Obinna Daraojimba, and Ogugua Chimezie Obi. 2024. «Artificial Intelligence in Cybersecurity: Protecting National Infrastructure: A USA Review». *World Journal of Advanced Research and Reviews* 21, no. 1: 2263–2275. <https://doi.org/10.30574/wjarr.2024.21.1.0313>
2. M., V., and R. Murugan. 2024. «The Role of Artificial Intelligence in Cyber Security». *International Journal of Innovative Research in Computer and Communication Engineering* 12, no. 03: 1635–1641. <https://doi.org/10.15680/ijircce.2024.1203044>

**Вікторія Кошинець**

Аспірант Інституту права Київського  
національного університету імені Тараса  
Шевченка

## **СИМБІОЗ ШТУЧНОГО ІНТЕЛЕКТУ ТА РОЗПОДІЛЕНИХ ТЕХНОЛОГІЙ ДЛЯ ЗАХИСТУ ОБОРОННИХ ДАНИХ**

У сучасному світі одна із головних тенденцій є активне впровадження інформаційних технологій у різні сфери діяльності. Цифровізація стає важливим фактором підвищення продуктивності та покращення якості життя, зокрема безпека даних в оборонній сфері стає питанням національного виживання. В умовах гібридних війн, кібератак і постійних спроб втручання в критичну інфраструктуру держав, традиційні методи захисту інформації вже не гарантують повної надійності. Саме тут на перетині двох революційних технологій – блокчейну та штучного інтелекту (ШІ) - народжується нова концепція цифрового захисту.

Серед найперспективніших напрямів впровадження інформаційних систем в Україні та світі є блокчейн технології. Blockchain – термін утворений поєднанням двох англійських слів: «block» – блок та «chain» - ланцюг. Блок – файл, який в залежності від системи має певні характеристики: частота створення, розмір, дані. Дані у блоках зашифровані з використанням криптографічних алгоритмів. Ланцюг – по'єднання блоків у логічну послідовність, кожен новостворений блок містить посилання на попередній блок. Такий спосіб зберігання інформації майже унеможливорює підробку даних. Зміна даних в будь-якому блоці призведе до зміни його посилання, таким чином буде запущена «ланцюгова реакція» і зміна буде відхилена [1].

Блокчейн, який здобув популярність через криптовалюти, є технологією, що має потенціал революціонізувати не тільки фінансовий світ, але й багато інших секторів, включаючи оборонну промисловість. У військових інформаційних системах блокчейн може пропонувати безпеку, прозорість та надійність, що є критично важливими для військових операцій [2].

Тому основою цієї технології є прозорість. Кожен має можливість подивитись інформацію про будь-який блок. Блокчейн використовує

криптографічні алгоритми для забезпечення роботи системи. Вони гарантують незмінність блоку транзакцій та керують аутентифікацією.

Ця система зберігає та передає цифрову книгу даних за допомогою криптографії для забезпечення конфіденційності та цілісності. В результаті блокчейнові мережі не тільки зменшують ймовірність загрози з боку атаки противника, але суттєво збільшують витрати супротивника на її створення. Важливість цієї технології полягає у створенні довіри до цифрових даних та може призвести до нових винаходів, придатних для застосування у оборонній галузі.

Перевагами використання блокчейну в військових системах є [3]:

*Безпека даних:* Однією з ключових характеристик блокчейну є його здатність забезпечити високий рівень безпеки. Це особливо важливо для військових систем, де захист інформації має життєво важливе значення.

*Незмінність записів:* Записи в блокчейні не можуть бути змінені або видалені без відома всієї мережі, що забезпечує відстежуваність та прозорість даних.

*Децентралізація:* Децентралізована природа блокчейну робить його менш вразливим до атак або збоїв, що можуть бути катастрофічними в централізованих системах.

*Комунікація та координація:* Блокчейн може вдосконалити комунікацію та координацію між різними військовими підрозділами, забезпечуючи надійний обмін даними.

Наразі існує безліч напрямів, де може бути запроваджена блокчейн технологія в оборонній сфері, особливо такі як командування, управління конфіденційними даними, військова логістика та безпека.

По-перше, командна інформаційна система має централізовані мережі та бази даних. Будучи важливою ціллю у воєнний і навіть мирний час, він вразливий перед ворогами чи хакерами. Вони можуть використовувати різні інформаційні засоби для проведення мережових і електричних атак, спричиняючи паралізацію всієї інформаційної системи, або для викрадення та фальсифікації ідентифікаційної інформації, подробиць важливих даних. Тому учасники бойових дій стикаються з великими ризиками щодо достовірності даних і навіть можуть приймати помилкові рішення зі шкідливими даними. По-друге, коли командири покладаються на командну інформаційну систему

для прийняття рішень, існуватиме ризик віддавати неправильні команди чи фальшиві накази, спричинені піддробкою даних ворогом [4].

Тому вчені пропонують впровадити інформаційну систему команд, засновану на технології блокчейн, в яку входять різні рівні, кожний з яких надає певний гарант. Рівень даних гарантує надійність даних, військової інформаційної системи; мережевий рівень - самоорганізацію та децентралізацію функцій військових інформаційних систем; рівень заохочення використовує програмований механізм, щоб уникнути всіх видів неправильної поведінки та досягти стимулів позитивної; договірний рівень допомагає автоматизувати інтелектуальні військові інформаційні системи, зменшуючи невизначеність, різноманітність і складність, яку люди та інші фактори привносять у бойове командування та військове управління.

Використовуючи такий механізм можна досягти достовірного бойового командування, що означає, що команди, віддані командирами на всіх рівнях, можуть бути повністю записані, а переписані дані є доказами та простежуються. Лише коли ворог або хакери модифікують більше 51% інформації вузла одночасно, дані командної інформаційної системи можуть бути змінені.

Сучасні тенденції розвитку цифрових технологій свідчать, що найбільший ефект у сфері безпеки дає поєднання блокчейну та штучного інтелекту (ШІ). Ці технології не конкурують між собою, а взаємно підсилюють можливості одна одної.

Штучний інтелект забезпечує аналіз великих обсягів даних, прогнозування та автоматичне ухвалення рішень, тоді як блокчейн гарантує достовірність, прозорість і захищеність цих даних. Крім того, застосування технологій машинного навчання у блокчейнових мережах допомагає оптимізувати енергоспоживання, маршрути передачі даних та швидкість транзакцій. Це особливо важливо для оборонних систем, що працюють у реальному часі та функціонують у середовищі з обмеженими ресурсами.

Таким чином, взаємодія блокчейну й штучного інтелекту створює основу інтелектуальних оборонних екосистем, здатних не лише зберігати інформацію, а й адаптивно реагувати на загрози, аналізувати ризики та забезпечувати стабільність управління в умовах невизначеності. Окрім того, ми можемо відстежувати в режимі реального часу, чи була піддроблена база

даних, чи відстежувалась військова система, і виключити ризик атак противника, притаманний командним інформаційним системам.

Агентство передових оборонних дослідницьких проєктів США (DARPA) вже працює над таким проєктом та намагається розробити безпечну та надійну інформаційну платформу на основі технології блокчейн, що буде здатна ефективно захищати конфіденційні дані та запобігати злому.

Не менш важливою сферою є **військова логістика**, де важко вирішити проблеми пов'язані з мережею, зберіганням даних, обслуговуванням системи, відстежування та контролем якості під час пакування, завантаження та розвантаження, транспортування та розбирання.

Використовуючи технологію блокчейн, створюється безпечний, спільний і постійний запис, який можна перевірити, відстежувати та перевіряти транзакції різних ланцюгів постачання та між усіма операційними партнерами, а також здійснювати ефективне управління життєвим циклом ланцюга поставок оборонного призначення та системи закупівлі та логістика.

Так, у квітні 2016 року Міністерство оборони США та його союзники по НАТО почали звертати увагу на потенційне застосування технології блокчейн в обороні, включаючи автоматичне виконання смарт-контрактів, безпечне зберігання конфіденційних файлів і зменшення помилок під час виконання оборонного контракту. У червні 2017 року ВМС США скористалися технологією блокчейн для підвищення безпеки систем адитивного виробництва. Вони записували весь процес проектування компонентів, виготовлення прототипу, тестування, виробництва та остаточної обробки, щоб користувачі могли переглядати будь-які конкретні дані та повідомляти про пошкодження компонентів або в кінці їх життєвого циклу [4].

І нарешті, одна з найважливіших сфер є військова безпека, де має бути забезпечена конфіденційність передачі даних на полі бою. Ґрунтуючись на розподілених характеристиках технології блокчейн, можна побудувати систему зв'язку із широким охопленням, стійку до аварій та з високим рівнем безпеки, щоб досягти безперебійного зв'язку в складному середовищі на полі бою.

У травні 2017 року компанія Technology and Manufacturing Company зі штату Індіана, фінансована DARPA, використала технологію блокчейн для розробки «недоступної платформи обміну повідомленнями та торгівлі» для військових. Ця платформа розділяє створення та передачу інформації, щоб

гарантувати, що надіслані та отримані дані неможливо зламати, і забезпечити безпечний зв'язок між штабом і наземними силами, а також між Міністерством оборони та офіційними особами розвідки [4].

Отже, блокчейн технологія є новим підходом в таких напрямках як інформаційна безпека, автентифікація, цілісність та стійкість даних та інших. Маючи величезний потенціал, блокчейн є перспективним в оборонній галузі, оскільки важливість шифрування, кібербезпеки та логістики процесів є важливою при забезпеченні конфіденційності інформації, комунікації та управлінні на полі бою, що так само є актуальним для Збройних Сил України та інших силових структур держави.

Таким чином, технології блокчейн та штучного інтелекту є одним із найперспективніших інструментів у сфері інформаційної безпеки, автентифікації, захисту даних та логістики, зокрема її поєднання зі штучним інтелектом створює принципово новий рівень автоматизації та безпеки військових систем. Впровадження таких технологій в оборонну галузь відкриває можливості для створення стійких, прозорих та децентралізованих систем, здатних гарантувати цілісність і надійність військової інформації, а також це відіграє основну роль для Збройних Сил України, оскільки ця технологічна синергія може стати ключовим чинником розвитку систем кіберзахисту, захищеного обміну повідомленнями та управління логістикою. Тому в найближчому майбутньому науково спільнота оборонної галузі буде шукати нові досліджувати нові шляхи впровадження програм для військових, засновані на технологіях блокчейн, з акцентом на таких сферах як кіберзахист, захищений обмін повідомленнями, комунікації, підтримка логістики тощо.

### **Список використаних джерел:**

1. Невмержицький Є.І., Мальков Є.В. Використання технологій блокчейн. Перспективи її розвитку: курсова робота. Київ, 2021. С.7-8
2. 10 технотрендів для армії майбутнього: чого не вистачає ЗСУ. URL: <https://softline.org.ua/news/10-tekhnotrendiv-dlia-armii-maibutnoho-choho-ne-vystachaie-zsu.html> (дата звернення: 03.03.2024)
3. Блокчейн та військові інформаційні системи: у «Айкюжн ІТ» розповіли про перспективи використання. URL: <https://softline.org.ua/news/blokcejn-ta-vijskovi-informacijni-sistemi-u-ajkuzn-it-rozpovili-pro-perspektivi-vikoristanna.html>
4. Опірський І., Васишин С. Перспективи військового застосування технології блокчейн. Український науковий журнал інформаційної безпеки. 2022. №28. С. 57-66.

**Марина Григор'єва**

кандидат юридичних наук, доцент, старший  
викладач кафедри кримінального права

Національного юридичного університету імені  
Ярослава Мудрого

ORCID: <https://orcid.org/0000-0003-1258-2063>

## **ФОРМУВАННЯ ЕТИЧНИХ НОРМ І СТАНДАРТІВ ПРИ ВИКОРИСТАННІ ШТУЧНОГО ІНТЕЛЕКТУ В УМОВАХ ВОЄННОГО СТАНУ**

Штучний інтелект кардинально змінює ведення сучасних війн роблячи їх швидкими і технологічними. Зростає небезпека автоматизованих систем, що не контролюються людиною. У зв'язку з цим використання штучного інтелекту в умовах воєнного стану набуває своїх особливостей, які можуть варіюватися в залежності від країни та специфіки конфлікту. Очевидно, що у воєнний час зосереджується увага на використанні штучного інтелекту у військових цілях. Зокрема, це стосується розробки автономних систем, дронів, систем розвідки та аналізу даних. Відповідно постає питання щодо урегулювання суспільних відносин, що забезпечують безпеку таких технологій та їх відповідність міжнародним нормам. У цьому ж контексті необхідно визначитися хто несе відповідальність за дії автономних систем, особливо якщо вони можуть призвести до цивільних втрат. Важливо окреслити етичні рамки використання штучного інтелекту в контексті кібербезпеки. Це включає в тому числі встановлення відповідальності за дії систем, що використовують штучний інтелект, особливо у випадках помилкових спрацьовувань або зловживань. Достатньо актуальними є питання пов'язані з тими ситуаціями, коли штучний інтелект використовується для прийняття ключових рішень у військових операціях. Щодо цього необхідно забезпечити нормативно-правову базу, яка регламентує прозорість алгоритмів та механізмів прийняття рішень, щоб уникнути зловживань і помилок у зазначеній сфері діяльності. У воєнний час інформація може бути особливо чутливою, впливати на багато процесів, які мають вирішальне значення для національної безпеки країни, її обороноздатності, недоторканості, суверенітету, територіальної цілісності, державної, економічної чи

інформаційної безпеки держави, тому нормативне упорядкування таких процесів необхідно направити на забезпечення врахування захисту даних та конфіденційності, особливо в контексті збору та обробки інформації за допомогою штучного інтелекту. Регулювання повинно створювати захист персональних даних, які використовуються для навчання та роботи алгоритмів штучного інтелекту, що включає в себе дотримання норм конфіденційності та захисту персональних даних. При цьому необхідно враховувати, що регулювання штучного інтелекту повинно бути гнучким і адаптивним до швидко змінюваного середовища воєнного стану, щоб мати можливість оперативно реагувати на нові виклики і загрози кібербезпеки, включаючи нові технології та методи нападу.

Таким чином, слід підкреслити важливість комплексного підходу до регулювання штучного інтелекту в умовах воєнного стану, що має в своєму змісті правові, етичні та технологічні вимоги.

Необхідно звернути увагу на те, що штучний інтелект може бути не лише інструментом для підвищення кібербезпеки, але й потенційною загрозою у тому випадку, коли він використовується для здійснення кіберзлочинів. Зловмисники можуть використовувати штучний інтелект для автоматизації процесів атаки, таких як злом паролів, фішинг або DDoS-атаки, що дозволяє їм швидше і ефективніше здійснювати напади на великі масштаби цілей[3, 52]. Штучний інтелект також використовується для генерації шкідливого контенту, в тому числі для створення переконливих фішингових листів або шкідливих веб-сайтів, які важче розпізнати[2, 68]. Наприклад, генеративні моделі можуть створювати тексти, які виглядають дуже правдоподібно і які важко відрізнити від оригіналу. Штучний інтелект використовується в кримінальних правопорушеннях для аналізу систем безпеки та адаптації своїх нападів та впливів, щоб обійти існуючі заходи захисту, включаючи вивчення поведінки системи та виявлення слабких місць. Наразі в сучасному світі вже достатньо широко розповсюджується явище, яке носить назву – соціальна інженерія[4, 22], коли штучний інтелект застосовується для збору даних про потенційні жертви через соціальні мережі та інші джерела, що дозволяє зловмисникам здійснювати більш персоналізовані та переконливі напади. При цьому створення фальшивих відео або аудіозаписів, як правило, ускладнює ідентифікацію справжніх фізичних осіб і може призвести до таких кримінальних правопорушень, як крадіжка, шахрайство, вимагання тощо.

Начебто позитивна ознака штучного інтелекту, яка теж досить часто використовується при наукових дослідженнях в різних галузях – здатність аналізувати великі обсяги даних для виявлення вразливостей у системах безпеки – зловмисниками використовується для планування кримінальних правопорушень, що завдають тяжкі наслідки. Також і системи на базі штучного інтелекту можуть бути використані для обходу традиційних механізмів захисту, таких як антивірусні програми або системи виявлення вторгнень, шляхом модифікації шкідливого коду.

Важливо зазначити, що всі ці загрози потребують активного реагування з боку фахівців з кібербезпеки, які повинні адаптувати свої стратегії та технології для протидії новим викликам, пов'язаним із використанням штучного інтелекту при вчиненні кримінальних правопорушень.

Протидіяти використанню штучного інтелекту в злочинних цілях, можна виключно із застосуванням комплексного підходу, що включає в себе як технологічні рішення, так і стратегії управління ризиками. Важливу роль в цьому відіграє розвиток адаптивних систем захисту. Наприклад, створення систем, які можуть виявляти аномалії в поведінці користувачів та аномалії чи нестабільність самої системи, що дозволить швидко реагувати на потенційні загрози впливу на певну сукупність суспільних відносин. Також інтеграція систем виявлення вторгнень (IDS) з алгоритмами штучного інтелекту для автоматичного аналізу та реагування на підозрілі дії. Важливе значення має також впровадження політик безпеки, що включають багатоетапну аутентифікацію (MFA) та регулярну зміну паролів; застосування технологій SIEM (Security Information and Event Management) для збору та аналізу даних з різних джерел; застосування блокчейну для забезпечення прозорості та незмінності записів, що може допомогти у боротьбі з шахрайством і підробками.

Створення етичних норм і стандартів для розробників штучного інтелекту, щоб запобігти зловживанню технологією безпосередньо пов'язане із впровадженням правових норм, які регулюють використання штучного інтелекту в кіберзлочинності. Саме з цією метою необхідно забезпечувати обмін інформацією про загрози між підприємствами, державними установами та правоохоронними органами для створення єдиної бази знань про нові загрози; приймати участь у міжнародних ініціативах і програмах, спрямованих на боротьбу з кіберзлочинністю. Необхідно забезпечувати регулярне

проведення тестування на проникнення (pentesting) для виявлення вразливостей у системах і їх усунення до того, як зловмисники зможуть їх використати, проводити інвестування в дослідження та розробку алгоритмів, які можуть виявляти фальшиві відео та аудіозаписи, щоб захистити користувачів від маніпуляцій.

Зараз вже стало нагальною проблемою вирішення питання щодо того хто несе відповідальність за дії штучного інтелекту у випадку помилок або несправностей, що є особливо актуальним тоді, коли алгоритми приймають автономні рішення? Не менш важливим є урегулювання проблемних моментів щодо використання військових технологій. Особливо якщо штучний інтелект застосовується в контексті військової кібербезпеки виникають питання про етику використання таких технологій у конфліктах і війнах[7, 424].

Які моральні межі можна встановити для використання штучного інтелекту в агресивних діях? Це питання є складним і багатогранним. Перш за все, слід звернути увагу на використання принципу пропорційності, відповідно якого застосування штучного інтелекту в агресивних діях має бути пропорційним суспільно-небезпечному посяганню або небезпеці, що загрожує правоохоронюваним інтересам. Такими насамперед є піддані небезпеці інтереси особи (напр., життя, здоров'я, тілесна недоторканність, особиста свобода, статевая свобода, майнові, житлові, політичні та інші охоронювані законом права та інтереси). Небезпека може загрожувати інтересам держави: зовнішній безпеці, обороноздатності, порядку управління, інтересам правосуддя, збереженню державної таємниці, майну тощо. В таких умовах необхідно оцінити, чи є застосування технології виправданим у конкретній ситуації і чи не призведе це до перевищення меж допустимої шкоди[8, 21].

Важливо підкреслити, що штучний інтелект необхідно розробляти так, щоб максимально зменшити ризик завдання шкоди цивільним особам, що має включати в себе точність націльності, визначення мети та можливість розрізняти військові та цивільні об'єкти. З цим тісно пов'язане питання щодо заборони автономних систем знищення [1, 15]. Серед науковців досить гостро ведеться дискусія щодо повністю автономних систем озброєння [5, 38], які можуть приймати рішення про вбивство без людського контролю. Багато із них вважають, що такі системи повинні бути заборонені. При цьому також висловлюється думка, що подібні питання щодо використання штучного інтелекту в агресивних діях слід зробити прозорими, а рішення, прийняті на

основі алгоритмів, мають бути підзвітними, що створить можливість перевірки рішень і дій, які приймаються штучним інтелектом самостійно [6, 12]. Безперечно, необхідно обмежити та ретельно контролювати дослідження і розробку технологій штучного інтелекту, які можуть бути використані для агресії, щоб запобігти їх використанню в небезпечний спосіб. Важливо забезпечити, щоб людина залишалася в центрі прийняття рішень у військових операціях, навіть якщо використовуються технології штучного інтелекту. Людський контроль має бути обов'язковим у критичних ситуаціях.

Очевидно, що вирішення таких питань має безпосередній вплив на міжнародну безпеку. Використання штучного інтелекту в агресивних діях може змінити баланс сил у світі, тому вже зараз необхідно враховувати, які це може мати наслідки для міжнародної безпеки та стабільності. Важливо визначити, хто несе моральну відповідальність за дії штучного інтелекту: розробники, військові командири чи уряди держав? Це питання потребує чітких відповідей для забезпечення справедливості та широкого обговорення з узгодженням на міжнародному рівні. Сьогодні на сучасному етапі розвитку суспільства терміново необхідно формувати етичні норми і стандарти у використанні штучного інтелекту в агресивних діях, в умовах воєнного стану. Саме вони мають створювати разом з цим також рівний доступ до технологій штучного інтелекту, що повинно стати важливим етичним зобов'язанням для країн і організацій. Вони мають прагнути до забезпечення справедливого доступу до технологій штучного інтелекту, щоб не посилювати існуючі соціальні та економічні нерівності, включаючи в себе надання можливостей для всіх верств населення, незалежно від їхнього соціального статусу, географічного положення чи економічних умов. Країни повинні співпрацювати для створення рівного доступу до технологій штучного інтелекту, обмінюючись знаннями, ресурсами та технологіями, при цьому важливу роль зобов'язані грати міжнародні організації, що будуть опікуватися саме цими питаннями [9, 57]. Саме вони мають запроваджувати політику відкритого доступу до даних і технологій, що буде дозволяти більшій кількості людей і організацій використовувати ці ресурси для свого розвитку та розвитку суспільства. Важливо залучати громади до процесу прийняття рішень щодо впровадження технологій штучного інтелекту, щоб враховувати їхні потреби та побоювання, що буде безпосередньо сприяти більш демократичному підходу до використання технологій. Кожен із тих, хто має

безпосереднє відношення до штучного інтелекту повинен визнавати свою соціальну відповідальність за вплив технологій штучного інтелекту на суспільство і працювати над тим, щоб цей вплив був позитивним і конструктивним.

### **Використані джерела**

1. Богом'я В., Гудзь А. Штучний інтелект: сучасний стан і перспективи застосування. Військова кібернетика та системний аналіз. 2023. Т. 46. № 1. С. 13-17. URL: <https://doi.org/10.33099/2311-7249/2023-46-1-13-17>
2. Вінникова Н.А. Штучний інтелект у контексті глобального управління. Politicus. 2022. Вип. № 3. С. 65-70. URL: <https://doi.org/10.24195/2414-9616.2022-3.10>
3. Корнєєва С.Р. Теоретичні підходи до визначення поняття та правового регулювання штучного інтелекту. Науковий вісник Ужгородського національного університету. Серія: Право. 2021. Т. 66. С. 50-55.
4. Кронівець Т.М. Правові та ціннісні особливості феномену штучного інтелекту як елементу правової дійсності. Науковий вісник Херсонського державного університету. Серія. Юридичні науки. 2022. Вип. 2. С. 21-23.
5. Свешніков С., Бочарніков В. Використання методів штучного інтелекту для розв'язання наукових задач у сфері оборони. Наука і оборона. 2023. № 4. С. 29-40.
6. Скіцько О., Складанний П., Ширшов Р., Гуменюк М. Загрози та ризики використання штучного інтелекту. Кібербезпека: освіта, наука, техніка. 2022. № 2 (22). С. 6-18.
7. Тимошенко Є.А. Правова природа штучного інтелекту: перспективи і проблеми. Юридичний науковий електронний журнал. 2023. № 4. С.424-425.
8. Хаустова В.Є., Решетняк О.І., Хаустов М.М., Зінченко В.А. Напрямки розвитку технологій штучного інтелекту в забезпеченні обороноздатності країни. БІЗНЕСІНФОРМ. 2022. № 3 С.17-26.
9. Цяпа С.М. Огляд зарубіжних законодавчих ініціатив стратегічного використання технологій штучного інтелекту в сучасних умовах. Інформація і право. № 2(37)/2021. С. 51-59.

**Федорченко О.С.**

молодший науковий співробітник,

Український науково-дослідний інститут

спеціальної техніки та судових експертиз СБУ

ORCID: <https://orcid.org/0009-0007-7358-7753>

## **ШТУЧНИЙ ІНТЕЛЕКТ ПРОТИ КІБЕРЗЛОЧИНЦІВ: ЕВОЛЮЦІЯ ЗАХИСТУ В ЕПОХУ СКЛАДНИХ КІБЕРАТАК**

Сучасний ландшафт кіберзагроз істотно розширився. У епоху цифрової трансформації кібератаки стають все складнішими, включаючи AI-генеровані фішинг, шкідливе програне забезпечення (ransomware) та zero-day вразливості. Питома вага кіберзагроз зростає і ця тенденція в міру розвитку інформаційних технологій та їх конвергенції з технологіями штучного інтелекту (ШІ) в найближче десятиліття посилюватиметься[1]. У 2025 році кібератаки досягли безпрецедентного рівня складності, включаючи AI-генерований фішинг, шкідливе програне забезпечення, zero-day експлойти та спонсоровані державою атаки.

Кіберпростір залишається одним із ключових напрямів гібридної агресії, а кібератаки дедалі частіше координуються із ракетно-дроновими ударами.

У першому півріччі 2025 року в Україні було зафіксовано 3018 кіберінцидентів, що на 17% більше, ніж у другій половині 2024 року. За даними CERT-UA, в середньому фіксується близько 15 кібератак на день [2]. Також відзначається зростання масштабів та складності кібератак..

Згідно з даними Асоціації галузі інформаційних технологій (CompTIA), глобальні витрати на кіберзлочини перевищать \$10.5 трлн щорічно, з річним зростанням на 10%. World Economic Forum повідомляє, що 72% респондентів відзначають зростання кіберризиків за останній рік, з акцентом на фішинг та соціальну інженерію [3].

Витрати, пов'язані з кіберзлочинністю, включають пошкодження та знищення даних, викрадені гроші, втрату продуктивності, крадіжку інтелектуальної власності, крадіжку особистих і фінансових даних, розкрадання, шахрайство, порушення роботи бізнесу після атаки, судові розслідування, відновлення та видалення скомпрометованих даних і систем, а також шкоду репутації [3].

Статистика свідчить про зростання кібератак на 30–50% щорічно, що вимагає еволюції захисних механізмів від кібератак різного ступеня складності.

На початку 2000-х років кіберсистеми базувалися на файрволах, антивірусах та сигнатурах, які реагували на відомі кіберзагрози. Водночас, такі традиційні системи були ефективними лише щодо простих кібератак. При цьому об’єкти кібератак залишалися вразливими до поліморфних вірусів та АРТ (Advanced Persistent Threats).

Перехід до машинного навчання (ML) у 2010-х роках дозволив аналізувати поведінку в реальному часі. Запровадження машинного навчання ML дозволило аналізувати поведінку в реальному часі, що зменшувало час реагування з днів до секунд. Як приклад, можна навести впровадження SIEM-систем з ML для кореляції логів.

Нині у сучасних умовах ІІІ еволюціонував до генеративних моделей, які прогнозують кібератаки. Американська компанія Fortinet, що спеціалізується на пристроях мережевої безпеки, зазначає, що ІІІ автоматизує виявлення загроз та посилює кібероборону проти еволюціонуючих ризиків. У 2025 році ІІІ стає основою для захисту об’єктів критичної інфраструктури, що актуалізує його ключову роль у виявленні та запобіганні кібератак. При цьому поширюється інструментарій, що передбачає накопичення великих масивів інформації щодо поведінки людини, соціальних груп та використання сучасних досягнень у сфері штучного інтелекту[1].

На нашу думку, перш за все, актуальним є застосування ІІІ у виявленні аномалій у мережевому трафіку, прогнозуванні кібератак та автоматизованому реагуванні на них.

У першому випадку ІІІ здатний використовувати нейромережі для аналізу трафіку в реальному часі, наприклад, у системах SIEM.

Що стосується прогнозування на основі аналізу великих масивів даних та автоматизованого реагування на кібератаки в режимі реального часу, то слід звернути увагу, що сьогодні ІІІ в змозі моделювати поведінку зловмисників для симуляцій. Тут можна згадати такі важливі інструменти, як SOAR (Security Orchestration, Automation and Response), котрі дозволяють збирати дані та попередження про небезпеку з різних джерел, а також допомагають інтеграції ІІІ для симуляції кібератак. В цьому контексті фахівці прогнозують, що у 2025

році ШІ сприятиме виявленню нових кіберзагроз та пом'якшенню наслідків вже існуючих.

Водночас, ШІ не обмежується лише цими напрямками застосування. Зокрема, ШІ у сфері кібербезпеки може бути застосований для: автоматичного аналізу подій у кіберпросторі; захисту від несанкціонованого доступу до інформаційних ресурсів; аналізу і структуризації інформації про кіберзагрози; розробки передових алгоритмів, які допомагають виявляти кібератаки та запобігання їх негативним проявам; машинного навчання на підставі набутого досвіду.

На сучасному етапі ШІ демонструє переваги над традиційними методами забезпечення кібербезпеки. Серед його переваг доцільно виділити: високу швидкість обробки даних; можливість одночасного аналізу великої кількості факторів; автономне реагування на загрози без участі людини; наявність механізмів самонавчання; прогнозування кіберризиків.

Слід також відзначити, що ШІ легко адаптується до нових кіберзагроз (без ручного оновлення), обробляючи значні масиви даних щодня. Вибіркові дослідження свідчать, що ШІ також зменшує помилкові результати аналізу масиву даних на 70–80%.

Поряд з вражаючими перспективами, які відкриваються з використанням ШІ, багато фахівців звертають увагу на ризики, пов'язані з розвитком та масштабами використання ШІ. Найбільш поширеним є підхід, згідно з яким зміст ризиків застосування ШІ зумовлює: високу залежність від якості тренувальних даних; маніпуляції, інші злочинні дії з боку хакерів (adversarial attacks); проблеми з поясненням рішень ШІ (black box problem); порушення права на конфіденційність та захист персональних даних; етичні питання, пов'язані з автономними системами [4, с. 43].

Одним із головних викликів є те, що системи ШІ самі можуть стати мішенню кібератак. Наприклад, змагальні атаки на алгоритми машинного навчання, спрямовані на модифікацію вхідних даних, щоб викликати помилки в їх роботі [5, с. 23]. Зловмисники можуть навмисно маніпулювати або підміняти ці дані, що призводить до небажаних результатів. Такі атаки становлять серйозну проблему для впровадження машинного навчання в критично важливі системи [6, с. 13]. Часто зловмисники використовують AI для обману систем (наприклад, data poisoning або evasion attacks). Вразливими є окремі моделі ШІ до маніпуляцій, зокрема AI-моделі.

Одним із викликів для України у сфері кібербезпеки є змагальний характер розвитку засобів кібербезпеки в умовах швидких прогресуючих змін інформаційно-комунікаційних технологій, зокрема хмарних та квантових обчислень, 5G-мереж, великих даних, Інтернету речей, штучного інтелекту тощо [1].

Але, не зважаючи на ризики та загрози використання ШІ у сфері кібербезпеки, інвестиції у нього з кожним роком збільшуються. Глобальний ринок ШІ-рішень може сягнути \$1 трлн у 2027 році, а фахівці прогнозують зростання витрат на AI-рішення до \$135 млрд до 2030 р. [8].

Викладене дозволяє зробити висновок, що ШІ є одним із ефективних заходів забезпечення кібербезпеки з потужним потенціалом у майбутньому.

На нашу думку, найбільш перспективним з точки зору забезпечення кібербезпеки є гібридний підхід, який поєднує ШІ з людським інтелектом для протидії складних кібератак.

#### **Список використаних джерел:**

1. Стратегія кібербезпеки України: затвердж. Указом Президента України від 26 серпня 2021 р. № 447. URL: <https://zakon.rada.gov.ua/laws/show/447/2021#Text>.
2. З початку 2025-го в Україні фіксується близько 15 кібератак щодня — Держспецв'язку URL: <https://suspilne.media/1075891-z-pocatku-2025-go-v-ukraini-fiksuetsa-blizko-15-kiberatak-sodna-derzspeczvazku/>
3. Якщо уявити кіберзлочинність як країну це була б третя економіка світу. URL: <https://10guards.com/ua/blog/2022/08/08/cybercrime-as-empire-would-be-the-worlds-third-largest-economy/>
4. Федорченко О.С. Роль штучного інтелекту у забезпеченні кібербезпеки України: сучасний стан та перспективи розвитку. *Інформація і право*. №3 (54). 2025. С. 139-146.
5. Dilek S., Cakır H., Aydın M. Applications of Artificial Intelligence Techniques to Combating Cyber Crimes: A Review. *International Journal of Artificial Intelligence & Applications*. 2015. Vol. 6, no. 1. P. 21–39. DOI: <https://doi.org/10.5121/ijaia.2015.6102>.
6. Класифікації моделей застосування машинного навчання у кібербезпеці. А.В. Антоненко та ін. *Таврійський науковий вісник*. Серія: Технічні науки. 2023. №4. С. 11–22. DOI: <https://doi.org/10.32782/tnv-tech.2023.4.2>.
7. Фішинг, Zero Click-експлойти, AI-атаки: тенденції кіберзагроз в Україні у 2025 році. URL: <https://itukraine.org.ua/fishing-zero-click-eksplotji-ai-ataki-tendentsiyi-kiberzagroz-v-ukrayini-u-2025-rotsi/>
8. Глобальний ринок ШІ-рішень може сягнути \$1 трлн у 2027 році URL: <https://denovo.ua/blog/one-trillion-ai-market-2027>

**Ганна Форос**

кандидат юридичних наук, доцент, завідувачка  
кафедри кримінального аналізу та  
інформаційних технологій

Одеський державний університет внутрішніх  
справ

ORCID: <https://orcid.org/0000-0002-9504-3681>

## **АВТОМАТИЗОВАНИЙ МОНІТОРИНГ КІБЕРЗАГРОЗ ЗА ДОПОМОГОЮ СИСТЕМ МАШИННОГО НАВЧАННЯ**

Активна цифровізація економічних процесів, державного управління та соціальної взаємодії зумовила суттєве зростання залежності суспільства від стабільного функціонування інформаційних систем. Збільшення кількості пристроїв, що підключені до мережі, формування розподілених хмарних середовищ і розвиток Інтернету речей призвели до різкого ускладнення ландшафту кіберзагроз. Кіберзлочинність набуває все більш структурованого та технологічно оснащеного характеру, а її інструменти швидко еволюціонують. У цих умовах питання забезпечення кібербезпеки перетворюється на ключовий аспект інформаційної політики держави та корпоративного управління.

Традиційні методи виявлення атак, що ґрунтуються на сигнатурному аналізі або фіксованих правилах реагування, поступово втрачають ефективність. Причиною є динамічність появи нових шкідливих кодів, варіативність поведінки зловмисників і використання складних технік ухилення. Тому виникає потреба у впровадженні адаптивних інтелектуальних систем, здатних виявляти відхилення від нормальної поведінки та прогнозувати потенційні інциденти без необхідності попереднього знання сигнатур атак.

Автоматизований моніторинг кіберзагроз є процесом безперервного збору, фільтрації, аналізу й кореляції подій у мережевому трафіку, системних журналах та інших джерелах телеметрії з метою своєчасного виявлення підозрілої активності. Системи машинного навчання забезпечують новий рівень аналітики, оскільки здатні опрацьовувати великі масиви даних, визначати приховані залежності між подіями та виявляти аномальні патерни,

що можуть свідчити про спробу несанкціонованого доступу або порушення політики безпеки.

Одним із найбільш результативних підходів є застосування алгоритмів класифікації, які дозволяють віднести вхідні дані до певної категорії ризику. До них належать Random Forest [1], Decision Trees [2], Support Vector Machines [3], а також ансамблеві методи, які комбінують результати кількох моделей для підвищення точності прогнозування. Для задач виявлення складних, багатоступеневих атак ефективно використовуються методи глибинного навчання, що базуються на штучних нейронних мережах. Автоенкодери допомагають ідентифікувати невідомі раніше типи загроз, а рекурентні мережі застосовуються для аналізу послідовностей подій у часових рядах журналів безпеки.

Методи кластеризації, зокрема k-means, DBSCAN і Hierarchical Clustering, дозволяють групувати схожі події, виділяючи потенційні аномалії без попереднього маркування даних. Це особливо важливо для систем із великим обсягом трафіку, де ручне маркування або формування навчальних вибірок є неможливим. Таким чином, поєднання різних алгоритмічних підходів забезпечує створення багаторівневої системи аналізу даних, здатної адаптуватися до змін кіберсередовища.

Окрему роль відіграють методи обробки природної мови (Natural Language Processing) [4], що використовуються для автоматичного аналізу текстових повідомлень, звітів про вразливості та відкритих джерел інформації. Завдяки цим технологіям система може виявляти згадки про нові експлойти, індикатори компрометації або ознаки підготовки кібератаки ще до її фактичної реалізації. Це суттєво підвищує проактивність реагування та дозволяє прогнозувати потенційні ризики.

Побудова автоматизованої системи моніторингу кіберзагроз потребує розроблення архітектури, що включає кілька взаємопов'язаних рівнів. Перший рівень забезпечує збір та агрегацію даних із мережевих сенсорів, систем контролю доступу, журналів безпеки, SIEM-платформ [5] і зовнішніх джерел. Другий рівень відповідає за попередню обробку, нормалізацію та очищення даних від шумів. Третій рівень реалізує аналітичні моделі машинного навчання, які виконують класифікацію, прогнозування та детектування аномалій. Четвертий рівень охоплює компоненти візуалізації, формування звітів і автоматизоване інформування про потенційні інциденти безпеки.

Інтеграція автоматизованих систем моніторингу з іншими компонентами кіберзахисту — міжмережевими екранами, IDS/IPS, платформами керування інцидентами — забезпечує створення єдиного інформаційного простору, у якому рішення приймаються на основі комплексного аналізу даних. Це дозволяє значно скоротити час виявлення атак і зменшити навантаження на аналітиків безпеки.

Подальші дослідження у цьому напрямі орієнтовані на створення гібридних моделей, що об'єднують алгоритми класичного машинного навчання, глибинних нейронних мереж та генеративних моделей. Використання генеративно-змагальних мереж (Generative Adversarial Networks, GAN) відкриває можливість імітації поведінки зловмисників і формування навчальних вибірок із синтетичних даних, що сприяє підвищенню стійкості систем до нових типів атак. Особливої уваги заслуговує аспект інтерпретованості результатів — здатність пояснити рішення, прийняті моделлю. Це критично важливо для забезпечення довіри до автоматизованих систем кіберзахисту та для впровадження їх у критичну інфраструктуру.

Отже, автоматизований моніторинг кіберзагроз на основі методів машинного навчання формує нову парадигму забезпечення кібербезпеки, орієнтовану на аналітичну обробку великих обсягів даних, самонавчання та адаптацію до змін середовища. Реалізація подібних систем дозволяє підвищити точність виявлення загроз, зменшити час реагування на інциденти, оптимізувати використання ресурсів служби безпеки й створити динамічну модель захисту інформаційних ресурсів у мінливому кіберпросторі.

### Список використаних джерел

1. Salman, Hasan Ahmed, Ali Kalakech, and Amani Steiti. "Random Forest Algorithm Overview." *Babylonian Journal of Machine Learning* (June 2024): 69–79. <https://doi.org/10.58496/BJML/2024/007>. 2024. Вип. 30. С. 67–77. DOI: <https://doi.org/10.32447/20784643.30.2024.07>
2. «Decision Tree – an Overview». *ScienceDirect Topics: Computer Science*. <https://www.sciencedirect.com/topics/computer-science/decision-tree-model>
3. «Support Vector Machine». *ScienceDirect Topics: Engineering*. Accessed October 30, 2025. <https://www.sciencedirect.com/topics/engineering/support-vector-machine>

4. Abuya, Teresa Kwamboka, Christal Anto V, Alexander Mutiso Mutua, and Richard Rimiru. *Natural Language Processing*. August 2025. Zenodo. <https://doi.org/10.5281/zenodo.16917578>. ISBN 978-93-6786-839-3.

5. Mahajan, Shilpa. «The Role of SIEM in Modern Cybersecurity: A Comprehensive Analysis». In *Cybersecurity and Privacy in the Era of Smart Technologies*, September 2025. <https://doi.org/10.4018/979-8-3373-2282-7.ch010>

## **ЦИФРОВИЙ СУВЕРЕНІТЕТ І НАЦІОНАЛЬНА СТРАТЕГІЯ РОЗВИТКУ У СФЕРІ ШТУЧНОГО ІНТЕЛЕКТУ**

**Авдіюк С. С.**

Студентка 1 курсу ОС «Магістр»

ННІ права КНУ ім. Тараса Шевченка

**Габелко В. О.**

Студентка 1 курсу ОС «Магістр»

ННІ права КНУ ім. Тараса Шевченка

ORCID: <https://orcid.org/0009-0009-3947-8873>

**Драчук Є.М.**

Студентка 1 курсу ОС «Магістр»

ННІ права КНУ ім. Тараса Шевченка

### **ПРАВОВІ МЕХАНІЗМИ ЗАБЕЗПЕЧЕННЯ ЦИФРОВОГО СУВЕРЕНІТЕТУ УКРАЇНИ В ЕПОХУ ШТУЧНОГО ІНТЕЛЕКТУ**

У сучасну епоху цифрових технологій здатність держави контролювати власну інфраструктуру, дані, алгоритми та програмне забезпечення перетворюється в один із ключових елементів національної безпеки. Якщо раніше суверенітет держави розглядався переважно у фізичних і територіальних вимірах, то нині він трансформується у нову площину – у форму цифрового суверенітету, тобто здатності держави до автономного управління своїми цифровими ресурсами, прийняття рішень щодо них і захисту [1].

В умовах глобалізації, стрімкого розвитку мережевих технологій і масового впровадження штучного інтелекту (далі – ШІ) у всі сфери суспільного життя цифровий суверенітет набуває стратегічного значення. Аналітики Всесвітнього економічного форуму зазначають, що цифровий суверенітет означає «володіння власним цифровим майбутнім», включаючи контроль над даними, апаратним та програмним забезпеченням [2].

Роль штучного інтелекту у стратегічних і безпекових процесах стрімко зростає: алгоритми приймають рішення, аналізують великі масиви інформації, функціонують у хмарних інфраструктурах, що створює нові вектори залежності держави від іноземних технологій і постачальників. Відповідно,

контроль над інфраструктурою ІІІ стає невід'ємною частиною реалізації державного цифрового суверенітету [3].

Глобальна тенденція свідчить, що дедалі більше держав упроваджують політику локалізації даних, створюють національні хмарні середовища, формують нормативні рамки для забезпечення цифрової автономії. У наукових дослідженнях підкреслюється, що цифрова автономія вже не лише технічне, а й політичне питання, пов'язане із забезпеченням суверенності та національного контролю [4].

Операціоналізація цифрового суверенітету України у 2025 році забезпечується урядовим Планом заходів, що передбачає створення кібервійськ, налагодження міжвідомчої взаємодії та кібероборони, регулярні навчання з НАТО/ЄС, розвиток національних CSIRT-спроможностей, впровадження ризик-орієнтованої оцінки кіберризиків, розроблення національних стандартів і посилення криптографічного захисту державних ресурсів [5].

Таким чином, актуальність теми дослідження зумовлюється тим, що цифровий суверенітет стає обов'язковою складовою національної безпеки та стратегічної автономії держави. У сучасних умовах глобальної цифровізації, зростання ролі штучного інтелекту та поширення гібридних кібератак контроль над національними даними, цифровою інфраструктурою й алгоритмами набуває вирішального значення для збереження державного суверенітету. Для України це питання є особливо гострим через воєнний стан, кіберагресію та значну залежність від іноземних технологічних рішень.

Фундаментом системи цифрового суверенітету є Конституція України, яка у статті 17 визначає забезпечення інформаційної безпеки як найважливішу функцію держави, а у статті 32 гарантує захист конфіденційної інформації, дозволяючи обмеження лише в інтересах національної безпеки. Стаття 92 підкреслює, що питання національної безпеки мають регулюватися виключно законами [6].

На рівні галузевого законодавства центральне місце посідає Закон України «Про основні засади забезпечення кібербезпеки України» (далі – Закон) від 2017 року. Цей акт заклав ключові поняття та розподілив суб'єктний склад системи кібербезпеки [7]. Проте дослідники зазначають, що Закон залишається переважно рамковим документом, а спроба розподілити сфери відповідальності державних органів виявилася не досить вдалою [8, с. 279].

Актуальною проблемою є наявність колізій і прогалин, що вимагає гармонізації українського законодавства з міжнародними стандартами. Доповнює базу Закон України «Про національну безпеку України», який інтегрує кібербезпеку в загальну систему оборони держави [9].

Реалізація законодавства відбувається через підзаконні акти — рішення РНБО, постанови КМУ та накази ДССЗІ. Особливо важливою є Постанова КМУ №518 від 19.05.2022 про державну систему реагування на кіберінциденти. Водночас дослідники наголошують на необхідності розробки спеціального порядку забезпечення кібербезпеки державних ресурсів у хмарних технологіях [8, с. 284].

Інституційна архітектура цифрового суверенітету охоплює низку ключових суб'єктів. РНБО здійснює стратегічну координацію та контроль [7]. ДССЗІ відповідає за реалізацію державної політики у сфері захисту державних інформаційних ресурсів, а підпорядкований їй CERT-UA виконує функцію центру реагування на кіберінциденти [7]. Міністерство цифрової трансформації впроваджує AI Governance та створює державні хмарні платформи для локалізації критичних даних [10, 11].

Проте система стикається зі значними проблемами: недостатній рівень координації між суб'єктами та неузгодженість їхньої компетенції. Ця фрагментарність знижує ефективність реагування на загрози. Важливим вектором посилення стійкості є співпраця з приватним сектором через механізми державно-приватного партнерства та залучення вітчизняних ІТ-компаній до захисту критичної інфраструктури.

Стратегія цифрового розвитку інноваційної діяльності України на період до 2030 року визначає досягнення цифрового суверенітету як стратегічну ціль, що передбачає розвиток національних технологій, зокрема штучного інтелекту, та мінімізацію використання іноземних компонентів для захисту критичної інфраструктури та державних ресурсів [12]. Концепція розвитку штучного інтелекту в Україні від 2021 року окреслила необхідність правового контролю за системами ШІ, розробленими за кордоном, особливо у сферах національної безпеки та оборони [13]. Нерегульоване використання іноземних алгоритмів створює ризики несправедливості та обмежень прав людини [14, с. 37].

В той же час, відсутність чітких процедур огляду стану кіберзахисту ускладнює реалізацію стратегічних цілей [8, с. 279].

Найгострішим викликом є використання іноземних ІІІ-сервісів державними установами. Несанкціоноване використання таких сервісів створює прямий ризик витоку конфіденційних даних до юрисдикцій, які можуть мати конфлікт інтересів з Україною. Делегування функції зберігання критичних даних іноземним провайдером фактично обмежує суверенні права держави.

Масштабні кібератаки на «Дію», реєстри Мін'юсту та Prozorro підкреслюють роль CERT-UA у оперативному реагуванні [15]. Водночас вони висвітлюють системну проблему: фрагментарність організаційно-правових механізмів та відсутність спеціальних заходів припинення протиправної поведінки у сфері кібербезпеки.

Процес міждержавної взаємодопомоги набирає все динамічнішого розвитку, все більша кількість країн Європейського Союзу охоче підписують документи, угоди і меморандуми, що закріплюють стосунки у сфері цифрового суверенітету, кібербезпеки та цифрових технологій. Найбільш показовими вбачаємо угоду «Working Arrangement» з Європейським агентством з кібербезпеки (ENISA), що має на меті розширити правове та технічне забезпечення шляхом взаємного обміну досвідом та національним законодавством [16]. Відмітимо, що Україна ратифікувала Будапештську конвенцію про кіберзлочинність (2005 р.) (EST no.185), яка дозволяє державі інтегруватися в міжнародні зусилля з боротьби з кіберзлочинністю та забезпеченням електронних доказів [17]. У 2022 році Україна підписала Другий додатковий протокол до Будапештської конвенції, який в стосується обміну цифровими доказами [18]. Крім того, Україна підписала Рамкову конвенцію Ради Європи щодо штучного інтелекту, прав людини, демократії та верховенства права у 2023 році, що дозволило розробляти правові норми для використання ІІІ у відповідності з міжнародними стандартами безпеки [19]. Україна має тривалу співпрацю з Організацією економічного співробітництва та розвитку (OECD) в сферах цифровізації та кібербезпеки. В рамках цієї співпраці Україна отримує аналітичну підтримку щодо цифрових ризиків та розвитку політик в галузі кібербезпеки. У 2021 році Україна підписала "Ukraine Country Programme" з OECD, що включає рекомендації щодо вдосконалення національних стандартів кібербезпеки та розвитку інфраструктури [20].

Євроінтеграційний вектор розвитку та забезпечення кібербезпеки України полягає не лише в імплементації профільного європейського законодавства, а й, насамперед, в його застосуванні громадянами та державою. 16 жовтня 2025 року Європейський Союз, за участі служб Європейської Комісії, її Агентства з кібербезпеки (ENISA), Європейського коледжу безпеки та оборони та Консультативної місії Європейського Союзу (EUAM), Європолу та Україна – Рада національної безпеки і оборони, Державна служба спеціального зв'язку та захисту інформації, Центр кібербезпеки Служби безпеки та Генеральний штаб Збройних Сил та інші представники провели свій 4-й Кібердіалог у Києві, Україна, підтвердивши свою відданість співпраці в аспекті кібербезпеки та кібердипломатії. Актуальні проблеми, які були винесені на обговорення, стосувалися поточної ситуації в галузі кіберзагроз, визначення основи для подальшої співпраці в сфері кібербезпеки. Сторони обговорили можливості більш тісного узгодження своїх політик у цій галузі, зокрема, щодо цифрового розділу Плану України, захисту критичної інфраструктури, а також імплементації Директиви ЄС про мережеву та інформаційну безпеку (NIS2).

Діалог також зосередився на оперативних питаннях, зокрема на відповідних передових практиках та отриманих уроках щодо стримування та реагування на кіберінциденти та кібератак. Сторони розглянули механізми реагування, що складаються з санкційних режимів, а також співпрацю у розслідуванні кіберзлочинів. Крім того, було обговорено статус доступу України до Резерву кібербезпеки ЄС та співпрацю в рамках розвитку кіберпотенціалу, зокрема через Талліннський механізм [21].

У висновку зазначимо, що наразі, зі стрімким розвитком глобалізованого цифрового середовища, Україні необхідно виконати наступні дії:

По-перше, закріпити поняття «цифрового суверенітету» в національному законодавстві, окреслити його поняття, елементи, способи захисту, що надасть змогу чітко окреслити правила використання іноземних технологічних надбань, при цьому дозволить розмежовувати національні пріоритети зі світовою інтеграцією в глобальні кіберпростори.

По-друге, доцільним вважаємо створення профільного законодавства у контексті використання іноземних програм та застосунків у державних установах з метою мінімізації ризиків шпигунства іноземними державами. Вважаємо, що законодавець має запровадити чітку структуру з оцінювання

безпеки програмного забезпечення і національного суверенітету, а також, його сертифікацію.

По-третє, Україна має розробити стратегії цифрової автономії, синхронізованою з політикою Європейського Союзу. Разом з тим, така стратегія має включати механізми посилення кібербезпеки, забезпечення національної безпеки завдяки власним технологічним інфраструктур. Положення мають відображати прагнення держави до співпраці у межах Європи без шкоди власній незалежності.

По-четверте, одним з найважливіших принципів, яким має керуватися правотворець – баланс інтересів, що полягає у відкритості до міжнародного досвіду, практики та елементів законодавства з безпекою національного кіберпростору. Пріоритетним є встановлення чітких критеріїв для імпорту іноземних технологій; розробка першочергових заходів для моніторингу та оцінки їх впливу на національну кібербезпеку та кіберсуверенітет. На нашу думку, створення механізмів, які дозволять інтегрувати іноземні технології, одночасно мінімізує можливі ризики для національної безпеки, включаючи механізми перевірки сумісності з політиками кібербезпеки та захищає критичної інфраструктури.

### **Список використаних джерел:**

1. Timmers P. Cybersecurity and digital sovereignty – bridging the gaps [Електронний ресурс] / Peter Timmers. – TNO, 2024. – URL: <https://publications.tno.nl/publication/34643188/DvSKsfCM/timmers-2024-cybersecurity.pdf> (дата звернення: 03.11.2025 р.)
2. What is digital sovereignty and how are countries approaching it? [Електронний ресурс] // World Economic Forum. – 2025. – URL: <https://www.weforum.org/stories/2025/01/europe-digital-sovereignty/> (дата звернення: 03.11.2025р.)
3. Couture S., Toupin S. Digital sovereignty and artificial intelligence: a normative approach [Електронний ресурс] // Ethics and Information Technology. – Springer, 2024. – URL: <https://link.springer.com/article/10.1007/s10676-024-09810-5> (дата звернення: 03.11.2025 р.)
4. Frontiers in Political Science. Digital autonomy and state sovereignty in the era of AI [Електронний ресурс]. – 2025. – URL: <https://www.frontiersin.org/articles/10.3389/fpos.2025.1548562/full> (дата звернення: 05.11.2025 р.)

5. Про затвердження плану заходів на 2025 рік з реалізації Стратегії кібербезпеки України: розпорядження Кабінету Міністрів України від 07.03.2025 р. № 204-р. Дата оновлення: 08.03.2025. URL: <https://www.kmu.gov.ua/npas/pro-zatverdzhennia-planu-zakhodiv-na-2025-rik-z-realizatsii-stratehii-kiberbezpeky-ukrainy-204r> (дата звернення: 03.11.2025)

6. Конституція України (за редакцією від 01.01.2020) URL: <https://www.president.gov.ua/documents/constitution> (дата звернення: 06.11.2025 р.)

7. Закон України «Про основні засади забезпечення кібербезпеки України» (за редакцією від 19.10.2025). URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text> (дата звернення: 06.11.2025 р.)

8. Семенченко А.І., Плєскач В.Л., Заярний О.А., Плєскач М.В. Організаційно-правові механізми державного управління забезпеченням кібербезпеки та кіберзахисту України: сутність, стан та перспективи розвитку. ISSN 1727-4907. Проблеми програмування. 2020. № 2–3. Спеціальний випуск. URL: <https://doi.org/10.15407/pp2020.02-03.278> (дата звернення: 06.11.2025 р.)

9. Закон України «Про національну безпеку України» (за редакцією від 05.10.2025). URL: <https://zakon.rada.gov.ua/laws/show/2469-19#Text> (дата звернення: 05.11.2025 р.)

10. Про схвалення Концепції розвитку штучного інтелекту в Україні: розпорядження Кабінету Міністрів України від 2 грудня 2020 р. № 1556-р (за редакцією від 29.12.2021) URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text> (дата звернення: 05.11.2025 р.)

11. Про схвалення Концепції розвитку цифрових компетентностей та затвердження плану заходів з її реалізації: Розпорядження Кабінету Міністрів України від 3 березня 2021 р. № 167-р. URL: <https://zakon.rada.gov.ua/laws/show/167-2021-%D1%80#Text> (дата звернення: 04.11.2025 р.)

12. 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 грудня 2024 р. № 1351-р (за редакцією від 14.07.2025) URL: <https://zakon.rada.gov.ua/laws/show/1351-2024-%D1%80#n14> (дата звернення: 05.11.2025 р.)

13. Про затвердження плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2021-2024 роки від 25.07.2023. URL: <https://zakon.rada.gov.ua/laws/show/438-2021-%D1%80#Text> (дата звернення: 05.11.2025 р.)

14. Заярний О.А. До питання удосконалення способів захисту інформаційних прав фізичних осіб у відносинах, пов'язаних із застосуванням технологій штучного інтелекту. ISSN 1728-3817. Вісник Київського національного університету імені Тараса Шевченка. – Юридичні науки. 1(120)/2021 – С. 36-39. URL: <https://doi.org/10.17721/1728-2195/2022/1.120-7> (дата звернення: 06.11.2025 р.)

15. Національний координаційний центр кібербезпеки. Cyber Digest. Огляд подій в сфері кібербезпеки, серпень 2023. URL: [chrome-extension://efaidnbmnnnibpcajpcgklclefindmkaj/https://www.rnbo.gov.ua/files/2023\\_YEAR/CYBERCENTER/september/Cyber%20digest\\_August\\_2023\\_UA.pdf](chrome-extension://efaidnbmnnnibpcajpcgklclefindmkaj/https://www.rnbo.gov.ua/files/2023_YEAR/CYBERCENTER/september/Cyber%20digest_August_2023_UA.pdf) (дата звернення: 04.11.2025 р.)

16. Разом – на захисті кіберпростору: Держспецзв'язку, НКЦК та ENISA підписали угоду про співпрацю. URL: <https://cip.gov.ua/ua/news/razom-na-zakhisti-kiberprostoru-derzhspeczv-yazku-nkck-ta-enisa-pidpisali-ugodu-pro-spivpracyu> (дата звернення: 04.11.2025 р.)

17. Конвенція про кіберзлочинність, ратифікований від 07.09.2005. URL: [https://zakon.rada.gov.ua/laws/show/994\\_575#Text](https://zakon.rada.gov.ua/laws/show/994_575#Text) (дата звернення: 04.11.2025 р.)

18. Додатковий протокол до Конвенції про кіберзлочинність, який стосується криміналізації дій расистського та ксенофобного характеру, вчинених через комп'ютерні системи, ратифікований від 27.07.2006. URL: [https://zakon.rada.gov.ua/laws/show/994\\_687](https://zakon.rada.gov.ua/laws/show/994_687) (дата звернення: 06.11.2025 р.)

19. Рамкова конвенція Ради Європи щодо штучного інтелекту, прав людини, демократії та верховенства права URL: [chrome-extension://efaidnbmnnnibpcajpcgklclefindmkaj/https://vkksu.gov.ua/sites/default/files/rye\\_ramkova\\_konvenciya\\_zi\\_shtuchnogo\\_intelektu\\_prav\\_lyudyny\\_demokratiyi\\_ta\\_verh.\\_prava.pdf](chrome-extension://efaidnbmnnnibpcajpcgklclefindmkaj/https://vkksu.gov.ua/sites/default/files/rye_ramkova_konvenciya_zi_shtuchnogo_intelektu_prav_lyudyny_demokratiyi_ta_verh._prava.pdf) (дата звернення: 06.11.2025 р.)

20. Ukraine Country Programme (OECD). URL: <chrome-extension://efaidnbmnnnibpcajpcgklclefindmkaj/https://www.oecd.org/content/dam/oecd/en/about/programmes/grc/grc-eurasia/Ukraine-Country-Programme.pdf> (дата звернення: 06.11.2025 р.)

21. EU and Ukraine deepen cooperation on cyber security at 4th Cyber Dialogue. URL: <https://digital-strategy.ec.europa.eu/en/news/eu-and-ukraine-deepen-cooperation-cyber-security-4th-cyber-dialogue> (дата звернення: 06.11.2025 р.)

**Людмила Заславська**

старший науковий співробітник Державної  
наукової установи «Інститут інформації,  
безпеки і права Національної академії правових  
наук України»

ORCID: <https://orcid.org/0000-0001-6951-2627>

## **НАЦІОНАЛЬНА ВЕЛИКА МОВНА МОДЕЛЬ ЯК ІНСТРУМЕНТ ЦИФРОВОГО СУВЕРЕНІТЕТУ УКРАЇНИ**

Сучасний глобальний розвиток передбачає не лише підвищення соціально-економічного розвитку та вирішення комплексу важливих питань життєдіяльності людини, суспільства і держави, а насамперед – стрімкий розвиток цифрових технологій. Технології штучного інтелекту (ШІ) це наша реальність, якою ми користуємося вже сьогодні.

Для пересічної людини розвиток технологій штучного інтелекту значно допомагає у вирішенні різнобічних питань професійної діяльності, навчанні, обслуговуванні, медицині, бізнесі тощо. ШІ – це не просто технологія, а рушійна сила цифрової трансформації, яка змінює спосіб нашої роботи, навчання та взаємодії зі світом. Проте, мало хто задумується, яку небезпеку, ризики та загрози можуть нести в собі ці технології.

GPT-4, Claude, Gemini, LLaMA та ін. системи ШІ побудовані з допомогою великих мовних моделей (LLM), що дає змогу використовувати їх в чат-ботах, пошукових системах, аналітиці тощо. Велика мовна модель – це неймережа, яка працює за принципом людського мозку і здатна аналізувати та генерувати тексти. Україна, впевнено тримаючи курс цифрової трансформації, має на меті створити свою україномовну модель національного масштабу. На сьогодні розробкою національної великої мовної моделі займається Міністерство цифрової трансформації України спільно з компанією Київстар. Для України це важливо, тому що власна LLM :

- розумітиме українців краще за іноземні аналоги (знатиме наш контент, діалекти, терміни, історію та культуру, тож зможе відповідати більш якісно);
- зробить ШІ доступним для державного сектору (покращить сервіси для людей і зробить державу ефективнішою);

- стане основою для оборонного ШІ (забезпечить технологічну незалежність і захист даних – без виведення за межі України) [1].

Власна LLM модель дозволить зберігати дані в межах України, не передаючи їх закордонним корпораціям, дасть змогу уникати залежності від іноземних ШІ-рішень, які можуть не враховувати український контекст або бути недоступними в критичні моменти, а також сприятиме забезпеченню кібербезпеки та захисту персональних даних та приватного життя громадян.

Слушними видаються висновки, що *«Українська LLM це не просто технологічний проєкт. Це цифровий код нації, який має навчитися мислити і відчувати українською. Як і кожен інтелект, вона засвоює світ через тексти, якими її “годують”. Тому якість датасетів – це не лише питання точності перекладу, а питання моральної та культурної безпеки»* [2]. Якщо мовна модель навчається на русифікованих або спотворених джерелах, вона не просто повторює ці викривлення – вона вбудовує їх у мислення тих, хто з нею взаємодіє. Отже, створення української LLM – це насамперед акт державного самозахисту від культурної окупації.

Під час повномасштабного вторгнення РФ, для України надзвичайно важливим є можливість ефективно протистояти ворожій пропаганді та маніпулятивним наративам, що активно поширюються в інформаційному просторі. Розробники стверджують, що власна мовна модель допоможе уникнути вбудованих упереджень, притаманних іноземним системам, і гарантуватиме, що ШІ не лише коректно оброблятиме українську мову, а й не відтворюватиме шкідливих наративів, вигідних державі-агресору.

Розробники української LLM обговорюють такі напрями як, *технології, лінгвістика, етика, право*. Щодо права, то ця модель повинна не лише його знати, але й діяти відповідно до його принципів, а саме: *відкритості, прозорості, недискримінації, пропорційності, дотримання принципів верховенства права та прав людини*.

Розглянемо основні правові виклики, що можуть виникнути при створенні та використанні національної LLM.

### ***1. Питання інтелектуальної власності.***

Мовні моделі потребують величезних обсягів текстових даних для навчання. Ці дані можуть містити матеріали, захищені авторським правом – літературні твори, журналістські тексти, наукові статті тощо. Основна правова дилема полягає у визначенні, чи є використання таких даних для навчання

"добросовісним використанням" (fair use), чи порушенням авторських прав. Відсутність чіткого законодавчого регулювання використання таких даних у процесах машинного навчання може створювати ризики порушення авторських прав і вимагати розроблення спеціальних ліцензійних механізмів або виключень для науково-дослідних цілей.

## ***2. Захист персональних даних.***

Мовні моделі можуть ненавмисно опрацьовувати персональні дані, що містяться у відкритих джерелах – наприклад, у соціальних мережах чи публічних базах даних. Це може порушувати Закон України «Про захист персональних даних» [3] та Регламент ЄС GDPR [4], якщо такі дані не були знеособлені. Відтак, необхідним є запровадження механізмів анонімізації, псевдонімізації та контролю доступу до персональних даних.

## ***3. Кібербезпека та державний суверенітет***

Створення національної мовної моделі має стратегічне значення для інформаційного суверенітету держави. Розробники можуть стикнутися з такими викликами :

- *забезпечення захисту від несанкціонованого доступу до моделі та даних;*
- *прозорість джерел даних, щоб уникнути використання ворожих або маніпулятивних інформаційних ресурсів;*
- *дотримання національних стандартів кібербезпеки (відповідно до Закону України «Про основні засади забезпечення кібербезпеки України» [5]), Стратегії кібербезпеки України (2021–2025) [6].*

## ***4. Мовно-культурні та правові колізії.***

Мовна модель, створена в національному контексті, має враховувати мовну політику держави (Закон України «Про забезпечення функціонування української мови як державної» [7]). Водночас вона не повинна порушувати права національних меншин чи інших мовних спільнот. Це створює баланс між мовною ідентичністю та правом на культурне різноманіття, який потребує чіткого нормативного регулювання.

## ***5. Регуляторні та етичні аспекти.***

Постає питання: хто несе юридичну відповідальність за наслідки використання ШІ – розробник, користувач чи держава? Європейська комісія у Регламенті про штучний інтелект (AI Act, 2024) [8] запровадила концепцію "ризик-орієнтованого підходу", відповідно до якого мовні моделі відносяться

до категорії "загального призначення" (GPAI). Для них передбачено обов'язки щодо прозорості, маркування контенту, ведення технічної документації та оцінки ризиків. Україна, у контексті євроінтеграції, повинна гармонізувати національне законодавство з AI Act, що створює нові правові рамки.

### **6. Необхідність правового регулювання ШІ**

Україна наразі не має спеціального закону про штучний інтелект. Однак у процесі гармонізації з європейським законодавством необхідно прийняти національний акт про ШІ, який визначатиме юридичний статус мовних моделей, вимоги до прозорості, сертифікації та етичних принципів, а також створити регуляторний орган з питань ШІ, який здійснюватиме моніторинг і сертифікацію мовних моделей. Необхідно розробити державну політику у сфері відповідального ШІ, що враховуватиме міжнародні стандарти.

Також варто зазначити, що розробкою національної великої мовної моделі займається Міністерство цифрової трансформації України спільно з приватною компанією Київстар. Основною небезпекою такої взаємодії можуть бути ризики конфіденційності, прозорості та потенційного впливу іноземного капіталу на критичну цифрову інфраструктуру. Особливо це актуально, коли приватна компанія має складну структуру власності або зв'язки з країною-агресором. За даними ЗМІ, серед власників Київстару досі є російські олігархи [9]. Тому це викликає занепокоєння щодо національної безпеки, особливо в умовах війни.

Загалом, створення національної мовної моделі LLM вимагає комплексного правового підходу, який поєднує авторське, інформаційне, етичне, мовне та кібербезпекове законодавство. Без належного правового регулювання існує ризик порушення прав людини, втрати даних, дискримінаційних практик і навіть загроз національній безпеці.

### **Література**

1. Українська національна велика мовна модель URL: <https://llm.thedigital.gov.ua/#about> , дата звернення 23.10.2025 р.
2. Українська національна велика мовна модель – шанс для цифрового суверенітету URL: <https://blog.liga.net/user/ssidorenko/article/57995>, дата звернення 23.10.2025 р.

3. Закон України «Про захист персональних даних» від 01.06.2010 № 2297-VI. URL: <https://zakon.rada.gov.ua/laws/show/2297-17#Text>, дата звернення 23.10.2025 р.
4. Regulation (EU) 2016/679 (General Data Protection Regulation, GDPR). URL: <https://www.gdpr.org.ua/>, дата звернення 23.10.2025 р.
5. Закон України «Про основні засади забезпечення кібербезпеки України» від 05.10.2017 № 2163-VIII. URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text>, дата звернення 23.10.2025 р.
6. Указ Президента України № 447/2021 «Про Стратегію кібербезпеки», URL: <https://www.president.gov.ua/documents/4472021-40013>, дата звернення 23.10.2025 р.
7. Закон України «Про забезпечення функціонування української мови як державної» від 25.04.2019 № 2704-VIII. URL: <https://zakon.rada.gov.ua/laws/show/2704-19#Text>, дата звернення 23.10.2025 р.
8. Регламент про штучний інтелект (AI Act, 2024) URL: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>, дата звернення 23.10.2025 р.
9. Мінцифри доручило розробку національної мовної моделі "Київстару". Серед власників оператора досі є російські олігархи Джерело: <https://censor.net/n3558458> URL: <https://censor.net/biz/news/3558458/mintsyfyrydoruchylo-rozrobku-natsionalnoyi-movnoyi-modeli-kyuyivstaru>, дата звернення 28.10.2025 р.

**Василь Орищук**

доктор філософії у галузі публічного  
управління та адміністрування, науковий  
співробітник Київського національного  
університету імені Тараса Шевченка,  
ORCID: <https://orcid.org/0000-0003-4362-2785>

**Максим Валін**

фахівець з освітнього контенту, співпрацює з  
EPAM Systems  
ORCID: <https://orcid.org/009-0002-2966-2887>

## **LLM ТА «ДІЯ.AI» ЯК ОСНОВА ПОБУДОВИ SMART CITY-ЕКОСИСТЕМ**

Україна переживає інтенсивну цифрову трансформацію державних і регіональних сервісів, навіть в умовах військового стану. Цифрові рішення, зокрема застосування технологій штучного інтелекту (ШІ), розглядаються як ключові інструменти для післявоєнного відновлення та розвитку міст [1]. Концепція Smart City передбачає застосування сучасних ІТ-рішень для підвищення якості життя громадян та ефективності управління. Уже сьогодні ці елементи реалізовані через екосистему «Дія», яка забезпечує доступ до десятків цифрових послуг. Наступним етапом стало впровадження великої мовної моделі (LLM) та державного ШІ-асистента «Дія.AI» – першого у світі національного AI-агента, інтегрованого у державний портал для надання послуг громадянам через чат-інтерфейс [2]. Такі інновації мають значний потенціал для розвитку Smart City-екосистем, але водночас створюють нові етичні та правові виклики, що потребують комплексного регулювання й гарантій захисту прав людини.

**Етичні аспекти впровадження ШІ-асистентів у Smart City.** Використання LLM та ШІ-асистентів у міських цифрових сервісах актуалізує питання етики та захисту персональних даних. Оскільки моделі потребують великих обсягів інформації, існує ризик несанкціонованого доступу чи зловживань. У «Дія.AI» реалізовано принцип zero trust: дані автоматично маскуються перед передачею до хмарної моделі, а LLM не має доступу до незашифрованої інформації. Запити користувачів проходять guardrail-фільтри,

що замінюють ідентифікатори на теги, а історія взаємодії шифрується унікальним ключем, мінімізуючи ризики витоку [3].

Другий етичний виклик — безпека та надійність ШІ-асистента. Система має стабільно працювати, витримувати великі навантаження й бути стійкою до кібератак. Розробники визнають, що безпека залишається постійним викликом: інтеграція з державними реєстрами, масштабування та розвиток guardrail-моделей потребують безперервного моніторингу [3]. Під час тестування користувачі намагалися обдурити «Дія.АІ», що стимулювало вдосконалення фільтрів небезпечного контенту. Це підкреслює важливість захисту від неправомірного використання, прозорості алгоритмів, усунення упередженості (bias) та регулярного аудиту моделей. Подібні принципи закріплені в AI Bill of Rights (США, 2022) — праві на захист від алгоритмічної дискримінації, прозорість і безпеку систем [4], а також у рекомендаціях OECD і UNESCO, які варто враховувати під час впровадження LLM-асистентів у державні сервіси.

Третій аспект — довіра та суспільне прийняття ШІ. Щоб інновації Smart City справді покращували життя, громадяни мають довіряти технологіям і розуміти їхню логіку. Взаємодія з асистентом повинна бути зрозумілою й гуманно орієнтованою. Інтерфейс «Дія.АІ» виконано у форматі дружнього чату, подібного до популярних ботів, що полегшує сприйняття користувачами [5]. Водночас держава має гарантувати інклюзивність: сервіси на базі ШІ повинні бути доступними для людей похилого віку, маломобільних осіб і користувачів із низьким рівнем цифрової грамотності. Збереження альтернативних офлайн-каналів є важливою умовою поваги до вибору людини — одного з базових принципів етичного цифрового врядування.

**Правові виклики використання LLM та «Дія.АІ».** Законодавче регулювання ШІ у публічних сервісах в Україні лише формується. Наразі відсутній спеціальний закон про штучний інтелект, а чинні норми, зокрема Закон «Про захист персональних даних», не враховують специфіку машинного навчання. Як підкреслює О. Заярний, у сфері Smart City досі немає єдиних стандартів і нормативної бази: бракує загальнонаціональної концепції, технічних регламентів і кадрової підтримки, тоді як у ЄС уже напрацьовано комплексні підходи до кібербезпеки, відкритих даних і використання ШІ [1].

Одним із ключових питань залишається визначення правового статусу та відповідальності ШІ-асистентів. Закон має чітко закріпити, що

використання ШІ не звільняє органи влади від відповідальності за якість наданих послуг, а громадянин повинен бути поінформований про взаємодію з автоматизованою системою та мати можливість звернутися до фахівця у разі спірної ситуації. Такий підхід відповідає європейському принципу *human-in-the-loop*, який забезпечує людський контроль над рішеннями ШІ.

З огляду на євроінтеграційний курс України, доцільно враховувати досвід ЄС у сфері регулювання ШІ. У 2024 р. Європарламент ухвалив AI Act — перший комплексний закон, що запровадив ризик-орієнтований підхід: системи класифікуються за рівнями ризику, а високоризикові (медицина, транспорт, правосуддя, публічні послуги) підлягають сертифікації, моніторингу та оцінці відповідності, тоді як соціальний скоринг і неконтрольоване стеження — заборонені [4]. Для України це означає необхідність визначити категорії AI-систем у державних сервісах і запровадити вимоги до безпечності, прозорості та людського нагляду.

Перші кроки у цьому напрямі вже зроблено. У червні 2024 р. Міністерство цифрової трансформації презентувало «Білу книгу» з баченням державної політики у сфері ШІ, а 14 IT-компаній підписали Добровільний кодекс етичного використання [4]. Такий *soft law*-підхід формує основу майбутнього Закону України «Про штучний інтелект», який має спиратися на EU AI Act і рекомендації Ради Європи. Водночас потребує оновлення суміжне законодавство — передусім Закон «Про захист персональних даних» (з урахуванням GDPR) та стандарти кіберзахисту Smart City. Як підкреслює О. Заярний, слід запровадити сертифікацію та стандартизацію AI-рішень за принципами *secure-by-design* і *privacy-by-design*, щоб гарантувати кібербезпеку та захист прав людини [1].

Показовим прикладом практичного застосування цих підходів є «Дія.AI» — перший у світі національний AI-агент, запущений наприкінці 2023 р. на порталі «Дія» [2]. Асистент здатен виконувати багатокрокові процедури, зокрема автоматично формувати довідки за запитом користувача, і вже до середини 2025 р. ним скористалися понад 27 тис. громадян [3]. Водночас інтеграція ШІ в мобільний застосунок для 19 млн користувачів виявила технічні та організаційні виклики, зокрема потребу

масштабування інфраструктури й оптимізації взаємодії з державними реєстрами.

Досвід «Дія.АІ» має стратегічне значення для розвитку Smart City в Україні. Єдина платформа для державних і муніципальних послуг створює основу «єдиного вікна» взаємодії держави та громад, а дані користувацьких звернень можуть стати аналітичним інструментом для управління містом. Такі LLM-асистенти поєднують функції надання послуг, аналітики та зворотного зв'язку, підвищуючи проактивність і адаптивність міського врядування. Подібні рішення вже впроваджуються в містах ЄС, тоді як в Україні громади тестують чат-боти ЦНАПів і системи відеоаналітики на базі ШІ. Важливо, щоб правове поле розвивалося синхронно з технологіями, формуючи єдині стандарти взаємодії держави, бізнесу й громад.

**Соціальні наслідки та вплив на права громадян.** Впровадження LLM та «Дія.АІ» має багатовимірний соціальний ефект, який необхідно оцінювати крізь призму захисту прав людини. Найпомітніший позитив — це підвищення доступності та швидкості отримання послуг. Для мешканців віддалених територій чи осіб з інвалідністю АІ-асистенти стають «вирівнювачем» можливостей, забезпечуючи дистанційний доступ без черг і бюрократії. Це сприяє реалізації права на інформацію та адміністративні послуги, а також зменшує корупційні ризики завдяки мінімізації контакту з посадовцями.

Водночас існують соціальні ризики, зокрема посилення *цифрової нерівності*: люди похилого віку чи користувачі з низьким рівнем цифрової грамотності можуть опинитися поза цифровими сервісами. Тому держава має інвестувати в розвиток цифрових навичок, проводити інформаційні кампанії щодо «Дія.АІ» та забезпечувати альтернативні, офлайн-способи отримання послуг. Це відповідає праву на рівність можливостей і недискримінацію.

Окремої уваги потребує захист приватності. Хоча технічні рішення «Дія.АІ» забезпечують анонімізацію даних користувачів [3], необхідні й законодавчі гарантії, що персональна інформація не використовуватиметься без згоди особи. Із розширенням інтеграції ШІ — зокрема у відеоаналітику чи міські бази даних — постає питання дотримання права на приватне життя у публічному просторі. Уже сьогодні варто розробляти місцеві політики ethical

AI, щоб запобігти ризику перетворення розумних міст на системи цифрового контролю. Важливим етичним принципом є прозорість і пояснюваність ШІ: громадяни мають розуміти, як працює система, її обмеження та можливі похибки. У «Дія.AI» це реалізовано через зворотний зв'язок: користувач може оцінити відповідь ШІ, залишити коментар, а команда розробників аналізує відгуки й удосконалює модель [5]. Такий механізм формує основу державної системи контролю якості ШІ та передбачає можливість створення незалежного омбудсмена з цифрових прав для захисту користувачів від зловживань автоматизованих систем.

Не менш важливими є культурні та психологічні аспекти взаємодії людини й штучного інтелекту. Масове використання AI-асистентів змінює очікування громадян щодо швидкості та проактивності держави, але може знизити рівень живого спілкування й емпатії. Тому необхідно зберігати баланс між автоматизацією та людяністю сервісів. Визначальним має залишатися людиноцентричний підхід, закріплений у концепції Society 5.0 (Японія), де ШІ розширює можливості людини, а не замінює її [4]. Для України це означає, що технології LLM і «Дія.AI» мають підвищувати комфорт і безпеку громадян, не порушуючи їхніх основоположних прав і свобод.

**Висновки.** Впровадження LLM та державного ШІ-асистента «Дія.AI» стало важливим кроком у формуванні сучасної Smart City-екосистеми України. Ці технології підвищують доступність, швидкість і зручність публічних послуг, роблячи державу ближчою до громадян. Запуск першої національної AI-системи у публічному секторі перетворює Україну на одного з лідерів цифрового врядування, що особливо актуально в умовах післявоєнного відновлення. Досвід «Дія.AI» уже привернув міжнародну увагу та може стати основою для масштабування рішень на рівні громад.

Водночас аналіз виявляє низку етичних і правових ризиків — від загроз приватності та кібербезпеки до нерівного доступу та відсутності комплексного правового регулювання.

Для безпечного й сталого розвитку Smart City на основі LLM і «Дія.AI» необхідні такі кроки:

- *Розробка нормативно-правової бази у сфері ШІ — із чіткими стандартами, процедурами сертифікації та вимогами до безпечності й етичності AI-систем [1].*

– *Імплементація європейських практик і стандартів*, включно з рекомендаціями для місцевого самоврядування щодо розвитку Smart City під час війни та у відновний період [1].

– *Створення ефективних механізмів контролю та відповідальності*, включно з незалежним моніторингом і аудитом алгоритмів.

– *Посилення етичних засад і саморегулювання* — інтеграція модулів перевірки на упередженість, фільтрації контенту, навчання держслужбовців і розробників з етики ШІ.

– *Прозорість і підзвітність* — регулярні звіти про роботу ШІ-систем, відкриті дані щодо звернень і результатів для підвищення довіри суспільства.

– *Інвестиції в кібербезпеку та інфраструктуру* — захист цифрових мереж, дата-центрів і хмарних сервісів, що підтримують LLM-рішення.

Підсумовуючи, LLM та «Дія.АІ» відкривають нову сторінку в розвитку публічних послуг і цифрового врядування. Їхній потенціал може бути реалізований лише за умови дотримання принципів прозорості, безпеки, етичності та людиноцентричності. Інтегруючи найкращі європейські стандарти й власний досвід, Україна має шанс стати флагманом у сфері правового регулювання ШІ, забезпечивши баланс між інноваціями, правами людини та довірою суспільства.

### **Список використаних джерел:**

1. Zaiarnyi O. Legislative support for the implementation of the "Smart City" concept in the European Union and proposals for Ukraine during wartime and post-war recovery. *Bulletin of Taras Shevchenko National University of Kyiv. Legal Studies*. 2024. No. 128. P. 28–32. URL: <https://doi.org/10.17721/1728-2195/2024/2.128-5>.

2. President of Ukraine (Press Office). The Ministry of Digital Transformation Team Presented New Digital Service Projects to the President, Including the World's First AI Assistant Providing Government Services. Official website of the President of Ukraine, News, 1 September 2025. URL: <https://www.president.gov.ua/en/news/komanda-mincifri-prezentovala-prezidentu-novi-proyekti-u-sfe-99885>.

3. Мамченко Н. Безпека залишається постійним викликом – у Мінцифри пояснили, як працюють з персональними даними при створенні національної LLM-моделі та Дія.АІ. *Судово-юридична газета*. 23 вересня 2025 р. URL: <https://sud.ua/uk/news/publication/341776-bezopasnost-ostaetsya-postoyannym-vyzovom-v-mintsifry-obyasnili-kak-rabotayut-s-personalnymi-dannymi-pri-sozdanii-natsionalnoy-llm-modeli-i-diyai>.

4. Kazantsev D. Ethical implementation of AI: How are Ukraine going to regulate the new industry? Mezha (IT news portal). 17 September 2024. URL: <https://mezha.ua/en/articles/yak-v-ukrajini-planuyut-regulyuvati-shi-galuz-304743/>.

5. Fedorov M. Diia.AI: The World's First National AI-Agent That Delivers Real Government Services. *Digital State UA: Ukrainian Tech for Future Societies*. URL: <https://digitalstate.gov.ua/news/govtech/diiaai-pershyy-u-sviti-derzavnyy-ai-ahent-iakyy-ne-prosto-konsultuye-a-nadaye-posluhy-iak-pratsiuye-shtuchnyy-intelekt-na-portali>.

**Николина К.В.**

кандидат юридичних наук, доцент кафедри  
теорії та історії права та держави Навчально-  
наукового інституту права Київського  
національного університету імені Тараса  
Шевченка

ORCID: <https://orcid.org/0000-0002-8603-1013>

## **НАЦІОНАЛЬНА ВММ ЯК ОСНОВА ЦИФРОВОГО СУВЕРЕНІТЕТУ УКРАЇНИ**

Глибока трансформація суспільних відносин, зумовлена стрімким розвитком цифрових технологій, ставить перед загальною теорією права фундаментальне завдання переосмислення ключових державно-правових категорій. Однією з таких категорій, що зазнає найбільш відчутних змін, є державний суверенітет. Класичне вестфальське розуміння суверенітету як верховенства державної влади в межах власної території та її незалежності у зовнішніх відносинах виявляється недостатнім для адекватного опису реалій глобалізованого цифрового світу.

Поява та швидкий розвиток великих мовних моделей (ВММ) кардинально змінює інституційно-правові механізми втілення цифрового суверенітету. Експерти визначають *великі мовні моделі* як окремий клас мовних моделей, які використовують алгоритми глибокого навчання та навчаються на великих наборах даних, котрі можуть містити не тільки текст, але й інші модальності (зображення, аудіо тощо) [1, с.7]. Аналіз наукової літератури свідчить, що вони мають деякі загальні характеристики: генеративні за своєю природою, використовують самостійне навчання та адаптуються до різних завдань[2, с.55].

В сучасних умовах великі мовні моделі перестають бути суто технологічним інструментом і перетворюються на стратегічну інфраструктуру, що формує інформаційний, культурний та семантичний простір нації. Контроль над цією інфраструктурою стає новим, невід'ємним виміром державної влади.

Для України, яка веде екзистенційну боротьбу за свій державний суверенітет в умовах повномасштабного вторгнення та водночас реалізує

стратегічний курс на європейську інтеграцію, питання цифрового суверенітету набуває особливої гостроти. Ініціатива створення національної великої мовної моделі від Міністерства цифрової трансформації є прямою відповіддю на цей виклик. Вона розглядається не лише як технологічний проєкт, а як інструмент забезпечення інформаційної безпеки, збереження культурно-мовної ідентичності та стимулювання інноваційної економіки [3]. Більше того, 25 жовтня 2025 року команда українських дослідників із Факультету прикладних наук УКУ, AGH University of Krakow, КПІ ім. Ігоря Сікорського та Львівської Політехніки вже представила Lora LLM — передову відкриту велику мовну модель з фокусом на обробку української мови.

У відповідь на глобальні виклики технологічного прориву в сучасній правовій науці формується концепція "цифрового суверенітету" як ключового інструменту утвердження влади правової держави у віртуальному середовищі. Цілком обґрунтованою є думка, що суть цифрового суверенітету полягає в необхідності контролю над цифровими технологіями на фізичному рівні (ресурси, інфраструктура, пристрої), на рівні коду (стандарти, правила, дизайн) та на інформаційному рівні (контент, дані) [4, с. 157].

З позиції загальної теорії права, фундаментальні моделі, і зокрема ВММ, слід розглядати не просто як об'єкт права інтелектуальної власності. За своєю функціональною роллю, на нашу думку, вони являють собою стратегічну інфраструктуру загальнонаціонального значення, оскільки не просто передають інформацію, а й генерують, інтерпретують та контекстуалізують її, впливаючи на формування суспільної думки, поширення знань та збереження культурних кодів.

Усвідомлення інфраструктурної ролі великих мовних моделей дозволяє обґрунтувати тезу про те, що володіння власною, національною мовною моделлю стає необхідним атрибутом сучасної суверенної держави. Залежність від іноземних ВММ, контрольованих кількома глобальними корпораціями, створює прямі загрози для реалізації функцій держави.

Проєкт створення української ВММ, ініційований державою у партнерстві з бізнесом, є практичним кроком до реалізації цифрового суверенітету. У науковому дискурсі поки не існує єдиного, уніфікованого визначення цифрового суверенітету, однак більшість підходів сходяться на тому, що він визначається як здатність держави створити автономну цифрову інфраструктуру та самостійно здійснювати керування нею, управляти

інформаційними потоками, обмежувати їх і претендувати на ексклюзивне право визначати національний дискурс та інформаційне поле, а також забезпечувати можливість контролювати і встановлювати правові межі транснаціональним технологічним компаніям, через які перерозподіляються інформаційні потоки[5, с. 27].

Цілком погоджуємося, що його стратегічне значення виходить за межі суто технологічних переваг і охоплює ключові функції держави. Насамперед, це *забезпечення національної безпеки та оборони*. Національна ВММ дозволяє аналізувати розвіддані, обробляти таємну інформацію, використовуючи захищені саме локальні сервери, що є критично важливим в умовах війни.

Глобальні моделі, навчені переважно на іноземних даних, не лише погано розуміють український контекст, а й можуть відтворювати ворожі наративи. Саме тому національна ВММ, навчена на ретельно відібраних українських джерелах, здатна також виконувати *функцію захисту культурного та інформаційного простору*, даючи коректні відповіді на чутливі питання історії, культури та війни.

Окрім цього, реалізація *економічної функції держави* проявляється в тому, що розробка власної ВММ стимулює інновації, створює нові можливості для бізнесу, підвищує ВВП та зменшує технологічну залежність від закордонних постачальників. Як зазначають самі розробники, оптимізація вартості обробки україномовних запитів робить AI-сервіси більш доступними для національної економіки.

Водночас реалізація цього амбітного проєкту виводить на поверхню низку складних теоретико-правових проблем, що потребують глибокого наукового осмислення та виважених законодавчих рішень. Зокрема, для створення потужної моделі необхідні величезні масиви текстів, значна частина яких захищена авторським правом. Це створює фундаментальну правову колізію. З одного боку, Конституція України та цивільне законодавство гарантують непорушність права інтелектуальної власності. З іншого боку, суспільний інтерес, що полягає у забезпеченні національної безпеки та цифрового суверенітету, вимагає використання цих даних для навчання національної моделі.

Ще одним складним теоретичним питанням є визначення суб'єкта відповідальності за шкоду, завдану внаслідок функціонування LLM. Хто нестиме відповідальність, якщо модель згенерує дифамаційний текст,

небезпечну інструкцію або поширить дезінформацію? Розробник фундаментальної моделі? Розробник кінцевого застосунку? Чи кінцевий користувач? Очевидним стає факт, що традиційні підходи до деліктної відповідальності, що вимагають доведення вини та прямого причинно-наслідкового зв'язку, виявляються неефективними в умовах складних алгоритмів функціонування штучного інтелекту.

Підсумовуючи зазначене, можемо констатувати, що такі системи, як GPT-5 від OpenAI, Llama від Meta та Gemini від Google, є не просто черговим етапом еволюції машинного навчання; вони являють собою нову парадигму розробки систем штучного інтелекту, що характеризується безпрецедентним масштабом, універсальністю та здатністю до розуміння та генерування текстів. Ця трансформація виводить технологічну конкуренцію на новий геополітичний рівень, де контроль над розробкою, навчанням та розгортанням великих мовних перетворюється на ключовий атрибут національної могутності, економічної стійкості та, що найважливіше, цифрового суверенітету.

На наше глибоке переконання, успішна реалізація проєкту національної ВВМ та утвердження цифрового суверенітету України залежать не стільки від технологічних досягнень, скільки від формування адекватної, продуманої та збалансованої правової рамки на основі синтезу найкращих світових практик та врахування українського контексту (воєнний стан, необхідність захисту мовної ідентичності, усунення русифікованих та упереджених джерел в навчанні ВВМ тощо).

### **Список використаних джерел:**

1. Словник термінів у сфері штучного інтелекту / упорядники: Чумаченко Д., Мішкін Д., Андрієнко О., Краковецький О., Турута О., Дубно О., Хрушова Д., Кобрін А., Авдєєва Т., Кравець І., Герасимяк В., Шабанов О., Бистрицька А. Київ: Міністерство цифрової трансформації України, 2024. 37с.
2. Юрчак, І., Хіч, А., & Оксентюк, В. Розуміння великих мовних моделей: майбутнє штучного інтелекту. *Computer design systems. Theory and practice*. № 6(2). 2024. 51-60.
3. Україна намагається розробити національну мовну модель – Мінцифри та Київстар підписали меморандум про співпрацю. Кабінет міністрів України. Офіційний вебсайт. URL: <https://www.kmu.gov.ua/news/ukraine-pochynaie-rozrobku->

natsionalnoi-movnoi-modeli-mintsyfry-ta-kyivstar-pidpysaly-memorandum-pro-spiivpratsiu (дата звернення: 30.10.2025).

4. Милосердна І. М. Цифровий суверенітет держави: наукова риторика та реальні зміни. Політикус: наук. журнал. 2024. № 6. С.154-160.

5. Єфремова К.В. Забезпечення цифрового суверенітету держави в умовах індустрії 4.0. Конституційні засади розвитку інноваційного суспільства: зб. наук. пр. за матеріалами інтернет-конференції (м. Харків, 25 червня 2021 року) / за ред. А. В. Стріжкової. Харків: НДІ ПЗІР НАПрН України, 2021. С. 25–30.

URL: <https://openarchive.nure.ua/server/api/core/bitstreams/f981db39-23a5-470f-bf25-e1af2f3701dc/content>

**Ярослав Мануїлов**

старший науковий співробітник Українського  
науково-дослідного інституту спеціальної  
техніки та судових експертиз СБУ

ORCID: <https://orcid.org/0000-0001-8149-2745>

## **СТРАТЕГІЯ РОЗВИТКУ МІЖНАРОДНОГО КІБЕРПРОСТОРУ ТА ЦИФРОВОЇ ПОЛІТИКИ США**

США послідовно обстоює доктрину **цифрової солідарності** (*digital solidarity*), основною ідеєю якої є створення «відкритої, стійкої та безпечної» глобальної цифрової екосистеми, в якій важливе місце займає цифровий захист персональних даних дітей.

6 травня 2024 р. була представлена нова Стратегія розвитку міжнародного кіберпростору та цифрової політики (далі – Стратегія), яку підготувала Бюро кіберпростору та цифрової політики Державного департаменту США (US Department of State’s Bureau of Cyberspace and Digital Policy) у співпраці з іншими федеральними відомствами та експертами у сфері кіберзахисту.

Зміст цієї Стратегії узгоджується з раніше прийнятими в США Стратегією національної безпеки (2022) та Стратегією кібербезпеки (2023), а також відповідає іншим програмним правовим актами глави держави в галузі ІТ та кіберпростору.

Преамбула Стратегії, яка має назву «Назустріч інноваційному, безпечному цифровому майбутньому, з повагою до прав» (Towards an Innovative, Secure, and Rights-Respecting Digital Future), відтворює інноваційний та новаторський підхід, зміст якого спрямований на формування образу майбутнього захищеного кіберпростору.

Запропонована в Стратегії модель майбутнього кіберпростору передбачає чотири напрями діяльності (сфери дії):

1. Розбудова й підтримка відкритої, інклюзивної, безпечної та стійкої міжнародної цифрової екосистеми. Це, зокрема, передбачає стимулювання конкурентного ринку ІТ та інновацій, зменшення залежності від авторитарних режимів, загальнодоступну, надійну та безпечну цифрову

інфраструктуру, належний захист телекомунікацій та цифровий захист персональних даних громадян, у т.ч. дітей.

2. Узгодження з міжнародними партнерами підходів, стандартів та політик цифрового управління й управління даними на засадах поваги до права, безпеки та підтримки міжнародної торгівлі.

3. Просування відповідальної поведінки держав у кіберпросторі, протидія кіберзагрозам шляхом створення коаліцій, сприяння партнерській співпраці, обміну інформацією та підзвітності між країнами.

4. Зміцнення та розбудова цифрового та кіберпотенціалу країн-партнерів, зокрема підвищення їхніх технічних і операційних можливостей та спільну боротьбу з кіберзлочинністю [1].

Стратегія містить низку ініціатив Державного департаменту США для колективного обговорення та ухвалення рішень, реалізацію яких визнають особливо актуальною. Основними з них є такі: розроблення рамок управління штучним інтелектом – установлення міжнародних стандартів етичного використання штучного інтелекту із залученням партнерів через G7 та інші форуми для створення керівних принципів розвитку технологій ШІ; підвищення безпеки ланцюга постачання – передбачає співпрацю з партнерами для диверсифікації та безпеки ланцюга постачання для критично важливих технологій; США готові інвестувати в рішення, які зменшують залежність від авторитарних режимів; просування кібернорм – спільне узгодження та захист глобальних норм, які визначають прийнятну поведінку держави в кіберпросторі та притягують порушників до відповідальності; розробка та впровадження механізмів захисту, які унеможливили б усі ризики, з якими нині зіткнулися соцмережі. Пропоновано майданчик ООН та інших міжнародних організацій [1].

Запропонована США модель розвитку міжнародного кіберпростору просуває ідею цифрового суверенітету, популярність якого зростає у світі протягом останніх років. Стратегія передбачає внесення фундаментальних змін в основну динаміку цифрового суверенітету, надавши перевагу захисту національного кіберпростору.

Водночас, з цього приводу висловлюються різні точки зору.

Одна з них зводиться до того, що така ідея у майбутньому може призвести до: надмірної фрагментації міжнародного управління інтернетом; ізоляції цифрових екосистем; додаткових бар'єрів для міжнародної співпраці,

, поширення інновацій і зрештою – для просування технологічного прогресу як такого.

Прихильники іншої точки зору вважають, що орієнтація національних урядів на цифровий суверенітет значно збільшує ризики кібербезпеки, ускладнює міжнародну співпрацю, потрібну для боротьби з глобальними кіберзагрозами – обмін інформацією та розвідувальними даними між правоохоронними органами іноземних держав, координацію реагування на кіберінциденти, спільне з міжнародними партнерами вироблення підходів до забезпечення кіберзахисту тощо.

Важливою складовою цифрової політики є інформаційний захист дітей, зміст якого охоплює різноманітні рекомендації: від обмеження доступу дітей до за віком і часом користування, створення окремих цифрових продуктів для дітей, до повної відмови від смартфонів (до з'ясування деструктивних наслідків використання соцмереж). У 2023 році в штаті Юта було підписано закон, згідно з яким використання соцмереж підлітками можливе лише після згоди батьків чи опікунів. Одночасно із цим підліткам заборонялося використовувати соцмережі вночі — після 22:30. У 2024 році законодавці штату Нью-Йорк представили законопроект, який заборонятиме алгоритмічні стрічки соцмереж[2]. Ці приклади свідчать про спроби регулювання цифрового простору і соцмереж для неповнолітніх користувачів.

Слід також зазначити, що політика нового президента США Д. Трампа змістила акценти у забезпеченні кіберзахисту об'єктів критичної інфраструктури. Так, у лютому 2025 року Д. Трамп заморозив і припинив фінансування програм USAID, включно з "Кібербезпекою критичної інфраструктури в Україні", а його адміністрація акцентує захист критичної інфраструктури від кібератак великих державних акторів – Китаю, Росії, Ірану. Як зазначають експерти, політика Д.Трампа у кібербезпеці поєднує проактивність і дерегуляцію, а її успіх залежатиме від балансу між скороченням витрат і реальними заходами проти кіберзагроз [3].

Між тим нова політика Д.Трампа не змінює проголошені Стратегією базові три керівні принципи: 1) чітка й послідовна зорієнтованість (affirmative vision) на безпечність та інклюзивність кіберпростору, керованого міжнародним правом; 2) цілісне бачення кібербезпеки; 3) комплексний підхід у формуванні політичних дій з використанням інструментарію дипломатії, державного управління та міждержавних відносин [1].

В будь-якому разі міжнародна цифрова стратегія США зумовлює потребу активних зусиль, спрямованих на нейтралізацію загроз у кіберпросторі, та міжнародну співпрацю у цій сфері.

Важливою складовою «відкритої, стійкої та безпечної» глобальної цифрової екосистеми залишається інформаційний захист персональних даних дітей.

### **Список використаних джерел:**

1. Гнатюк С. Геополітичне цифрове візіонерство США та ЄС: цифрова солідарність, цифровий суверенітет і нова американська стратегія розвитку міжнародного кіберпростору. URL: [https://niss.gov.ua/sites/default/files/2025-01/az\\_cifrove\\_vizionerstvo\\_06012025.pdf](https://niss.gov.ua/sites/default/files/2025-01/az_cifrove_vizionerstvo_06012025.pdf).

2. Цифрові кордони для неповнолітніх. Як світ намагається обмежити інтернет для дітей. URL: <https://tyzhden.ua/tsyfrovi-kordony-dlia-nepovnolitnikh-iak-svit-namahaietsia-obmezhyty-internet-dlia-ditej/>

3. Політика Дональда Трампа у сфері кібербезпеки: від першого терміну до 2025. URL: <https://cybersec.net.ua/statti/824-polityka-donalda-trampa-u-sferi-kiberbezpeky-vid-pershoho-terminu-do-2025-roku.html>.

## **ПРАВА ЛЮДИНИ, ПРИВАТНІСТЬ ТА ПИТАННЯ ІНТЕЛЕКТУАЛЬНОЇ ВЛАСНОСТІ**

**Пашинський В.Й.**

доктор юридичних наук, доцент, доцент  
кафедри адміністративного права та процесу  
Навчально-наукового інституту права  
Київського національного університету імені  
Тараса Шевченка

ORCID: <https://orcid.org/0000-0001-7349-2227>

**Цьоменко А.В.**

PhD, асистент кафедри адміністративного  
права та процесу Навчально-наукового  
інституту права Київського національного  
університету імені Тараса Шевченка

ORCID: <https://orcid.org/0000-0002-5615-7838>

## **АДМІНІСТРАТИВНО-ПРАВОВЕ ЗАБЕЗПЕЧЕННЯ ЗАХИСТУ ПЕРСОНАЛЬНИХ ДАНИХ В УМОВАХ РОЗВИТКУ ШТУЧНОГО ІНТЕЛЕКТУ**

Штучний інтелект в епоху цифровізації людства широко впроваджується всюди: від побуду, користувацького досвіду в мережі Інтернет й електронної комерції до забезпечення правопорядку, правосуддя, оборонної сфери та навіть дипломатії. Новації створюють нові можливості та виклики для безпеки в кіберпросторі, одні з яких можуть спровокувати дезінформацію, підрив довіри до уряду з боку суспільства та перетворення людей на «прозорих громадян», а великі обсяги їхніх персональних даних підштовхнути до незаконного використання, де інформація у відкритому потоці економіки даних використовується без її дозволу та, більше того, проти неї. Система ШІ вчиться на даних, використовує дані та заснована на даних. Це алгоритм, і дані різного характеру є основою його функціонування. Саме тому завдяки доступності великих наборів даних з інтернету система досягла суттєвого прогресу в навчанні [1, с. 174].

Однак нові дослідження показують, що зростає так звана «криза згоди» у веб-доменах, які вжили заходів для запобігання збирання штучним інтелектом даних, які ті зберігають, і натомість можуть вимагати плату за доступ до них [2].

Згідно з Постановою КМУ від 30 квітня 2024 № 476 «Про затвердження переліку пріоритетних тематичних напрямів наукових досліджень і науково-технічних розробок на період до 31 грудня року, наступного після припинення або скасування воєнного стану в Україні року» системи штучного інтелекту, глибоке навчання, великі дані (big data) та нейроподібні мережі є актуальними до опрацювання зокрема у правовій доктрині, оскільки штучний інтелект в епоху цифровізації людства створює нові виклики для безпеки в кіберпросторі, одні з яких створюють загрозу витоків великих наборів персональних даних, зловживання ними та несанкціонованого доступу до них.

У доповіді Уповноваженого за 2023 рік резюмується, що законодавство про захист персональних даних потребує оновлення у зв'язку зі стрімким технологічним розвитком і використанням технологій штучного інтелекту. Застосування цих сучасних технологій у різних сферах створює необхідність у чіткому правовому регулюванні для забезпечення захисту прав і свобод людини, що в сукупності робить тему актуальною для досліджень у юридичній науці.

Урядом України поставлена мета створити цифрову державу, де штучний інтелект – важлива частина цього завдання. Навесні 2024 року Міністерство закордонних справ України представило першу у світі цифрову дипломатку, створену нейромережею на основі реальної людини [3]. Існує чітка вірогідність ризиків, пов'язаних з цифровими аватарами, наприклад, широке поняття кризи згоди, яке теж відноситься до необхідності надання однозначно закріпленої згоди людини на обробку її зображення (біометричних даних), а також можливість створення цифрової підробки аватара та поширення дезінформації серед суспільства.

В останні роки важливість адміністративно-правового забезпечення захисту персональних даних стає особливо актуальною як у країнах ЄС, так і в Україні. Це відбувається на тлі широкої дискусії щодо впровадження в різні сфери життя суспільства інформаційних технологій та програмних комплексів, які об'єднуються терміном штучний інтелект [4]. У процесі їх

застосування, як показує сучасна практика, виникають проблеми з надійним захистом персональних даних користувачів.

З метою правового врегулювання проблем використання штучного інтелекту при обробці персональних даних громадян та з метою їх забезпечення парламентом ЄС 13 березня 2024 року було прийнято перший нормативно-правовий акт «Artificial Intelligence Act» («Закон про штучний інтелект») [5], який врегулював питання використання штучного інтелекту та захисту персональних даних при його використанні. Заявлена мета цього нормативно-правового акту полягає в захисті фундаментальних прав, демократії, принципу верховенства закону та екологічної стійкості від потенційно небезпечного використання штучного інтелекту. Також його метою є сприяння інноваціям та зміцнення позицій Європи як провідного гравця в цій галузі [6]. Зазначений нормативно-правовий акт знаменує собою ключовий момент у глобальному регулюванні використання штучного інтелекту, встановлюючи прецедент для відповідальної та етичної розробки і його практичного впровадження. Прийняття зазначеного акту знаменує собою початок нового етапу нормативно-правового регулювання, як питань використання штучного інтелекту, так і питань захисту персональних даних при його використанні.

«Artificial Intelligence Act» передбачає: заходи для розробки і використання технології штучного інтелекту враховуючи пріоритет прав людини, безпеки та прозорості; вимоги щодо надійних гарантій захисту від алгоритмічних упереджень і дискримінації, суворі вимоги до конфіденційності та безпеки даних, а також механізми підзвітності та контролю.

Згідно з новими правилами особлива увага приділяється використанню штучного інтелекту у сферах, які вважаються високоризиковими, щодо забезпечення прав людини. Наприклад ті, що використовуються в охороні здоров'я, на транспорті та в правоохоронних органах, підлягатимуть суворій регулятивній перевірці та вимогам до тестування. Нові принципи містять заборони застосування штучного інтелекту, які можуть порушувати права громадян, включаючи системи біометричної ідентифікації на основі чутливих ознак, а також непризначене збирання зображень осіб з Інтернету або відеозаписів з камер відеоспостереження для створення баз даних для розпізнавання осіб [5].

Згідно з новим законом, використання систем біометричної ідентифікації правоохоронними органами «в принципі заборонено, за винятком конкретно визначених та обмежених випадків». Використання таких систем у реальному часі можливе лише за умови суворого дотримання заходів безпеки і «потребує спеціального попереднього судового або адміністративного дозволу». Це може стосуватися наприклад ситуацій пошуку зниклої людини або запобігання терористичним загрозам [6].

Яскравим прикладом необхідності створення організаційних та правових рамок для захисту прав людини в сфері обігу персональних даних є швидке поширення штучного інтелекту, наприклад, чат-бота на основі штучного інтелекту ChatGPT, розробленим компанією OpenAI. Останніми роками штучний інтелект набув широкого використання, що в свою чергу породжує низку правових питань, в тому числі щодо обробки штучним інтелектом персональних даних, які містяться в вільному доступі в світовій мережі. Італія стала першою країною, яка ввела обмеження на використання штучного інтелекту через проблеми, пов'язані, в тому числі, з обробкою персональних даних користувачів [7].

Причиною цього рішення була втрата персональних даних та платіжної інформації користувачів чат-бота 20 березня 2023 року. За думкою Національного управління із захисту персональних даних Італії, масштабна діяльність зі збору персональних даних, проведена розробниками чат-бота на основі штучного інтелекту ChatGPT, не може бути виправдана режимом тестування алгоритмів, що лежать в його основі, через відсутність нормативно-правових засад. Національне управління із захисту персональних даних Італії вказує на те, що недостатність будь-якого фільтра для перевірки віку користувачів ставить неповнолітніх під ризик неприпустимої поведінки відносно їхнього розвитку і самосвідомості. OpenAI були зобов'язані повідомити про вжиті заходи з усунення виявлених технічних проблем [8].

У квітні 2023 року Міністри цифрових технологій країн «Великої сімки» (G7) зреагували на проблему правового регулювання використання штучного інтелекту. Вони узгодили необхідність прийняття відповідних законів для цього. Зазначено, що таке правове регулювання має зберігати відкрите та сприятливе середовище для розвитку технологій штучного інтелекту. Крім того, воно повинно ґрунтуватися на демократичних цінностях, враховуючи занепокоєння щодо конфіденційності та ризиків для безпеки [9].

Наразі забезпечення захисту персональних даних фізичних осіб у контексті розвитку штучного інтелекту стає важливим завданням для органів публічного управління і повинно бути врегульоване законодавством. Потреба у створенні ефективних механізмів захисту особистих даних обумовлена також широким використанням інноваційних проектів, таких як «цифрове урядування» та «смарт-міста», що ґрунтуються на створенні обширних баз даних користувачів. Необхідно регулювати питання стосовно використання та обробки персональних даних у процесі «навчання» штучного інтелекту [10].

Враховуючи, що нині Україна є одним із лідерів щодо запровадження інформаційних технологій в діяльність органів публічного управління, а також створення системи електронного урядування, та здійснює широкий доступ до застосування, в тому числі програм штучного інтелекту. Проте, із зростанням використання інформаційно-телекомунікаційних технологій та технологій штучного інтелекту виникає потреба у відповідних нормативно-правових механізмах, які гарантуватимуть ефективне забезпечення захисту прав та свобод громадян, включаючи захист персональних даних відповідно до стандартів ЄС.

Зазначені вище приклади з світової практики доводять необхідність вже зараз розпочати розробку комплексу правових та організаційно-управлінських заходів забезпечення захисту персональних даних громадян у контексті використання штучного інтелекту. Аналіз досвіду ЄС дає підстави зазначити, що основними завданнями щодо подальшого вдосконалення правового регулювання захисту персональних даних є здійснення заходів за такими напрямками: а) приведення національного законодавства щодо захисту персональних даних до стандартів ЄС, як невід’ємна складова європейської інтеграції України; б) визначення на законодавчому рівні системи органів публічного контролю за забезпеченням захисту персональних даних, яке відповідає законодавству ЄС.

### **Список використаних джерел:**

1. Юридичні інновації та стартапи: підручник. Київ: за заг. ред. Пряміцина В., 2024. 290 с.
2. Roose K. The Data That Powers A.I. Is Disappearing Fast. The New York Times. 2024. URL: <https://www.nytimes.com/2024/07/19/technology/ai-data-restrictions.htm>.

3. МЗС України призначило цифрову особу для інформування щодо консульських питань. Міністерство закордонних справ України. 2024. URL: <https://mfa.gov.ua/news/mzs-ukrayinipriznachilo-cifrovu-osobu-dlya-informuvannya-shchodo-konsulskih-pitan>
4. Kelvin Chan. Europe reaches a deal on the world's first comprehensive AI rules. *Apnews*. URL: <https://apnews.com/article/ai-act-europe-regulation-59466a4d8fd3597b04542ef25831322c>
5. Artificial Intelligence Act: MEPs adopt landmark law. URL: [https://www.europarl.europa.eu/news/en/press-room/20240308IPR19015/artificial-intelligence-act-meps-adopt-landmark-la.M5rp8rD6VJMckEwOVeo\\_](https://www.europarl.europa.eu/news/en/press-room/20240308IPR19015/artificial-intelligence-act-meps-adopt-landmark-la.M5rp8rD6VJMckEwOVeo_)
6. Європарламент схвалив закон про штучний інтелект. *Збруч*. URL: [https://zbruc.eu/node/117949\\_](https://zbruc.eu/node/117949_)
7. Італія першою із західних країн заблокувала ChatGPT. *Українська правда*. URL: <https://www.pravda.com.ua/news/2023/03/31/7395927/&>
8. Intelligenza artificiale: il Garante blocca ChatGPT. Raccolta illecita di dati personali. Assenza di sistemi per la verifica dell'età dei minori. *Garante Per La protezione dei dati personali*. URL: [https://www.gpdp.it/web/guest/home/docweb/-/docweb-display/docweb/9870847\\_](https://www.gpdp.it/web/guest/home/docweb/-/docweb-display/docweb/9870847_)
9. Пилипів І. Країни G7 виступають за прийняття законів для регулювання ШІ. *Українська правда*. URL: [https://www.epravda.com.ua/news/2023/04/30/699615/\\_](https://www.epravda.com.ua/news/2023/04/30/699615/_)
10. Пунда О.О., Арзянцева Д.А. Забезпечення захисту персональних даних фізичних осіб в умовах розвитку штучного інтелекту. *Наука і техніка перспективи. Серія «Право, економіка, педагогіка, техніка, фізико-математичні науки»*. 2024. № 2(30). С.132–143.

**Жанна Лупак**

аспірантка 1-го року навчання Державної наукової установи «Інститут інформації, безпеки і права Національної академії правових наук України», викладач кафедри інформаційного, господарського та адміністративного права Національного технічного університету України «Київський політехнічний інститут»

ORCID ID 0009-0006-9150-626X

**Науковий керівник:**

**Заярний О.А.**

д.ю.н., професор, професор кафедри інтелектуальної власності та інформаційного права Навчально-наукового інституту права Київського національного університету імені Тараса Шевченка

## **ПРАВО ОСОБИ НА СПРАВЕДЛИВЕ РІШЕННЯ В АДМІНІСТРАТИВНИХ ПРОЦЕДУРАХ ІЗ ЗАСТОСУВАННЯМ ШТУЧНОГО ІНТЕЛЕКТУ СУБ'ЄКТАМИ ПУБЛІЧНОГО АДМІНІСТРУВАННЯ**

Активна цифровізація публічної сфери зумовлює впровадження автоматизованих інформаційних систем та штучного інтелекту (далі – ШІ) у діяльність суб'єктів публічного адміністрування, зокрема у сферу адміністративних процедур. З одного боку, це підвищує ефективність і швидкість розгляду адміністративних справ, сприяє автоматизації адміністративних процесів, мінімізує людський фактор, зменшує корупційні ризики та забезпечує оперативний доступ до адміністративних послуг. А з іншого – безпосередньо впливає на захист прав і свободи людини, створюючи нові виклики для їх реалізації.

У контексті застосування ШІ особливого значення набувають гарантії здійснення адміністративних процедур, такі як обґрунтованість та законність адміністративних рішень; доступ до інформації; права особи на участь в

ухваленні рішення та ефективні засоби правового захисту. За своєю природою штучний інтелект функціонує на основі заданих інструкцій та алгоритмів, що дозволяє йому формувати висновки або рішення. Водночас на практиці такі рішення можуть бути необґрунтованими, непрозорими або позбавленими логічного мотивування, що ускладнює доступ особи до інформації, розуміння причин прийняття рішення, порушує принцип рівності учасників адміністративної процедури та породжує юридичну невизначеність.

Додаткову проблему становить потенційна «алгоритмічна дискримінація», коли система штучного інтелекту відтворює упередженість, що закладена у вихідні дані, або технічні помилки, включно з так званими «ШІ-галюцинаціями». Це може призводити до неправомірних рішень у сфері надання адміністративних послуг, здійснення реєстраційних, дозвільних, інспекційних та інших адміністративних процедур. Відсутність визначеної законом відповідальності за такі рішення створює ризик перекладення вини на «помилку програми», що суперечить принципу верховенства права та унеможливорює забезпечення захисту прав і свобод особи, а відтак і справедливості адміністративного рішення.

Міжнародна практика демонструє інший підхід. Так, Акт про штучний інтелект (AI Act) класифікує застосування систем штучного інтелекту органами державної влади або від їх імені як системи високого ризику. Такий тип систем підлягає суворому регулюванню, що передбачає підвищені вимоги щодо прозорості, точності, надійності, ведення технічної документації та обов'язкового людського контролю [1]. Крім того, відповідно до статті 22 Загального регламенту про захист даних (далі – GDPR) суб'єкт даних має право не підлягати рішенню, що ґрунтується винятково на автоматизованому опрацюванні, в тому числі, профайлінгу, що породжує правові наслідки для чи подібним чином істотно впливає на нього або неї [2]. У поєднанні із статтями 13-15 GDPR це включає право на отримання достовірної інформації про логіку опрацювання даних [2]. Таким чином, європейське регулювання закріплює обов'язковість людського контролю, прозорості алгоритмів та можливості оскарження рішень, ухвалених за участю ШІ. Натомість в Україні закріплення таких гарантій на законодавчому рівні перебуває на початковому етапі.

У жовтні 2019 року Україна приєдналася до Рекомендацій Організації економічного співробітництва і розвитку з питань штучного інтелекту, а вже в грудні 2020 року Уряд України схвалив Концепцію розвитку штучного

інтелекту в Україні (далі – Концепція). У ній зазначено визначення штучного інтелекту як організованої сукупності інформаційних технологій, із застосуванням якої можливо виконувати складні комплексні завдання шляхом використання системи наукових методів дослідження 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 визначити способи досягнення поставлених завдань [3].

Серед основних проблем, визначених у Концепції, – недосконалість механізмів прийняття управлінських рішень у публічній сфері, забюрократизованість системи надання адміністративних послуг, обмеженість доступу до інформації та її низька якість, недостатній рівень впровадження електронного документообігу між державними органами, а також низький ступінь оцифрованості даних, що перебувають у власності державних органів [3]. Відповідно до Плану заходів з реалізації Концепції у 2026 році планується розробка та подання Кабінету Міністрів України законопроекту щодо правового врегулювання у сфері розвитку штучного інтелекту [4].

Окрім цього, стаття 8 Закону України «Про адміністративну процедуру» (далі – Закон) закріплює принцип обґрунтованості, який покладає обов’язок на адміністративний орган обґрунтовувати адміністративні акти, які він приймає. Якщо ж адміністративний акт негативно впливає на права, свободи чи законні інтереси особи, він обов’язково повинен містити мотивувальну частину [5]. У контексті застосування штучного інтелекту, виникає питання, чи здатен ШІ забезпечити належне юридичне обґрунтування рішення, сформоване зрозумілою мовою, з безпомилковим застосуванням юридичної термінології та можливістю в майбутньому реалізувати особі, щодо якої приймався адміністративний акт, право на його оскарження. Таким чином, навіть за участі штучного інтелекту обов’язок належної мотивації рішення має залишатися за адміністративним органом, а відсутність зрозумілого обґрунтування ставить під сумнів законність та справедливість адміністративної процедури в цілому.

Також відповідно до частини 2 статті 71 Закону України «Про адміністративну процедуру» адміністративний акт має містити підпис та/або печатку (у тому числі електронні), якщо інше не передбачено законом, та повне ім’я відповідальної посадової особи адміністративного органу [5].

Навіть у разі здійснення адміністративного провадження в автоматичному режимі, відповідальність за адміністративні акти, ухвалені в автоматичному режимі, продовжує нести адміністративний орган, що прямо передбачено частиною 3 статті 62 Закону [5]. Це зумовлює логічне питання про можливість покладення юридичної відповідальності за прийняті рішення на штучний інтелект. Однак ШІ наразі не наділений правосуб'єктністю, у тому числі й деліктоздатністю, а тому остаточною відповідальність за зміст адміністративного акта покладається на відповідного суб'єкта владних повноважень.

Показовою є позиція Касаційного суду у складі Верховного Суду про те, що штучний інтелект може бути лише допоміжним інструментом у сфері правосуддя. А технологію слід використовувати лише для підтримки та посилення верховенства права. Водночас Верховний Суд наголошує, що ухвалення рішень має, явно чи неявно, здійснюватися лише суддями, без можливості делегування або виконання за допомогою технології [6].

Ця позиція Суду є релевантною і для сфери публічного адміністрування: в адміністративних процедурах штучний інтелект не може замінити адміністративний орган, оскільки це усуне людський контроль за законністю, справедливістю та вмотивованістю адміністративного акта. Безумовно ШІ може забезпечувати ефективність у діяльності суб'єктів публічного адміністрування шляхом прискорення обробки даних та зменшення навантаження, однак відповідальність за остаточне рішення та його правові наслідки повинна залишатися за суб'єктом владних повноважень. Такий підхід забезпечуватиме дотримання правовладдя та гарантії права особи на справедливе адміністративне рішення.

На підставі викладеного необхідним є законодавче визначення гарантій застосування штучного інтелекту в адміністративних процедурах. До них можуть належати:

- обов'язкове повідомлення особи про використання ШІ під час прийняття адміністративного акта;
- право особи отримати зрозуміле пояснення мотивів та логіки прийнятого рішення;
- забезпечення людського контролю над результатом роботи системи штучного інтелекту;
- чіткий механізм оскарження рішення адміністративного органу;

- закріплення відповідальності суб'єкта владних повноважень за прийняті рішення із застосуванням ШІ;
- можливість аудиту алгоритмів ШІ на предмет точності, недискримінації та коректності обробки даних.

Запровадження таких гарантій дасть змогу забезпечити прийняття справедливих та обґрунтованих адміністративних рішень навіть за умови використання ШІ.

Таким чином, впровадження штучного інтелекту в адміністративні процедури є неминучим етапом в цифровій трансформації публічного адміністрування, однак воно не може відбуватися без належних правових механізмів захисту. ШІ здатний підвищити ефективність діяльності суб'єктів публічної адміністрації, проте не забезпечує належної обґрунтованості рішень та не несе відповідальності за них. Це обумовлює необхідність збереження людського контролю за законністю та обґрунтованістю індивідуальних актів, а також формування законодавства у сфері застосування ШІ, спрямованого на захист прав і свобод людини. Гарантованість мотивованого адміністративного рішення, можливість його оскарження та чітке визначення відповідального суб'єкта є ключовими умовами інтеграції штучного інтелекту у сферу публічного управління та адміністрування відповідно до принципу верховенства права, що забезпечує прийняття справедливих адміністративних рішень.

### **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). URL: [https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=OJ:L\\_202401689](https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=OJ:L_202401689) (дата звернення: 04.11.2025 р.).

2. РЕГЛАМЕНТ ЄВРОПЕЙСЬКОГО ПАРЛАМЕНТУ І РАДИ (ЄС) 2016/679 від 27 квітня 2016 року про захист фізичних осіб у зв'язку з опрацюванням персональних даних і про вільний рух таких даних, та про скасування Директиви 95/46/ЄС (Загальний регламент про захист даних). URL: [https://zakon.rada.gov.ua/laws/show/984\\_008-16#Text](https://zakon.rada.gov.ua/laws/show/984_008-16#Text) (дата звернення: 04.11.2025 р.).

3. Концепція розвитку штучного інтелекту в Україні, схвалена розпорядження Кабінету Міністрів України № 1556-р від 02.12.2020 р. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text> (дата звернення: 04.11.2025 р.).

4. План заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2025-2026 роки, затвердженої розпорядженням Кабінету Міністрів України № 457-р від 09.09.2025 р. URL: <https://zakon.rada.gov.ua/laws/show/457-2025-%D1%80#Text> (дата звернення: 04.11.2025 р.).

5. Про адміністративну процедуру: Закон України № 2073-IX від 17.02.2022 р. URL: <https://zakon.rada.gov.ua/laws/show/2073-20#Text> (дата звернення: 04.11.2025 р.).

6. Ухвала Касаційного господарського суду у складі Верховного Суду у справі № 925/2000/22 від 08.02.2024 р. URL: <https://reyestr.court.gov.ua/Review/116984639> (дата звернення: 04.11.2025 р.).

**Наталія Дедюєва**

студентка 5 курсу

НТУУ «КПІ імені Ігоря Сікорського»

ORCID: <https://orcid.org/0009-0005-5065-1228>

**Ольга Головка**

медіатор, кандидат юридичних наук, старший дослідник, с.н.с.

Державна наукова установа «Інститут

інформації, безпеки і права НАПрН України»

ORCID: <https://orcid.org/0000-0001-8963-6598>

## **БАЛАНС ПРИВАТНОСТІ ТА СВОБОДИ ВИРАЖЕННЯ ПОГЛЯДІВ ЧЕРЕЗ ПРИЗМУ ЗАКОНОПРОЄКТУ № 14057**

У сучасних умовах цифровізації та глобального поширення інформації питання балансу між свободою слова і захистом приватності особи набуває особливої актуальності. В Україні це питання перебуває у фокусі сучасної правотворчої діяльності, зокрема у контексті Законопроєкту №14057 (далі — Законопроєкт). Він спрямований на оновлення положень Цивільного кодексу України (далі — ЦКУ) щодо захисту особистих немайнових прав у сфері інформаційних відносин. Запропоновані зміни охоплюють низку інноваційних положень від визначення меж достовірності публікацій до нового погляду на право на відповідь. Такі нововведення формально узгоджуються з існуючими тенденціями права ЄС. Водночас запропоновані новели викликають потребу у критичній оцінці їх пропорційності та відповідності принципу свободи вираження поглядів. Наявна редакція Законопроєкту створює додаткові заборони та правові підстави для понаднормового контролю [1]. Запропоновані Законопроєктом нововведення варто дослідити як з точки зору суспільної значущості, так і через призму потенційних загроз свободі слова та діяльності медіа.

Одна з найбільш дискусійних новел — **пропозиція автоматично вважати недостовірною будь-яку інформацію про особу, яка не підтверджена обвинувальним вироком суду**. Вона передбачена статтею 277 Законопроєкту [1]. Фактично йдеться про те, що публікація відомостей щодо можливих правопорушень до прийняття рішення судом порушує особисті

немайнові права особи, оскільки суперечить презумпції невинуватості. Для порівняння, ст. 277 ЦКУ передбачає право особи вимагати спростування недостовірної інформації. Проте дана стаття не визначає автоматично, яку інформацію вважати недостовірною [2]. Нова ж норма істотно зміщує баланс доказування, адже на медіа покладається обов'язок доводити правдивість відомостей через вирок суду, інакше такі відомості вважаються неправдивими. Зокрема, в частині 2 статті 277 Законопроєкту міститься наступне положення: «Недостовірною вважається інформація, яка порушує презумпцію невинуватості, доки вину особи не доведено у визначеному законом порядку і не встановлено обвинувальним вирок суду, що набрав законної сили» [1]. Такий підхід виходить за межі усталених підходів до питання дифамації і, по суті, суттєво звужує можливості журналістських розслідувань [3]. На наш погляд, це може призвести до надмірного навантаження на судову систему, а також створити додаткові перешкоди для публікацій, наприклад про зловживання владою. На нашу думку, запропонована в Законопроєкті редакція ст. 277. створює колізію із принципом свободи вираження поглядів, оскільки встановлює надмірно жорсткий критерій доказовості щодо розслідування журналістами суспільно значущих фактів. Відповідно до ЦКУ, захист честі і репутації має здійснюватися шляхом спростування, відповіді (стаття 277 ЦКУ), відшкодування (стаття 280 ЦКУ) у разі доведеної недостовірності повідомленого в публікації [2]. На нашу думку, більш доречно розглядати цей аспект не через превентивну заборону публікацій на тій підставі, що відсутній вирок суду. Таким чином, дане нововведення перекладає тягар доказування на медіа та створює істотний дисбаланс на користь «репутаційної безпеки» проти свободи слова.

Законопроєкт також істотно розширює **право на відповідь та спростування**. Практика Європейського суду з прав людини (далі — ЄСПЛ) щодо «права на відповідь» визнає це право як пропорційну та обмежену реакцію на дифамаційні твердження, спрямовану на забезпечення балансу між свободою вираження і захистом репутації. Зокрема, ЄСПЛ у справі *Melnychuk v. Ukraine* вказав, що позитивне зобов'язання держави може полягати у наданні заявнику реальної можливості реалізувати право на відповідь, тобто подати відповідь для публікації у відповідь на медіаматеріал [4]. Суд наголошує, що у певних випадках, позитивний обов'язок держави полягає саме у забезпеченні права на відповідь, а не у жорсткіших санкціях або вилученні матеріалів.

У справі *Hachette Filipacchi v. France* ЄСПЛ деталізував, що вимога опублікувати відповідь є менш обтяжливою для свободи вираження, ніж вилучення самого матеріалу або накладення більш суворих обмежень [5]. Відповідно до ст. 277 ЦКУ, фізична особа, чиї права порушено поширенням щодо неї недостовірної інформації, має право на відповідь та на спростування цієї інформації. В Законопроекті положення ст. 277 пропонують надавати право на спростування будь-якій особі, згаданій у медіа, навіть якщо її ім'я прямо не названо (достатньо, що її можна ідентифікувати або віднести до згаданої групи осіб). Це випливає із ч. 4 ст. 277 Законопроекту: «Інформація вважається такою, що поширена про особу, якщо з неї можна достеменно встановити, що вона стосується конкретної особи або ця особа включається до кола осіб, яких стосується інформація» [1].

Право на відповідь, передбачене у ст. 277-1 Законопроекту, виникає незалежно від того, чи є допущена в публікації інформація недостовірною. Кожна згадка особи у негативному контексті чи наявність критики в сторону особи може спричинити вимогу опублікувати її відповідь. Це створює загрозу знищення редакційного контролю і професійних стандартів журналістики, покладаючи на медіа обов'язок публікувати будь-які репліки у відповідь. Це виходить за рамки чинної ст. 277 ЦКУ, що гарантує право на відповідь у зв'язку із поширенням недостовірної інформації, і створює колізію з нормами про свободу вираження поглядів. На нашу думку, запропоноване в Законопроекті розширення права на відповідь призведе до зловживань ним та спричинить загрозу для редакційної автономії. ЄСПЛ наголошує, що право на відповідь має забезпечуватися переважно як засіб реагування на недостовірні або дифамаційні твердження, а не на будь-яку негативну чи критичну оцінку. У рішенні ЄСПЛ у справі *Eker v. Turkey* Суд розкрито зміст права на відповідь у контексті претензій на порушення свободи висловлювань. З практики ЄСПЛ випливає, що право на відповідь виникає в разі публікації недостовірної або образливої інформації, що порушує репутацію, і гарантує можливість якісного й оперативного виправлення цієї інформації у публічному просторі, забезпечуючи баланс між захистом права на честь і свободою слова [6]. Цей кейс засвідчує, що право на відповідь не є абсолютним та застосовується для виправлення конкретних неправдивих чи образливих фактів без обмеження загальної свободи висловлювань.

Окремо слід відзначити новелу, яка дозволяє **стягувати компенсацію моральної шкоди за оціночні судження, якщо форма їх подачі принижує гідність чи репутацію особи** (ст. 277 Законопроекту). Чинне законодавство виходить з того, що оціночні судження не є фактичними твердженнями, а отже не підлягають спростуванню чи відшкодуванню. Такий підхід узгоджується і з практикою ЄСПЛ [7]. Проект же пропонує наступну зміну якщо оцінка чи критична думка висловлена в образливій, принизливій формі, то може наставати відповідальність у вигляді відшкодування моральної шкоди. Така новела може суперечити принципу пропорційності втручання. Будь-яка емоційно забарвлена оцінка може кваліфікуватися як образлива, що суперечить європейським стандартам захисту репутації, де критика публічних осіб допускається у ширших межах. В цьому контексті доречно згадати рішення ЄСПЛ *Oberschlick v. Austria*, адже в ньому висвітлено питання різниці меж допустимої критики стосовно політичного діяча та приватної особи. Європейський суд з прав людини у цьому рішенні заклав фундаментальні підходи до визначення меж свободи вираження та захисту репутації публічних осіб. Зазначається, що у демократичному суспільстві преса виконує місію критичного аналізу ініціатив політичних діячів й повинна бути вільною у виборі такої форми коментування, яка найбільш адекватно відповідає її меті. Свобода слова, гарантована статтею 10 Конвенції, включає в себе не лише нейтральні або ж суспільно прийнятні твердження, а й ті, що можуть шокувати чи ображати державу або частину суспільства. Принципово, що саме до висловлювань щодо публічних осіб встановлюються розширені межі допустимої критики, адже політик свідомо піддає свої слова і дії суспільному контролю. ЄСПЛ зазначає, що форма подачі критичного матеріалу сама по собі не може вважатися підставою для обмеження свободи слова, якщо вона не вводить в оману читача щодо фактів і не містить прямих необґрунтованих звинувачень. Згідно з вищезгаданим рішенням, втручання держави у свободу слова в даній справі не було визнане «необхідним у демократичному суспільстві... для захисту репутації... інших осіб». З огляду на це, мало місце порушення статті 10 Конвенції. Саме такі підходи ЄСПЛ доречно імплементувати в українське законодавство. Зокрема, мова йде про встановлення спеціальних винятків для оціночних суджень, адресованих публічним особам, та ширших меж для політичної критики. Це сприятиме захисту свободи вираження в рамках європейських стандартів [8]. В контексті

нашого дослідження це ключовий момент, адже ми вважаємо за доцільне проводити законодавче розмежування допустимої критики для різних категорій осіб публічних та приватних. В тому числі це необхідно для зменшення навантаження на судову систему, адже чітке правове регулювання зможе зменшити кількість непорозумінь що можуть стати предметом позову на основі запропонованої в Законопроекті статті.

Зміст окремих запропонованих новел викликає сумнів щодо пропорційності обмежень. Запропонований у Законопроекті підхід до захисту особистих немайнових прав, зокрема, через автоматичне визнання інформації недостовірною без наявності обвинувального вироку створює суттєві ризики для свободи слова і журналістських розслідувань. Переклад тягаря доказування на медіа та встановлення жорстких критеріїв перевірки інформації загрожує надмірним цензурним контролем. У цьому контексті розширене право на відповідь без обмежень щодо достовірності інформації порушує баланс між захистом честі та свободою слова. Практика ЄСПЛ демонструє, що право на відповідь має бути пропорційним, обмеженим реакцією саме на недостовірні або дифамаційні твердження, а не використаним задля унеможливлення критики публічних осіб. З огляду на проведений аналіз Законопроекту, ЦКУ та рішень ЄСПЛ вважаємо за доцільне: 1) конкретизувати широкі формулювання (обмежити право на відповідь випадками недостовірної інформації або очевидного порушення приватності, виключити бездоказове ототожнення непідтверджених даних із недостовірними); 2) внести уточнення, що висвітлення питань суспільного інтересу є проявом свободи слова та не можуть бути обмежені; 3) зберегти діалог із медіа спільнотою та експертами при подальшому доопрацюванні Законопроекту.

### **Список використаних джерел**

1. Проект Закону про внесення змін до Цивільного кодексу України у зв'язку із оновленням (рекодифікацією) положень книги другої : Законопроект № 14057 від 21.09.2025 р. URL: <https://itd.rada.gov.ua/billinfo/Bills/Card/57355> (дата звернення: 25.10.2025).

2. Цивільний кодекс України : Закон України від 16.01.2003 № 435-IV // Відомості Верховної Ради України. 2003. № 40–44. Ст. 356. URL: <https://zakon.rada.gov.ua/laws/show/435-15#Text> (дата звернення: 25.10.2025).

3. Гуйван П. Д. Питання правового відмежування дифамації від допустимої критики в мас-медіа. Південноукраїнський правничий часопис. 2020. № 2. С. 60–70. URL: <http://www.sulj.oduvs.od.ua/archive/2020/2/12.pdf> (дата звернення: 25.10.2025).

4. Melnychuk v. Ukraine, Application no. 28743/03, European Court of Human Rights, 5 July 2005. URL: <https://hudoc.echr.coe.int/eng?i=001-70089> (дата звернення: 25.10.2025)

5. Hachette Filipacchi Associés v. France, Application no. 71111/01, European Court of Human Rights, 14 June 2007. URL: <https://hudoc.echr.coe.int/eng?i=001-81066> (дата звернення: 25.10.2025).

6. Eker v. Turkey, Application no. 24016/05, European Court of Human Rights, 24 October 2017. URL: <https://hudoc.echr.coe.int/eng?i=001-177928> (дата звернення: 25.10.2025).

7. Справа «Сірик проти України» (SIRYK v. UKRAINE), no. 001-104281, ECHR, 5 October 2010. URL: <https://hudoc.echr.coe.int/eng?i=001-104281> (дата звернення: 25.10.2025).

8. Oberschlick v. Austria, no. 001-57716, ECHR, 23 May 1991. URL: <https://hudoc.echr.coe.int/tur?i=001-57716> (дата звернення: 25.10.2025).

**Юрій Капіца**

доктор юридичних наук, директор Центру  
досліджень інтелектуальної власності та  
трансферу технологій НАН України  
ORCID: <https://orcid.org/0000-0002-9449-8422>

## **ПРАВОВІ АСПЕКТИ ВИКОРИСТАННЯ КОНТЕНТУ ДЛЯ НАВЧАННЯ ШТУЧНОГО ІНТЕЛЕКТУ: ПІДХОДИ В ЄС, УКРАЇНІ ТА США**

Розвиток генеративного штучного інтелекту (ШІ) обумовив правові виклики, пов'язані із використанням контенту для навчання моделей ШІ. Основним питанням є: чи є тренування ШІ на масивах даних, що містять твори, які охороняються авторським правом, правомірним без окремого дозволу правовласників. В Україні системний аналіз випадків використання контенту для навчання ШІ наведено в Рекомендаціях щодо відповідального використання штучного інтелекту, підготовлених Міністерством цифрової інформації України, Українським національним офісом інтелектуальної власності та інновацій, іншими організаціями у 2024 р. [1]. В тезах звертається увага на досвід ЄС та США, які демонструють відмінні правові підходи.

### **1. Європейський Союз**

Стаття 4 Директиви Європейського Парламенту і Ради (ЄС) 2019/790 від 17 квітня 2019 року про авторське право і суміжні права на Єдиному цифровому ринку та про внесення змін до директив 96/9/ЄС та 2001/29/ЄС (CDSM) передбачає винятки щодо відтворення та витягування даних із творів та інших об'єктів, до яких особи мають законний доступ, з метою проведення глибинного аналізу тексту та даних (TDM) за умови, що правовласники не заборонили це у належний спосіб, такий як за допомогою засобів машинного зчитування у випадку контенту, доведеного до загального відома публіки онлайн. Більший обсяг прав TDM надається науково-дослідними організаціями та установами зі збереження культурної спадщини (ст. 3) [13].

Регламент Європейського Парламенту та Ради (ЄС) 2024/1689 від 13 березня 2024 року, що встановлює гармонізовані правила щодо штучного інтелекту (AI Act), запроваджує зобов'язання щодо прозорості: провайдери

генеративного ШІ повинні публікувати «деталізовану інформацію» про використані дані [14].

В той же час низка дослідників (Quintais, Buick, Dornis та ін.) наголошує, що формально ст. 4 Директиви 2019/790 надає можливість навчання ШІ, однак на практиці межа між аналізом і відтворенням даних є розмитою [2-4]. Це ставить під сумнів застосовність винятку ст. 4 CDSM до навчання LLM-моделей. Dornis і Stober (2025) вказуються, що моделі ШІ зберігають і перетворюють твори, що виходить за межі TDM [4] .

У звітах Європейського Парламенту, EUIPO, Європейського товариства авторського права (ECS) за 2024–2025 роки наголошується на потребі уточнення положень законодавства та практики. Серед пропозицій – добровільні кодекси поведінки, розробка ліцензійних моделей для навчання, прозорість щодо датасетів [5-7].

На цей час в ЄС, за виключення одиничних справ, відсутня розвинути судова практика щодо навчання ШІ. Справою Robert Kneschke v. LAION e.V. (Hamburg Regional Court, Germany], Case No. 310 O 227/23, 2024) підтверджено застосування винятків TDM для науково-дослідних організацій. Поточна справа у Великій Британії Getty Images & others v. Stability AI Ltd (High Court, England & Wales (слухання 2024–2025) стосується використання візуальних творів (зібраних з Інтернет) для тренування ШІ. У справі розглядалися як «input» claims (завантаження контенту в датасет), так і «output» claims (створені AI-зображення, які, за твердженням позивачів, охоронювані елементи). У ході процесу з'ясувалося, що позивач не зміг довести, що «substantial part» її захищених творів була скопійована або використана так, щоб модель ШІ створювала outputs, які прямо містять чи повторюють суттєву частину її захищених зображень. Також, розгляд справи виявив юрисдикційні проблеми застосування права Великої Британії у зв'язку з тим, що тренування ШІ відбувалися поза межами країни.

## **2. США**

У США при навчанні ШІ застосовується доктрина справедливого використання (fair use). Відповідно до неї дозволено обмежене використання творів без згоди правовласника для трансформативних цілей, критики, освіти чи досліджень. На думку Lee (2025), тренування ШІ може кваліфікуватися як трансформативне використання, особливо якщо воно не замінює на ринку оригінальні твори [8]. У практиці (Torrance & Tomlinson, 2023)

виокремлюються чотири ключові критерії застосування доктрини: мета й характер використання, природа оригінального твору, обсяг використання і вплив на ринок [9].

У звіті U.S. Copyright Office 2025 року (USCO, Part 3) зазначено, що навчання генеративного ШІ потенційно порушує права на відтворення та створення похідних творів. USCO визнає, що законодавство США не враховує сучасних технологічних реалій і рекомендує здійснення аналізу на рівні окремих моделей і ситуацій [10].

У справі *The New York Times Company v. OpenAI Inc & Microsoft Corp.* (U.S. District Court for the Southern District of New York; ongoing dispute, initial complaint filed February 2023), розгляд якої на цей час триває, розглядається позов видавців щодо використання статей медіа видань для тренування моделей без ліцензії. Позов включає твердження про копіювання та відтворення. Очікується, що вирішення справи стане прецедентом для визначення меж fair use щодо ШІ.

Також мають значення рішення за справами *Authors Guild, Inc. v. Google, Inc.* (2d Cir., США, 2015) щодо «book scanning» та *Authors Guild v. HathiTrust* (2d Cir., 2014) стосовно HathiTrust digital library. Судами ухвалено, що сканування книг і індексація для створення пошукової бази даних здійснюється у відповідності з доктриною fair use. Хоча справи не стосувалися навчання ШІ, судова логіка щодо трансформативності часто цитується у аргументації щодо training data на користь провайдерів ШІ. В той же час зазначається, що технічна природа навчання ШІ, можливість відтворення фрагментів текстів є відмінним від Google Books.

Істотним в умовах нечіткості законодавства є позиція компаній-розробників генеративного ШІ. У політиках Microsoft (2023), Adobe (2024) та OpenAI (2024) декларується, що при навчанні ШІ використовуються лише правомірні дані (наприклад, Adobe – лише Adobe Stock). Крім того, компанії беруть на себе обов'язки щодо юридичного захисту користувачів – у разі претензій правовласників.

### **3. Україна**

В рамках наближення національного законодавства в Україні до законодавства ЄС проблематика використання творів для навчання ШІ є близькою до ситуації в ЄС та обумовлює потреби у вивченні за застосуванні

перспективних рішень, які на цей час пропонуються в ЄС, а також застосуванні підходів провідних світових компаній – розробників ІІІ.

Розроблені у 2024 р., Рекомендації щодо відповідального використання штучного інтелекту [1] звертають увагу на використання при навчання ІІІ об'єктів, які перейшли до суспільного надбання; які не охороняються авторським правом; використання на підставі правочину, у тому числі ліцензій відкритого доступу (якщо такими ліцензіями передбачаються застосовні ІІІ способи використання творів); на підставі винятків та обмежень, встановлених законом. Також надані конкретні рекомендації правовласникам.

#### **4. Перспективи врегулювання**

Передбачається, що практика застосування ІІІ провідними компаніями – розробниками ІІІ, судові справи в ЄС та США визначають напрямки правомірного використання контенту для навчання ІІІ, а також потреби зі змін законодавства. З врахуванням неможливості зупинити технологічний розвиток - законодавство буде підлаштовуватися під провідну практику, яка, з одного боку, дозволяє розвиток систем ІІІ, та, з іншого, не спричиняє шкоду суб'єктам авторського права.

На цей час, правовою основою використання творів для навчання ІІІ в ЄС є ст. 3–4 Директиви (ЄС) 2019/790, якщо відсутні технічні заборони. В Україні потребує імплементації в повному обсязі статті 3 та 4 директиви, що розрізняються обсяг прав для здійснення TDM науково-дослідними організаціями та установами зі збереження культурної спадщини та іншими особами.

У США використання має відповідати доктрині fair use, особливо якщо воно є трансформативним і не шкодить ринку (Lee, 2025) [8].

Також актуальним є вивчення впливу на ринок механізмів дотримання авторського права та захисту користувачів провідними компаніями – розробниками ІІІ - Microsoft, Adobe, OpenAI.

Стосовно перспектив врегулювання слід вказати на наступне.

1. Вдосконалення законодавства – Офіс авторського права США рекомендує адаптувати законодавство до викликів генеративного ІІІ [10]. Аналогічно, у ЄС обговорюється оновлення CDSM-директиви [6].

2. Кодекси поведінки – Єврокомісія планує розробку добровільних кодексів доброчесності для провайдерів ІІІ, які включатимуть обов'язкове розкриття джерел даних [5].

3. Централізоване ліцензування – пропонується створення нових систем колективного ліцензування контенту для навчання [2, 7].

4. Технічні засоби захисту (TPM) – застосування метаданих, протоколів opt-out (наприклад, noai.txt) як цифрового способу заборони використання (Adobe, 2024; OpenAI Policies, 2024).

5. Публічні реєстри даних – ініціативи зі створення прозорих реєстрів тренувальних датасетів [11].

6. Медіація і позасудове врегулювання – пропонується посилити механізми альтернативного вирішення спорів між правовласниками і AI-платформами [1].

7. WIPO Conversation – міжнародні обговорення, що здійснюються WIPO, орієнтовані на розробку глобальних принципів балансу між інноваціями і правами авторів[12].

**Висновки.** Регулювання використання авторського контенту для тренування ШІ перебуває на стику авторського права та цифрових інновацій. Підхід ЄС базується на нормативному регулюванні з винятками для TDM, тоді як США використовують гнучкий підхід fair use. Наразі немає єдиного стандарту «правомірного навчання ШІ», проте формуються кілька напрямків, що стосуються прозорості, саморегулювання, колективного ліцензування, технічних засобів заборони, медіації та оновлення законодавства.

### Список використаних джерел

1. Рекомендації щодо відповідального використання ШІ: питання права інтелектуальної власності. Міністерство цифрової трансформації України. УКРНОІВІ та інш. 2024. 51 с.

[https://nipo.gov.ua/wp-content/uploads/2025/01/Rekomendatsii\\_shchodo\\_ShI\\_ta\\_IV.pdf](https://nipo.gov.ua/wp-content/uploads/2025/01/Rekomendatsii_shchodo_ShI_ta_IV.pdf)

2. Quintais, J.P. Generative AI, Copyright and the AI Act. SSRN. 2025. URL: [https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=4912701](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4912701)

3. Buic, A. Copyright and AI training data — transparency to the rescue? *Journal of Intellectual Property Law & Practice*. 2025. <https://doi.org/10.1093/jiplp/jpae102>

4. Dornis, T. & Stober, S. Generative AI Training and Copyright Law. arXiv. 2025. URL: <https://arxiv.org/abs/2502.15858>

5. EUIPO. The development of Generative Artificial Intelligence from a copyright perspective. 2025. URL: <https://www.euipo.europa.eu/en/news/euipo-releases-study-on-generative-artificial-intelligence-and-copyright>

6. European Parliament. Generative AI and Copyright — Training, Creation, Regulation. European Parliament. 2025. URL: [https://www.europarl.europa.eu/thinktank/en/document/IUST\\_STU\(2025\)774095](https://www.europarl.europa.eu/thinktank/en/document/IUST_STU(2025)774095)
7. European Copyright Society. Opinion on Copyright and Generative AI. 2025. URL: [https://europeancopyrightsociety.org/wp-content/uploads/2025/02/ecs\\_opinion\\_genai\\_january2025.pdf](https://europeancopyrightsociety.org/wp-content/uploads/2025/02/ecs_opinion_genai_january2025.pdf)
8. Lee, E. Fair Use and the Origin of AI Training. SSRN. 2025. URL: [https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=5253011](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5253011)
9. Torrance, A.W. & Tomlinson, B. Training Is Everything: Artificial Intelligence, Copyright, and Fair Training. arXiv. 2023. URL: <https://arxiv.org/abs/2305.03720>
10. USCO. Copyright and Artificial Intelligence: Part 3 — Generative AI Training. U.S. Copyright Office. 2025. URL: <https://www.copyright.gov/ai/Copyright-and-Artificial-Intelligence-Part-3-Generative-AI-Training-Report-Pre-Publication-Version.pdf>
11. Kluwer IP Blog. Generative AI, copyright and the AI Act. 2023. URL: <https://copyrightblog.kluweriplaw.com/2023/05/09/generative-ai-copyright-and-the-ai-act/>
12. WIPO Conversation on IP and Frontier Technologies. URL: [https://www.wipo.int/meetings/en/details.jsp?meeting\\_id=84809](https://www.wipo.int/meetings/en/details.jsp?meeting_id=84809)
13. Directive (EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on copyright and related rights in the Digital Single Market and amending Directives 96/9/EC and 2001/29/EC.
14. 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)

**Олена Бахарева**

професіонал з інтелектуальної власності

1 категорії

ORCID: <https://orcid.org/0000-0002-4599-9781>

## **ШТУЧНИЙ ІНТЕЛЕКТ ЯК «ПОРУШНИК» АВТОРСЬКОГО ПРАВА**

Історія штучного інтелекту починається ще зі стародавнього світу, коли ширилися міфи, оповідання та чутки про створення штучних істот, яких обдаровували розумом чи свідомістю. «Зерна сучасного штучного інтелекту заклали філософи, які намагалися описати процес людського мислення як механічне маніпулювання символами». У 1940-х роках був винайдено програмований цифровий комп'ютер, машину, що ґрунтувалася на абстрактній сутності математичного міркування [1]. В умовах сьогодення штучний інтелект став невід'ємною частиною нашого повсякденного життя та вийшов за межі концепції з наукової фантастики. Генеративний штучний інтелект, такий як ChatGPT та MidJourney, зробили справжній бум на ринку цифрового контенту і продовжують стрімко розвиватись далі. З'являються нові інструменти, що будуються на базі штучного інтелекту, завданням яких є спрощення людського життя. Водночас існує чимало аналітичних прогнозів про те, що творчі професії можуть стати неактуальними для людей, адже штучний інтелект набагато швидше і якісніше, ніж звичайна людина, може виконати завдання щодо написання тексту чи створення зображень [2]. Штучний інтелект стає новим викликом для економіки та правової системи і відповідно новим об'єктом для правового регулювання. Штучний інтелект здатний генерувати та створювати різні твори – науки, літератури і мистецтва, що стає проблемою для визнання авторства, можливості розпорядження авторами своїми правами і використання ними механізмів правової охорони об'єктів ІВ [3]. Тому на сучасному етапі економічного розвитку конкурентних відносин, штучний інтелект може стати новою формою порушення прав на об'єкти інтелектуальної власності.

Разом із розвитком штучного інтелекту трансформується відповідно і авторське право. Виникають питання, хто є автором результатів роботи штучного інтелекту, на чийх творах штучний інтелект навчається, чи не порушуються при цьому права авторів? Поступово, широке поле для дискусій

та судова практика стануть основою для заповнення законодавчих прогалин [4]. Штучний інтелект нікого не залишає байдужим. Хтось бачить в ньому загрозу, хтось – можливості. Часто можна почути, що розвиток штучного інтелекту призведе до зростання безробіття, порушень приватності, і навіть «повстання машин». Натомість прибічники штучного інтелекту запевняють, що він буде досить ефективним у різних сферах: медицині, військовій справі, промисловості. Одне з питань, яке активно обговорюють сьогодні – регулювання штучного інтелекту. Адже він впливає на всі сфери нашого життя – від національної безпеки до освіти, від сфери оборони до медицини. За регулювання штучного інтелекту взялися США, Китай та ЄС [5]. Вплив штучного інтелекту на різні сфери нашого життя стає набагато потужнішим, і це вимагає врегулювати його на законодавчому рівні. Довгий час люди не надавали значення такому регулюванню, проте зараз воно необхідне як ніколи. Його відсутність може призвести до великої кількості порушень безпеки персональних даних, потенційної дискримінації користувачів, загрози створення штучного інтелекту для військової агресії та інших негативних наслідків [6]. Штучний інтелект увірвався у всі сфери нашого життя: від автовідповідача в банку до архітектури. За останні декілька років інновації у створенні нового контенту за допомогою штучного інтелекту сколихнули суспільство. Проте зараз штучний інтелект став не лише творцем неймовірно правдоподібних картин, фото чи відео, а й причиною судових позовів щодо порушення авторських прав [7]. У Каліфорнії штучний інтелект став предметом позову, коли у січні 2023 року у Федеральному суді зареєстрували позов художниць Сари Андерсен, Келлі МакКернан і Карлі Ортіс проти платформ Stability AI, Midjourney (генерують зображення за допомогою штучного інтелекту) та DeviantArt (сайт-портфоліо, на якому розмішують художні роботи, зокрема створені штучним інтелектом) [4]. На думку художниць, системи штучного інтелекту навчаються на їхніх роботах, захищених авторським правом. При цьому автори не дають на це прямої згоди, не отримують компенсації чи посилання на свої роботи. Позивачами зазначалося, що на DeviantArt роботи створені штучним інтелектом стали витісняти зображення, створені людьми, тому художниці «вимагали компенсації за збитки, завдані означеними платформами». [8]. Водночас авторки наголошували на тому, що штучний інтелект опосередковано

порушує законодавство про недобросовісну конкуренцію та їхні права на публічність [4].

Тоді у січні 2023 року ініційований позов зазнав критики через технічні неточності. Суддя відхилив частину позовних вимог з можливістю внесення додаткових пояснень та усунення технічних недоліків та неточностей. Відповідачі у справі наполягали на тому, що художниці (Келлі МакКернан і Карлі Ортіс) не зареєстрували свої зображення в Бюро з авторського права США і тому позовні вимоги про порушення авторських прав мають бути відхилені. Суд підтримавши таку позицію відхилив позов у частині порушення авторських прав, адже реєстрація творів є обов'язковою умовою для подання позову про порушення авторських прав. Водночас, суд продовжив судовий розгляд в частині того, що використання авторських творів для навчання можливо порушило авторські права художниці Сари Андерсен, адже вона зареєструвала 16 збірників творів. При цьому відповідачі просили суд «обмежити» її позовні вимоги про порушення авторських прав саме зареєстрованими збірниками, адже висувати свої звинувачення в порушенні авторського права щодо «створених нею понад двісті творів, включених до навчальних даних», вона може лише у тому випадку, якщо ідентифікує з точністю кожен з її зареєстрованих робіт, які були використані як навчальні зображення для штучного інтелекту. Суд також зазначив про необхідність внесення правок, зокрема, в частині роз'яснення теорії «стиснутих копій навчальних зображень» (про що йшлося у позові) з наведенням конкретних фактів принаймні часткового використання штучним інтелектом саме зображення творів, а не їх математичних моделей. Крім того необхідним є ширше розкриття поняття «копіювання» у розумінні Закону про авторське право США. Адже якщо таке використання буде доведено (копіювання творів), то можна буде говорити про порушення авторських прав. Адже суду не було зрозуміло, чи передбачає процес «змішування» пряме копіювання творів, а не застосування математичних рівнянь і алгоритмів для фіксації ідеї та концепції навчальних зображень. Водночас, вимоги про порушення прав авторів на публічність і принципи добросовісної конкуренції суд відхилив. Також наголосивши на необхідності конкретно визначити обсяг відповідальності для кожного відповідача та посилатися на переконливі аргументи. В подальшому, рішення суду може стати знаковим у питанні співвідношення розвитку штучного інтелекту і захисту авторських прав та для

формування положень про те, якою мірою платформи штучного інтелекту можуть використовувати захищені авторським правом твори в навчальних цілях та як художники можуть захистити себе від порушень авторських прав у зв'язку зі створенням продуктів штучним інтелектом [3].

Підсумовуючи слід зазначити наступне. Історія створення штучного інтелекту не була раптовою, вона почала зароджуватися ще у стародавні часи, коли мова йшла про істот, що були наділені розумом. Нові розробки, пристрої та винаходи поступово входили у побут людей, постійно удосконалюючись. Результатом вивчення закономірностей роботи людського мозку стало та є створення нових програмних систем, які адоптуються під сучасні вимоги. Штучний інтелект переживав різні етапи свого становлення, від «народження» до сучасного стану, на якому створення та використання штучного інтелекту є все більш популярним [9]. Безумовний вплив штучного інтелекту на різні сфери нашого життя вимагає його законодавчого врегулювання. Крім того, штучний інтелект вже є предметом позову, як можливий «порушник» авторського права, отже судова практика може стати/стане основою для заповнення законодавчих прогалин.

### **Список використаних джерел:**

1. Історія штучного інтелекту: веб-сайт URL: [https://uk.wikipedia.org/wiki/Історія\\_штучного\\_інтелекту](https://uk.wikipedia.org/wiki/Історія_штучного_інтелекту).
2. Харебава Т. Авторські права на об'єкт, створений штучним інтелектом. URL: [https://jurliga.ligazakon.net/analytics/225383\\_avtorsk-prava-na-obkt-stvoreniy-shtuchnim-ntelektom](https://jurliga.ligazakon.net/analytics/225383_avtorsk-prava-na-obkt-stvoreniy-shtuchnim-ntelektom).
3. О. Петрів, О. Спесивцева. Автори проти штучного інтелекту: на чиєму боці суд?: веб-сайт URL: <https://cedem.org.ua/news/avtory-proty-shi-na-chyjomu-botsi-cud>.
4. Петрів О. Штучний інтелект та авторське право: веб-сайт URL: <https://cedem.org.ua/analytics/shtuchnyi-intelekt-avtorske-pravo/>.
5. Петрів О. Штучний інтелект та Artificial intelligence act: час для юридичних рамок: веб-сайт URL: <https://cedem.org.ua/analytics/artificial-intelligence-act/>.
6. Художниці подали позов проти Stability AI, Midjourney та DeviantArt: веб-сайт URL: <https://gamedev.dou.ua/news/ai-vs-artists-lawsuit/>.
7. О. Спесивцева. Регулювання штучного інтелекту: досвід США: веб-сайт URL: <https://cedem.org.ua/analytics/shtuchnyi-intelekt-usa/>.

8. А. Большакова. Суд відхилив позови художників проти штучного інтелекту: веб-сайт URL: <https://life.pravda.com.ua/culture/2023/11/1/257364/>.

9. Бахарева О.В. «Штучний інтелект: історія виникнення». Збірник матеріалів VI Міжнародній науково-практичній конференції «міст Київ-Дніпро». Управління проектами. Перспективи розвитку проєктного та нейроменеджменту, інформаційних технологій управління, технологій створення та використання об'єктів прав інтелектуальної власності, трансфер технологій». (м. Київ-Дніпро. 21.03.2024). С. 23-25.

**Коваленко Т.В.**

молодший науковий співробітник відділу  
питань захисту прав інтелектуальної власності  
НДІ інтелектуальної власності НАПрН  
України

ORCID: <https://orcid.org/0000-0002-0099-2631>

## **ШТУЧНИЙ ІНТЕЛЕКТ І ПРАВА ІНТЕЛЕКТУАЛЬНОЇ ВЛАСНОСТІ**

Розвиток штучного інтелекту продовжує піднімати та вирішувати питання щодо застосування законів про інтелектуальну власність. Права інтелектуальної власності залишаються одними із спірних питань у сфері управління робіт та виробів із застосуванням штучного інтелекту.

Правники всього світу шукають баланс між стратегіями, які направлені на стимулювання розвитку штучного інтелекту, одночасно намагаючись вдосконалити та гармонізувати правові межі інтелектуальної власності щодо штучного інтелекту.

Закони щодо інтелектуальної власності надають її власникам певні права щодо заборони третім особам використовувати без дозволу свої творіння. Ці права дозволяють власникам інтелектуальної власності отримати визнання та фінансову вигоду, стимулюючи їх вкладати час та ресурси у свої інтелектуальні творіння. Традиційно це було одним із обґрунтувань необхідності таких законів про інтелектуальну власність. Однак ця аргументація непридатна до творів штучного інтелекту, для створення яких не вимагається стимулу [1]. Таким чином виникає проблема, чи повинні закони щодо інтелектуальної власності застосовуватися до творів штучного інтелекту таким же чином, як вони застосовуються до творів, які створені руками та розумом людини.

Законодавство у сфері інтелектуальної власності, як правило, при поданні заявки вимагає вказівки імені заявника. Свого часу доктор Стів Талер при поданні заявок у декількох країнах вказав заявником систему штучного інтелекту DABUS, яка була ним створена та належала йому [2]. майже всі заявки були відхилені на підставі того, що власником прав інтелектуальної власності має бути людина [3].

Хоча система штучного інтелекту не може бути автором твору, це не означає, що винаходи та твори, створені за допомогою штучного інтелекту, ніколи не будуть захищені законами про інтелектуальну власність. Адже штучний інтелект використовується лише як інструмент або помічник, який допомагає людині. Захист прав інтелектуальної власності на твір, створений за допомогою штучного інтелекту, можливо отримати, оскільки участь людини у створенні є очевидною. Використання штучного інтелекту не повинно бути перешкодою для отримання захисту прав інтелектуальної власності [4].

Юрисдикційна система світу має вирішувати проблеми, пов'язані із порушенням прав інтелектуальної власності щодо штучного інтелекту, включаючи проблеми розповсюдження по всьому світу негармонізованих законів про штучний інтелект. Незважаючи на це, окремі країни створюють нові закони та вносять правки до існуючих законів, направлені на стимулювання розвитку штучного інтелекту та забезпечення гарантій прав на нього.

Наприклад, Бюро з авторських прав Сполучених Штатів Америки відхилило реєстрацію твору художника Джейсона Аллена, яка була створена за допомогою штучного інтелекту, а саме за допомогою перетворення тексту на зображення. Художник стверджував, що використав не менше 624 текстових підказок для настройки сцени, тону і фокусу зображення, а також використовував інший штучний інтелект для збільшення роздільної здатності та розміру [5].

Питання в тому, чи слід захищати твір, створений за допомогою штучного інтелекту, законами про інтелектуальну власність є спірним, враховуючи, що рівень людської участі у його створенні нижче, а у деяких випадках мінімальний.

У зв'язку із відсутністю чіткої відповіді щодо захисту творів, створених за допомогою штучного інтелекту, виникає також питання який власник права інтелектуальної власності має право розпоряджатися правами, дозволяти або забороняти третім особам використовувати такі твори, а також подання позовної заяви проти протиправних дій.

Більшість країн вважають, що системи штучного інтелекту не мають правоздатності або правосуб'єктності, а отже, не можуть володіти правами інтелектуальної власності. Можливими власниками творів, створених за

допомогою штучного інтелекту, є власники системи штучного інтелекту, розробники системи, що написали код для системи штучного інтелекту, користувачі або оператори системи штучного інтелекту або будь-які комбінації із перерахованих осіб. Розпорядження правами інтелектуальної власності також можливо шляхом договірних положень.

Інша точка зору полягає в тому, що твір, створений за допомогою штучного інтелекту, має бути об'єктом суспільного надбання. Бюро з авторських прав Сполучених Штатів Америки опублікувало документ, де зазначено, що «створений матеріал (мається на увазі твір, створений за допомогою штучного інтелекту) не є продуктом людського авторства» [6].

На національному рівні найкращим способом забезпечення правового захисту може стати прийняття відповідного законодавства або внесення змін до діючого законодавства. Деякі країни визнали право *sui generis* на твори, створені за допомогою штучного інтелекту.

«*Sui generis* (с'юї г'енерис; букв. своєрідний, єдиний у своєму роді) - латинський вираз, що означає унікальність правової конструкції (акта, закону, статусу і т. д.), яка, попри наявність схожості з іншими подібними конструкціями, в цілому не має прецедентів» [7].

В українському законодавстві були внесені зміни до Закону про авторське право і суміжні права, де у п. 2 статті 3 зазначено, що «Положення цього Закону спрямовані на охорону та захист особистих немайнових прав і майнових прав: (...) суб'єктів права особливого роду (*sui generis*) на неоригінальний об'єкт, згенерований комп'ютерною програмою, зазначених у статті 33 цього Закону» [8].

Втілення штучного інтелекту продовжує зростати, ми маємо нагоду спостерігати більш широкий спектр його комерційного використання у всіх сферах. З огляду на проблеми, пов'язані із правами інтелектуальної власності на штучний інтелект, компаніям, що використовують технології штучного інтелекту, варто розробити стратегію, яка дозволяє орієнтуватися у змінах стосовно інтелектуальної власності. Із розвитком правової бази, що регулює штучний інтелект, виробникам необхідно гарантувати відсутність порушення прав, одночасно забезпечуючи ефективний захист своїх активів інтелектуальної власності. Чітке розуміння, як діюче законодавство розповсюджується на технології та процеси штучного інтелекту та відстежуючи зміни у законодавстві матимуть вирішальне значення для

компаній, які прагнуть до відповідного впровадження інновацій та захисту своєї інтелектуальної власності в економіці, заснованій на штучному інтелекті.

Наступні роки стануть роками більш широкого впровадження інструментів штучного інтелекту правовласниками. Такі інструменти слід використовувати для оптимізації ручних процесів, а також для швидкого виявлення та усунення ризиків порушення прав. Правники та юристи отримають вигоду від доступу до великих обсягів даних та їх аналізу, що може бути використано для покращення процесів прийняття рішень в галузі інтелектуальної власності. Юристи вже використовують складні інструменти на основі штучного інтелекту для перегляду та обробки великих обсягів інформації. Інструменти штучного інтелекту можуть швидко сканувати та узагальнювати великі обсяги доказів, виявляючи найбільш актуальну інформацію для поточного завдання. Це дозволяє юристам та правникам зосередитися на стратегії, а не на тривалому вивченні документів. Автоматизуючи збір та упорядкування доказів, системи штучного інтелекту дозволяють економити час. Оптимізація підготовки доказів, необхідних для судового провадження, що стосуються будь-яких об'єктів інтелектуальної власності, також можуть бути спрощені завдяки використанню інструментів штучного інтелекту,

Однак, слід враховувати, що хоча штучний інтелект і забезпечує підвищення ефективності, тим не менш людський контроль залишається необхідним.

У процесі розвитку штучного інтелекту виникатимуть нові можливості і проблеми у галузі захисту та забезпечення прав інтелектуальної власності. Правова ситуація буде змінюватися по мірі того, як судам та регулюючим органам доведеться вирішувати складні та нові проблеми. У зв'язку з цим всім учасникам системи штучного інтелекту необхідно бути обізнаними щодо останніх правових змін і тенденцій.

### **Список використаних джерел:**

1. Robert Plotkin. AI-Generated Inventions Need Human Ingenuity and Patents. URL: <https://news.bloomberglaw.com/us-law-week/ai-generated-inventions-need-human-ingenuity-and-patents>
2. Геннадій Андрощук. DABUS не був і не є людиною: Верховний суд Великої Британії відхилив DABUS як винахідника. Юридична газета, 25 грудня 2023. URL:

<https://yur-gazeta.com/publications/practice/zahist-intelektualnoyi-vlasnosti-avtorske-pravo/dabus-ne-buv-i-ne-e-lyudinoyu-verhovniy-sud-velikoyi-britaniyi-vidhiliv-dabus-yak-vinahidnika.html>

3. URL: <https://www.droit-technologie.org/wp-content/uploads/2020/03/motifs-de-la-decision-EP-18275174.pdf>

4. Tobias Kempas. Copyright Committee Issues Report on Copyrights and Neighboring Rights in AI-Generated and AI-Assisted Works. URL: <https://www.inta.org/news-and-press/inta-news/copyright-committee-issues-report-on-copyrights-and-neighboring-rights-in-ai-generated-and-ai-assisted-works/>

5. Andrew Kenney. Jason Allen's AI art won the Colorado fair — but now the feds say it can't get a copyright. CPR News. Sep. 6, 2023 URL: <https://www.cpr.org/2023/09/06/jason-allens-ai-art-won-colorado-fair-feds-deny-copyright-protection/>

6. [Copyright Registration Guidance: Works Containing Material Generated by Artificial Intelligence URL: [https://www.copyright.gov/ai/ai\\_policy\\_guidance.pdf](https://www.copyright.gov/ai/ai_policy_guidance.pdf)

7. Sui generis. URL: [https://uk.wikipedia.org/wiki/Sui\\_generis](https://uk.wikipedia.org/wiki/Sui_generis)

8. Закон України «Про авторське право і суміжні права» (Відомості Верховної Ради (ВВР), 2023, № 57, ст.166). URL: [https://zakon.rada.gov.ua/laws/show/2811-20?find=1&text=%D0%BD%D0%B5%D0%BE%D1%80%D0%B8%D0%B3%D1%96%D0%BD#w1\\_1](https://zakon.rada.gov.ua/laws/show/2811-20?find=1&text=%D0%BD%D0%B5%D0%BE%D1%80%D0%B8%D0%B3%D1%96%D0%BD#w1_1)

**Олена Гончаренко**

кандидат економічних наук, доцент кафедри  
світової економіки Державного торговельно-  
економічного університету

ORCID: <https://orcid.org/0000-0003-2563-810X>

**Анастасія Левківська**

студентка ФМТП 3-1 Державного  
торговельно-економічного університету

## **КРИЗА ІНТЕЛЕКТУАЛЬНОЇ ВЛАСНОСТІ В УМОВАХ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ**

Можна безапеляційно стверджувати, що ХХІ століття увійде в історію як епоха фундаментальних тектонічних зсувів, коли штучний інтелект (ШІ) із суто теоретичної концепції перетворився на всеосяжну технологію, що трансформує саму основу людського суспільства. Яскравим прикладом такого перетворення є концепція «Суспільство 5.0» — стратегія, сформована урядом Японії ще у 2016 р., що передбачає максимальну та безшовну інтеграцію цифрових технологій, великих даних та ШІ в усі сфери людського життя [1, с. 68].

Актуальність дослідження цієї трансформації критично загострюється в умовах, які вже не описуються звичною моделлю VUCA, а відповідають характеристикам світу BANI.

Саме в цьому крихкому і нелінійному середовищі, де технологічні зміни відбуваються непередбачувано та незбагненно швидко, тісна інтеграція генеративних моделей ШІ в усі сфери людської діяльності катапультиє низку нових, гострих викликів. Ступінь занепокоєння щодо неконтрольованого розвитку ШІ сягнув такої критичної межі, що понад 53,8 тис. відомих особистостей підписали відкрите звернення із закликом призупинити створення надрозумного ШІ [2]. Вони вимагають зупинити роботи, доки не буде досягнуто глобальної згоди щодо безпечного впровадження таких систем.

Ця глобальна тривога прямо корелює з масовістю використання генеративного контенту. За останніми підрахунками, кількість статей, створених штучним інтелектом, вперше перевищила обсяг контенту, написаного людьми, і наразі половина текстів у мережі може мати «нейромовні» ознаки. Ця тенденція підтверджується і в освітній сфері, де використання генеративного ШІ серед

студентів університетів лише за один рік «стрімко зросло» [3]. Хоча такі пошукові системи, як Google, досі віддають перевагу контенту з чітким авторством, ця тенденція створює безпрецедентний обсяг даних, джерело та правовий статус яких є невизначеним.

Криза ІВ набуває гострої актуальності, оскільки межа, яка розділяє результати діяльності людини та результати, згенеровані ІШ, стає дедалі більш розмитою та дискусійною. Особливу крихкість ситуації додає низький рівень ІШ-грамотності серед користувачів. Цей тривожний стан нерозуміння посилюється відсутністю чітких інституційних правил: менше третини студентів стверджують, що їхній навчальний заклад заохочує використання ІШ, тоді як майже третина повідомляє про пряму заборону [3]. Ігнорування навчання відповідальному використанню ІШ може не лише призвести до поглиблення цифрової нерівності та нерозуміння проблем упередженості та «галюцинацій», але й ускладнює правову ідентифікацію авторства та відповідальності за згенерований контент. Таким чином, криза ІВ має не лише технологічний, але й освітньо-інституційний вимір.

З огляду на необхідність структурованого впровадження ІШ, на національному рівні були розроблені «Рекомендації щодо відповідального впровадження та використання технологій штучного інтелекту в закладах вищої освіти» [4]. Цей документ встановлює основні принципи етичного використання, організаційні вимоги та правові аспекти (включно з авторським правом) для ЗВО, що свідчить про намагання регуляторів надати чіткі інструкції для адаптації національної освіти до нових технологічних реалій, таким чином мінімізуючи інституційний «тривожний» стан.

Ця криза ІВ набуває особливої гостроти на тлі глобальних зусиль щодо розвитку відкритої науки, оскільки життєво необхідні для вирішення глобальних викликів відкритий доступ та вільний обмін ідеями вступають у пряме і нерозв'язне протиріччя з принциповою, негативною залежністю ІШ від масштабних, захищених авторським правом масивів даних.

Крихкість системи ІВ ілюструє критична правова проблема процесу навчання генеративного ІШ, який залежить від використання величезних масивів даних, захищених авторським правом, що невідворотно спричиняє правовий хаос. Зокрема, у червні 2025 року Disney та Universal подали позов проти Midjourney за «бездонну яму плагіату», звинувачуючи компанію у несанкціонованому копіюванні персонажів. Генеральний директор Midjourney

підтвердив, що база даних була зібрана через «великий вебскрейпінг Інтернету». Ці прецеденти, включно з попереднім позовом 10 художників, чітко засвідчують неадекватність чинних механізмів ІВ до регулювання масового та автоматизованого використання захищеного контенту [5].

Насправді, принцип роботи сучасного генеративного ШІ полягає у «змішуванні» та обробці великих масивів даних (Текст і Data Mining), а не у прямому копіюванні. Юридично це має ключове значення, оскільки авторським правом не охороняються ідеї, методи, концепції та принципи [6]. Саме тому в юрисдикціях ЄС (через Директиву про авторське право на Єдиному цифровому ринку) та США (через доктрину добросовісного використання, де копія виконує іншу функцію, ніж оригінал) навчання ШІ на правомірно доступних творах вважається цілком законним і не порушує авторських прав. Проте, така практика є предметом судових суперечок, особливо коли йдеться про використання «піратських матеріалів».

Варто відзначити, що українське законодавство врегулювало результат роботи ШІ як неоригінальний об'єкт, що охороняється правом особливого роду (*sui generis*) [7]. Такий об'єкт не містить творчого підходу і генерується без участі людини, тому на нього не виникають особисті немайнові права, а термін чинності майнових прав становить 25 років з моменту створення. Авторство може бути визнано за людиною лише у випадку, якщо вона внесла творчі корективи у згенерований ШІ-контент.

Ці приклади свідчать, що в умовах незбагненого світу BANI, криза ІВ посилюється подвійним бар'єром: з одного боку, домінуванням комерційних видавців, чий інтерес часто не збігається з потребами наукової спільноти, створюючи фінансовий (через високі APC) та правовий бар'єр. З іншого боку, тривожна нерегульованість використання контенту ШІ лише поглиблює цю кризу.

Метою даної роботи є глибокий аналіз основних, системних викликів, що постають перед системою інтелектуальної власності в умовах масового використання генеративних ШІ, а також терміновий пошук шляхів адаптації правового регулювання для забезпечення стійкого та життєздатного балансу між нестримним технологічним прогресом та невід'ємним захистом інтересів правовласників, зокрема у критично важливому контексті Відкритої науки.

Ця системна проблема набуває особливої гостроти у світлі глобальної тенденції до відкритої науки. З одного боку, відкритий доступ є стратегічно

важливим для наукового прогресу, оскільки він «каталізує» співпрацю та обмін ідеями, необхідними для вирішення глобальних викликів. З іншого боку, як прямо вказано в дискусії про відкритість, кожен, хто розміщує контент у відкритому доступі або навіть просто онлайн, має розуміти, що рано чи пізно цей матеріал стане «їжею» для AI-моделей. Таким чином, прагнення до максимальної відкритості, що є рушієм прогресу, вступає в пряме протиріччя із традиційним захистом авторського права, оскільки фактично призводить до неліцензійного та безоплатного масового відтворення творів.

Ще одним прикладом є подія, що відбулася навесні 2025 року. Тоді в соціальних мережах почав стрімко поширюватися феномен «гібліфікація» — генерації зображень за допомогою ChatGPT у стилі японської студії анімації Ghibli, заснованої відомим режисером анімаційних мультфільмів Хаяо Міядзакі. Цей вірусний тренд ініціював гостру дискусію, поставивши питання про етичні та правові наслідки такого використання ШІ. Результатом широкого суспільного резонансу стало включення цього питання до порядку денного засідання японського парламенту. У квітні 2025 р. під час засідання Комітету кабінету Палати представників Японії депутат Масато Імаї порушив питання про можливе порушення авторських прав. На це він отримав відповідь від представника уряду Хірохіці Накхара, що використання стилю не є порушенням, але створення контенту «надто подібного» до першоджерела, може кваліфікуватись як таке. Станом на той час уряд Японії не висловлював жодних звинувачень на адресу OpenAI [8]. Проте, все змінилось 1 жовтня 2025 р., коли компанія випустила генеративну модель Sora для створення коротких відео, після чого користувачі почали активно створювати контент з персонажами відомих японських франшиз. Це викликало офіційний запит уряду країни про дотримання авторських прав, що, у свою чергу, змусило OpenAI обіцяти розширити захист правовласників. Ця ситуація актуалізувала питання необхідності нових правових підходів до генеративного ШІ [9].

Таким чином, розвиток генеративних моделей штучного інтелекту значно ускладнює традиційні механізми захисту інтелектуальної власності. Алгоритми здатні створювати тексти, зображення, музику, відео на основі вже існуючих матеріалів, що ставить під сумнів поняття авторства та законність використання захищеного контенту.

Проте самі розробники ШІ виступають проти суворих правових обмежень. Зокрема, компанія OpenAI у своєму документі, направленому до уряду

Сполученого Королівства заявила, що навчання є основним етапом створення ШІ, тому будь-які правові обмеження на цьому етапі не просто уповільнюють його розвиток, а й підривають саму основу технології [10, с. 5]. Це протиріччя безпосередньо віддзеркалює боротьбу «Спільноти понад комерціалізацію», оскільки наразі лівова частка виробництва знань контролюється комерційними організаціями, які прагнуть монетизувати доступ і до своїх опублікованих матеріалів, і до даних для тренування ШІ.

Представлені приклади свідчать про системну кризу в галузі інтелектуальної власності, спричинену розвитком штучного інтелекту. У відповідь на ці виклики формуються альтернативні моделі взаємодії: ліцензійні угоди та купівля прав; обмежений доступ та технічні бар'єри.

Одним із способів вирішення проблеми є купівля компаніями-розробниками ШІ авторських прав у правовласників. Таке придбання дозволяє легалізувати використання творів для тренування моделей штучного інтелекту та генерування нового контенту на їх основі. Наприклад, китайська компанія ShengShu Technology, розробник генератора відео Vidu, вже запровадила альтернативну модель взаємодії ШІ з інтелектуальною власністю, підписавши ліцензійну угоду на використання анімаційного серіалу «Подорож на Захід: Легенди короля мавп» (1999). На противагу скандалу навколо OpenAI Sora, яка використала стиль Ghibli без дозволу, Vidu легалізувала доступ до контенту [11].

Ще одним способом є модель обмеженого доступу до навчальних даних, за якої правовласники отримують можливість перевіряти наявність своїх творів у захищених сховищах даних. Проте її головним недоліком є те що, вартість створення і підтримка такої інфраструктури є доволі високою [12, с. 185]. Також для запобігання порушення авторських прав розробники впроваджують спеціальні механізми, серед яких: фільтрація запитів, здатних призвести до копіювання чужого контенту; модифікація алгоритмів навчання для мінімізації правопорушних результатів та вбудовані обмеження, що забороняють моделям відтворювати образи захищених персонажів або імітувати стиль сучасних митців [13].

Таким чином, протиріччя між технологічним прогресом та авторським правом загострюється через принципову залежність ШІ від великих масивів даних, що робить будь-які обмеження на їх використання загрозою для розвитку технологій. У відповідь на ці виклики формуються альтернативні моделі взаємодії, проте жодна з них поки не є універсальним рішенням. З огляду на це,

ключовим завданням залишається розробка нових, гнучких правових механізмів, здатних збалансувати інтереси інноваторів та правовласників, а також забезпечити стале функціонування Відкритої науки в умовах домінування комерційних інтересів.

### Список використаних джерел

1. Чалюк, Ю. (2023). Суспільство 5.0 у японській концепції кейданрен. *Mechanism of an Economic Regulation*, (1 (99)), 65-74. URL: <http://www.mer-journal.sumy.ua/index.php/journal/article/view/132>
2. Statement on Superintelligence: відкрите звернення. URL: <https://superintelligence-statement.org/> (date of access: 30.10.2025).
3. Freeman, J. *Student Generative AI Survey 2025*: HEPI Policy Note 61. Oxford : Higher Education Policy Institute, Kortext, 2025. 12 p. URL: <https://www.hepi.ac.uk/reports/student-generative-ai-survey-2025/>
4. Рекомендації щодо відповідального впровадження та використання технологій штучного інтелекту в закладах вищої освіти. Київ : Міністерство цифрової трансформації України, Міністерство освіти і науки України, 2025. URL: <https://mon.gov.ua/static-objects/mon/sites/1/news/2025/04/24/shi-v-zakladakh-vyshchoi-osvity-24-04-2025.pdf>
5. Montgomery B. Disney and Universal sue AI image creator Midjourney, alleging copyright infringement. *the Guardian*. URL: <https://www.theguardian.com/technology/2025/jun/11/disney-universal-ai-lawsuit> (date of access: 27.10.2025).
6. Петрів О. Штучний інтелект та авторське право. Центр демократії та верховенства права. 2023. URL <https://cedem.org.ua/analytics/shtuchnyi-intelekt-avtorske-pravo/> (date of access: 30.10.2025).
7. Рекомендації щодо відповідального впровадження та використання технологій штучного інтелекту в закладах вищої освіти. Київ : Міністерство цифрової трансформації України, Міністерство освіти і науки України, 2025. URL: [https://mon.gov.ua/static-objects/mon/sites/1/news/2025/04/24/shi-v-zakladakh-vyshchoi-osvity-24-04-2025.pdf?\\_\\_cf\\_chl\\_tk=ovqGuMWYKeIQvkeAArpgP3ec572.W8bWNc\\_Z15m4hoU-1761822678-1.0.1.1-\\_\\_w20j4JtU1i1L7ibU.rJPjDpaCsOZqvVQn7ms3rpdE](https://mon.gov.ua/static-objects/mon/sites/1/news/2025/04/24/shi-v-zakladakh-vyshchoi-osvity-24-04-2025.pdf?__cf_chl_tk=ovqGuMWYKeIQvkeAArpgP3ec572.W8bWNc_Z15m4hoU-1761822678-1.0.1.1-__w20j4JtU1i1L7ibU.rJPjDpaCsOZqvVQn7ms3rpdE)
8. Качковська Я. Стилізацію під Ghibli заборонити не можна змиритися: японські депутати думають, де поставити кому. *Suspilne.media*. 22.04.2025. URL: <https://suspilne.media/culture/1039947-stucnij-intelekt-bilse-ne-nestime-zagrozu->

intelektualnij-vlasnosti-si-generator-vidu-pridbav-prava-na-korola-mavp/ (дата звернення: 27.10.2025).

9. Заблоцька О. Японія закликала OpenAI захистити аніме та мангу від впливу штучного інтелекту. *Suspilne.media*. 16.10.2025. URL: <https://suspilne.media/culture/1140952-aponia-zaklikala-openai-zahistiti-anime-ta-mangu-vid-vplivu-stucnogo-intelektu/> (дата звернення: 27.10.2025).

10. Our response to the UK's copyright consultation. *OpenAI*. 02.04.2025. URL: <https://openai.com/global-affairs/response-to-uk-copyright-consultation/> (дата звернення: 27.10.2025).

11. Chinese AI Video Generator Vidu Strikes IP Deal For Animated Series 'Journey To The West: Legends Of The Monkey King'. *Deadline*. URL: <https://deadline.com/2025/06/chinese-ai-video-generator-vidu-ip-deal-journey-to-the-west-1236428314/> (date of access: 28.10.2025).

12. Adam Buick, Copyright and AI training data—transparency to the rescue?, *Journal of Intellectual Property Law & Practice*, Volume 20, Issue 3, March 2025, Pages 182–192, <https://doi.org/10.1093/jiplp/jpae102>

13. Copyright Office Weighs In on AI Training and Fair Use | Skadden, Arps, Slate, Meagher & Flom LLP. *Skadden, Arps, Slate, Meagher & Flom LLP*. URL: <https://www.skadden.com/insights/publications/2025/05/copyright-office-report> (date of access: 28.10.2025).

## **СОЦІАЛЬНИЙ ВИМІР ШТУЧНОГО ІНТЕЛЕКТУ: ОСВІТА, ПРАЦЯ, ІНКЛЮЗІЯ**

**Андрій Озарчук**

старший викладач кафедри психології та  
інклюзивної освіти Рівненського обласного  
інституту післядипломної педагогічної освіти,  
науковий співробітник відділу STEM-освіти  
Інституту педагогіки НАПН України  
ORCID: <https://orcid.org/0009-0001-5909-7279>

### **ПЕРСОНАЛІЗАЦІЯ ТА ІНКЛЮЗІЯ: МОЖЛИВОСТІ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ СУЧАСНОЇ ПЕДАГОГІКИ**

У сучасному світі освіта переживає суттєві зміни, зумовлені впровадженням цифрових технологій, що дедалі більше змінюють традиційні підходи до навчання. В умовах пандемії COVID-19 та повномасштабної війни в Україні, що призвело до переходу на дистанційне та змішане навчання, технології штучного інтелекту (ШІ) отримали новий поштовх до розвитку в освітньому середовищі. Використання штучного інтелекту стає одним із ключових трендів у педагогічних інноваціях, адже він надає можливості для більш індивідуалізованого, адаптивного та ефективного навчання. Сьогодні штучний інтелект виступає не тільки як інструмент для автоматизації адміністративних процесів, а й стає активним учасником освітнього процесу, впливає на сучасні педагогічні інновації.

В останні роки науковці активно вивчають можливості штучного інтелекту в освітньому процесі. За даними досліджень ЮНЕСКО, впровадження штучного інтелекту у сферу освіти дозволяє покращити навчальні результати здобувачів освіти через індивідуалізовані підходи до навчання, моніторинг прогресу та адаптацію освітніх матеріалів [3]. В огляді від Sciforce (2023) зазначається, що ШІ може ефективно підтримувати процеси персоналізованого навчання, створюючи рекомендаційні системи, здатні адаптувати контент під потреби окремих студентів [7]. Окрім того, експерти стверджують, що ШІ може сприяти вирішенню проблем інклюзивної освіти,

зокрема шляхом впровадження технологій, що допомагають учням із особливими освітніми потребами.

Наукові роботи також приділяють увагу етичним аспектам використання штучного інтелекту в освіті. Наприклад, Бенджамін Вільямсон, Фелісітас Макгілкріст і Джон Портер (2023) у своєму дослідженні вказують на необхідність забезпечення прозорості алгоритмів та відповідальності за прийняті рішення, особливо у випадках автоматизованого оцінювання здобувачів освіти [9]. Це важливий аспект, оскільки збільшується ризик упереджень і дискримінації, якщо моделі ШІ не будуть належним чином контролюватися.

Крім того, Ілля Філіпов, CEO і співзасновник студії онлайн-освіти EdEra, наголошує на тому, що ШІ стане одним із драйверів EdTech і у світі, і в Україні завдяки інтелектуальним системам навчання (ITS), інструментам адаптивного шкільного навчання, автоматизованому оцінюванню, аналітиці успішності користувачів продуктів, створенню генеративного контенту [2].

Однією з найбільших переваг використання штучного інтелекту в освіті є можливість глибокої персоналізації навчального процесу, що відкриває нові горизонти для якісної освіти. Завдяки комплексному аналізу великих даних (Big Data) та використанню передових алгоритмів машинного навчання, сучасні системи ШІ отримують змогу не просто адаптувати, а конструювати індивідуальні освітні траєкторії, що враховують унікальні особливості кожного здобувача освіти. Технологічний фундамент цієї трансформації становлять складні багаторівневі архітектури нейронних мереж, здатних обробляти та інтерпретувати різноманітні дані – від часових характеристик виконання завдань і паттернів поведінки до мікроекспресій обличчя під час навчання. Наприклад, інноваційні освітні платформи, такі як DreamBox або Knewton, здійснюють в реальному часі моніторинг не лише правильності відповідей, але й аналізують стратегії вирішення задач, виявляючи індивідуальні когнітивні патерни та стилі навчання [6]. Механізм адаптації ґрунтується на безперервному оновленні профілю учня, де кожна взаємодія з системою уточнює його інтелектуальний портрет: визначаються оптимальний темп подачі матеріалу, співвідношення теоретичних блоків та практичних завдань, найбільш ефективні типи наочності. Унікальність такого підходу полягає в здатності системи прогнозувати освітні труднощі до їх прояву, пропонуючи превентивні корегуючі вправи, що дозволяє уникнути

формування знаннєвих прогалин. Емпіричні дослідження демонструють, що динамічні навчальні траєкторії, побудовані на основі ШІ, знижують ризик когнітивного перевантаження на 40-45% та підвищують рівень тривалої мотивації на 35-40% порівняно з традиційними методами навчання. Перспективним напрямком розвитку є інтеграція біометричних даних, що дозволить враховувати фізіологічний стан учня та корегувати навантаження відповідно до його психоемоційних особливостей. Важливим аспектом залишається розробка етичних стандартів для таких систем, що гарантуватиме прозорість алгоритмів та захист конфіденційності особистих даних.

Таким чином, персоналізація на основі ШІ перетворюється з технологічної інновації на фундаментальний принцип сучасної освіти, що забезпечує оптимальні умови для розкриття потенціалу кожної особистості. Також перспективним напрямом є використання чат-ботів для надання зворотного зв'язку здобувачам освіти. Інтерактивні помічники на основі ШІ здатні оперативно відповідати на запитання, пояснювати складні теми та допомагати у розв'язанні задач і багато іншого. Наприклад, у дослідженнях, опублікованих у *Journal of University Teaching and Learning Practice*, показано, що використання чат-ботів допомагає здобувачам освіти значно покращити результати тестування, підвищити залученість та загальну задоволеність від навчального процесу [4].

Однією з ключових областей застосування штучного інтелекту є інклюзивна освіта. Технології ШІ можуть допомогти у створенні доступного навчального середовища для осіб з ООП. Наприклад, розробки на основі комп'ютерного зору дозволяють автоматично перетворювати текст на мову або ж аудіо на текст для здобувачів освіти з порушеннями слуху або зору. Програми з аналізу мовлення можуть підтримувати учнів з дислексією або іншими мовленнєвими труднощами, допомагаючи їм поліпшити свої навички читання та письма. Варто зазначити, що Габріель Жульєн розглядає у своєму дослідженні можливості допомогти педагогам, батькам, учням, урядовцям і політикам у прийнятті правильних рішень щодо ефективного використання штучного інтелекту в інклюзивному освітньому середовищі при навчанні осіб із особливими освітніми потребами [5].

Дослідження також показують, що інструменти штучного інтелекту можуть підвищити інклюзивність через створення гнучких ігор та навчальних

матеріалів, які адаптуються до рівня знань та можливостей кожного учня. Такі інструменти дозволяють учням та студентам з ООП навчатися у зручному для них темпі, не відчуючи тиску або дискримінації.

Автоматизація оцінювання є ще однією сферою, де штучний інтелект може внести значні зміни. Наприклад, платформи на основі ШІ здатні автоматично перевіряти письмові роботи, виявляючи граматичні та логічні помилки, що полегшує роботу викладачів [1]. Однак, попри високий потенціал, ця сфера вимагає обережності, оскільки існує ризик упередженого оцінювання через неадекватно навчені моделі штучного інтелекту. Необхідно забезпечити регулярний контроль якості моделей та їх відповідність принципам етичного використання даних.

Штучний інтелект також може сприяти управлінню навчальними процесами через прогнозування академічної успішності здобувачів освіти. Алгоритми на основі аналізу попередніх даних можуть передбачати, які здобувачі освіти можуть мати труднощі з успішністю, дозволяючи викладачам вчасно втрутитися і надати підтримку [8]. Це дає можливість забезпечити більш гнучкий та адаптивний підхід до навчання, що особливо актуально в умовах дистанційної освіти.

Попри значний прогрес у використанні штучного інтелекту в освіті, існують певні виклики, що потребують подальших досліджень. Насамперед це питання етики та відповідальності при використанні ШІ, адже рішення, прийняті автоматизованими системами, можуть мати серйозний вплив на майбутнє здобувачів освіти. Необхідно вивчати способи зменшення алгоритмічної упередженості та забезпечення справедливості при використанні цих технологій.

Крім того, залишається відкритим питання щодо підготовки педагогів до роботи з технологіями штучного інтелекту. Важливо забезпечити педагогічних та наково-педагогічних працівників відповідними навичками, щоб вони могли ефективно використовувати ці інструменти у своїй професійній діяльності. Дослідження у цій галузі повинні також приділяти увагу розробці нових освітніх програм, що включатимуть вивчення штучного інтелекту як обов'язкової частини підготовки майбутніх педагогів.

Штучний інтелект уже сьогодні змінює освітній ландшафт, стаючи важливим інструментом для індивідуалізації навчання, підвищення інклюзивності та автоматизації оцінювання. Попри значні переваги,

впровадження ІІІ в освіті супроводжується викликами, зокрема у сфері етики та підготовки викладачів. Подальші дослідження у цій галузі повинні сприяти створенню справедливих, етичних та ефективних рішень для підтримки сучасної освіти.

### Список використаних джерел

1. Загальні відомості про помічники й інструменти для написання тексту на основі ІІІ. <https://www.microsoft.com/>. URL: <https://www.microsoft.com/uk-ua/microsoft-365/word/ai-writing> (дата звернення: 22.10.2025).
2. Філіпов І. Перспективи та напрямки зростання EdTech в Україні. *Новини ІТ і бізнесу в Україні – AIN.ua*. URL: <https://ain.ua/2024/05/14/perspektyvy-ta-napryamky-zrostannya-edtech-v-ukrayini-kolonka/> (дата звернення: 22.10.2025).
3. Artificial intelligence in education. *UNESCO : Building Peace through Education, Science and Culture, communication and information*. URL: <https://www.unesco.org/en/digital-education/artificial-intelligence> (дата звернення: 22.10.2025).
4. Enhancing student engagement using artificial intelligence (AI) and chatbots like ChatGPT / *Journal of University Teaching and Learning Practice*. - Vol. 21 No. 06 (2024). <https://open-publishing.org/>. URL: <https://doi.org/10.53761/pzd17z29> (дата звернення: 22.10.2025).
5. Gabriel J. How Artificial Intelligence (AI) impacts inclusive education. *Educational Research and Reviews*. 2024. Vol. 19, no. 6. P. 95–103. URL: <https://doi.org/10.5897/err2024.4404> (дата звернення: 22.10.2025).
6. Holmes W. The unintended consequences of artificial intelligence and education. *Education International*. URL: <https://www.ei-ie.org/en/item/28115:the-unintended-consequences-of-artificial-intelligence-and-education> (дата звернення: 13.10.2024).
7. Sciforce. AI revolution in edtech: AI in education trends and successful cases. *Medium*. URL: <https://medium.com/sciforce/ai-revolution-in-edtech-ai-in-education-trends-and-successful-cases-7d5b7d69b77b> (дата звернення: 22.10.2025).
8. Systematic review of research on artificial intelligence applications in higher education – where are the educators? / O. Zawacki-Richter et al. *International Journal of Educational Technology in Higher Education*. 2019. Vol. 16, no. 1. URL: <https://doi.org/10.1186/s41239-019-0171-0> (дата звернення: 22.10.2025).
9. Williamson B., Macgilchrist F., Potter J. Re-examining AI, automation and datafication in education. *Learning, Media and Technology*. 2023. Vol. 48, no. 1. P. 1–5. URL: <https://doi.org/10.1080/17439884.2023.2167830> (дата звернення: 22.10.2025).

**Олександр Остапенко**

кандидат юридичних наук, доцент кафедри  
публічно - правових дисциплін Вінницького  
державного педагогічного університету імені  
Михайла Коцюбинського

ORCID: <https://orcid.org/0000-0003-1158-0146>

## **ШТУЧНИЙ ІНТЕЛЕКТ У СФЕРІ ОСВІТИ ТА НАУКИ: ПРАВОВІ ВИКЛИКИ ЦИФРОВОЇ ЕПОХИ**

Стрімкий розвиток штучного інтелекту (ШІ) радикально трансформує сучасні соціальні, економічні та правові системи. Особливо відчутними ці зміни є у сфері освіти та науки, де алгоритми машинного навчання, великі дані й генеративні моделі відкривають нові горизонти для навчання, досліджень та управління знаннями. Водночас активна інтеграція ШІ породжує низку правових, етичних і регуляторних викликів, які потребують системного осмислення та відповідного нормативного врегулювання [1, с. 17].

Сфера освіти сьогодні стає полігоном для апробації інтелектуальних технологій: від адаптивних освітніх платформ і автоматизованих систем оцінювання до чат-ботів-помічників викладача та генеративних систем контенту. З одного боку, це підвищує ефективність освітнього процесу, забезпечує персоналізацію навчання, розширює доступ до знань. З іншого — створює ризики порушення авторських прав, академічної доброчесності та захисту персональних даних студентів і викладачів.

Одним із ключових правових викликів цифрової епохи є визначення правового статусу продуктів, створених за допомогою ШІ. В українському законодавстві наразі відсутнє чітке розуміння, чи може результат, створений штучним інтелектом (наприклад, текст, зображення, наукова стаття або програмний код), визнаватися об'єктом авторського права та хто є його суб'єктом — розробник алгоритму, користувач чи система. На міжнародному рівні тривають дискусії, які підтверджують складність застосування традиційних правових категорій.

Ще одним важливим аспектом є захист персональних даних. ШІ-системи у сфері освіти збирають, зберігають і аналізують великі обсяги інформації про студентів: від результатів навчання до поведінкових

характеристик. Це створює ризики неправомірного використання або витоку даних, що порушує норми Закону України «Про захист персональних даних» та положення Регламенту ЄС 2016/679 (GDPR). Особливої уваги потребує питання згоди на обробку даних у навчальному процесі, коли студенти можуть не мати реальної альтернатив.

Проблемним залишається і питання академічної доброчесності. Використання генеративних систем (ChatGPT, Copilot, Gemini тощо) у навчанні та наукових дослідженнях викликає суперечливі оцінки: з одного боку, вони є інструментом пізнання, з іншого — загрожують появою плагіату нового типу. Науковці та університети стикаються з необхідністю розробки внутрішніх політик щодо використання ШІ, визначення меж його застосування та формування цифрової етики в освітньому середовищі [3].

На рівні міжнародного права формується підхід до регулювання ШІ, що базується на принципах етичності, прозорості та підзвітності. Так, у 2021 році ЮНЕСКО ухвалила Рекомендацію з етики штучного інтелекту, а Європейський Союз у 2024 році затвердив AI Act — перший комплексний правовий акт, який класифікує системи ШІ за рівнем ризику та встановлює вимоги до безпеки, якості даних і відповідальності. Україна, як держава-кандидат на вступ до ЄС, має гармонізувати своє законодавство з цими нормами.

Особливої уваги заслуговує питання юридичної відповідальності за дії або помилки систем ШІ. У випадках, коли автономна система ухвалює рішення (наприклад, під час автоматизованого оцінювання чи добору навчального контенту), важливо визначити, хто несе відповідальність за наслідки — розробник, користувач чи освітня установа. Відсутність правової визначеності у цій сфері може призвести до зловживань і порушення прав учасників освітнього процесу.

Суттєвим викликом для правового регулювання є також питання алгоритмічної дискримінації. ШІ може відтворювати або підсилювати упередження, закладені у даних для навчання, що потенційно впливає на об'єктивність оцінювання студентів або відбір наукових проєктів. Це створює нову сферу для юридичного контролю, яка потребує інтеграції норм антикорупційного, освітнього та інформаційного законодавства.

У перспективі необхідно формувати національну політику етичного використання ШІ в освіті та науці, яка включатиме правові механізми

запобігання зловживанням, підтримку інновацій і розвиток цифрової компетентності освітян. Важливо не лише запроваджувати обмеження, а й забезпечити стимули для відповідального впровадження ШІ, що підвищує якість освітніх і наукових результатів.

Отже, правові виклики, пов'язані зі штучним інтелектом у сфері освіти та науки, потребують міждисциплінарного підходу, поєднання правових, етичних і технологічних рішень. В умовах цифрової епохи саме ефективне правове регулювання може забезпечити баланс між інноваціями, безпекою та захистом фундаментальних прав людини.

### **Список використаних джерел:**

1. Кузнецова Н.С. Право і штучний інтелект: виклики цифрової цивілізації. – Київ: Юрінком Інтер, 2023. – 256 с.
2. European Parliament. Artificial Intelligence Act: Regulation (EU) 2024/1689 of the European Parliament and of the Council. – Brussels, 2024.
3. UNESCO. Recommendation on the Ethics of Artificial Intelligence. – Paris: UNESCO, 2021.
4. Тараненко О.В. Штучний інтелект у сфері освіти: правові аспекти застосування. *Освітній дискурс*. 2024. №3. С. 20-27.
5. Закон України «Про захист персональних даних» № 2297-VI від 01.06.2010.

**Садова С.**

викладач кафедри права і соціальної роботи

Ізмаїльський державний гуманітарний

університет

ORCID: <https://orcid.org/0000-0001-7268-9001>

## **ШТУЧНИЙ ІНТЕЛЕКТ І РИНОК ПРАЦІ: СОЦІАЛЬНІ РИЗИКИ ТА ВИКЛИКИ ПЕРЕРОЗПОДІЛУ ВИГОД**

Швидке поширення і поглиблення застосувань штучного інтелекту (ШІ) докорінно трансформує не тільки виробничі процеси, але й суспільні практики, інститути та самоусвідомлення праці. Сьогодні автоматизація перестала бути суто механічним заміщенням ручних операцій: алгоритми й машинне навчання беруть на себе дедалі складніші когнітивні функції - від аналізу документів і прогнозування ризиків до генерації текстів і підтримки прийняття рішень. Така зміна якісно змінює попит на навички, структуру зайнятості та розподіл економічних вигод у суспільстві.

Це дослідження спрямоване на системне виявлення ключових трендів впливу цифрових інтелектуальних систем на ринок праці та на розвиток нових форм соціальної нерівності, що виникають внаслідок алгоритмізації праці. Акцент зроблено на соціально-політичних наслідках: як зміни в організації праці й контролі за нею перетворюють економічні вигоди на джерело соціальної напруги або, навпаки, на можливість для інклюзивного розвитку.

Метою дослідження є розробка рекомендацій в царині соціальної політики, які сприятимуть справедливому перерозподілу вигод від ШІ: інструменти перекваліфікації, механізми соціального захисту для нестабільної зайнятості, податкові та регуляторні рішення для зменшення концентрації доходів. Для цього дослідження поєднує кількісні й якісні методи: систематичний огляд сучасної наукової та політичної літератури, аналіз статистичних показників.

Обмеження наукового доробку обумовлені швидкою еволюцією технологій і неоднорідністю доступних даних на національному рівні: багато показників потребують довшої часової серії для надійних висновків, а також існує ризик присутності прихованих упереджень у джерелах інформації. Водночас робота прагне запропонувати практично орієнтовані рішення, які

можуть бути застосовані в коротко- й середньостроковій перспективі, а також сформулювати напрями подальших досліджень.

Сучасні алгоритми штучного інтелекту вже переступили межу заміщення лише фізичної чи рутинної роботи - вони витісняють фахівців середнього рівня: бухгалтерів, перекладачів, аналітиків, юристів. Цей зсув означає, що ринок праці будується за новими правилами, що мають глибокі соціально-економічні наслідки.

По-перше: автоматизація завдань когнітивного характеру зменшує попит на професії зі стандартними повторюваними функціями. Як зауважено в українському дослідженні: «У рамках цифрової економіки автоматизація і роботизація праці в Україні створює умови, за яких працівники муситимуть опановувати інформаційні технології або ризикують бути витісненими» [1; с. 155]. По-друге: зростає вимога не лише бути спеціалістом у певній області, але й вміти працювати з алгоритмами, аналізувати їх результати, доповнювати людським досвідом та інтерпретацією. Тобто навичка автоматизації чи співпраці з ШІ стає критичною для середнього фахівця. По-третє: формально стабільна «середня ланка» праці, яка раніше забезпечувала соціальну мобільність, стає ризиковою. Якщо раніше спеціалізація середнього рівня давала впевненість у зайнятості і стабільному доході, то тепер ця ланка може бути зменшена або замінена алгоритмічною роботою.

Ця трансформація створює велику небезпеку появи «подвійного ринку праці»: з одного боку - високооплачувані, спеціалізовані ІТ-фахівці, аналітики, розробники алгоритмів; з іншого - працівники низькооплачуваних сфер обслуговування або ті, чиї завдання стали алгоритмізованими. Розрив між цими полюсами зростає: можливості руху з нижчого сегмента в середній щезають, а середній сегмент стискається або зникає як переходова ланка.

Аналіз вітчизняної ситуації підтверджує, що цифровізація праці змінює логіку зв'язку між працею, навичками і капіталом: «технології автоматизації і роботизації праці... потребують перебудови системи формальної та неформальної освіти, яка враховуватиме запити цифрової економіки» [4; с. 2]. Цей висновок підкреслює, що не лише технічна автоматизація має значення, але й соціальні інститути - освіта, державна політика, перекваліфікація визначають, чи стане автоматизація джерелом соціального прогресу чи загрозою.

Економічна вигода від автоматизації і впровадження систем штучного інтелекту (ШІ) не автоматично трансформується в суспільне благо - тобто не обов'язково призводить до покращення умов чи зменшення нерівності. Натомість існує ризик, що вигоди концентруються у вузькому колі, тоді як більшість працівників - особливо тих, чії функції автоматизовані, залишаються без компенсацій і підтримки. Тому питання: як саме перерозподілити ці вигоди, стає ключовим у соціальній політиці цифрової ери [1; с. 160].

Вигода від автоматизації включає підвищення продуктивності, зниження витрат, створення нових товарів і послуг, збільшення прибутків технологічних компаній. Але значна частина цих прибутків не «спускається» до працівників: наприклад, зростання продуктивності не завжди супроводжується зростанням заробітної плати чи зайнятості. Крім того, доступ до даних і алгоритмів, які породжують прибутки, часто контролюється невеликою кількістю глобальних гравців. Цей розрив між тими, хто володіє технологією, капіталом, алгоритмами, і тими, хто обслуговує або підпорядковується цим системам породжує новий виклик соціальної політики [2; с. 389].

Класичні інструменти соціального захисту: страхування від безробіття, пенсійні системи, допомога малозабезпеченим були розроблені для індустріальної економіки ХХ століття, де переважала стабільна зайнятість і чіткий поділ між роботодавцем та працівником. Натомість цифрова епоха, зокрема поширення штучного інтелекту і платформеної економіки, формує гібридні моделі зайнятості: часткова зайнятість, віддалена робота, гіг-контракти, самозайнятість через платформи. Такі форми не підпадають під класичні механізми соціального страхування, що створює так звану «зону невидимості» - мільйони людей опиняються поза межами соціальної підтримки. У цьому контексті дедалі більш актуальним стає концепт «справедливого цифрового переходу», який поєднує технологічну модернізацію з гарантіями соціальної справедливості [6; с. 92].

ШІ змінює саме уявлення про працю і соціальний статус. Алгоритми визначають, хто отримає роботу, яку оплату і коли; цифрові платформи, як Uber чи Upwork, стають фактичними «роботодавцями без відповідальності». В результаті соціальна політика, що базується на класичних трудових відносинах, втрачає ефективність. Наявні інструменти соціального захисту не

враховують нових типів ризиків, пов'язаних із диджиталізацією, гіг-економікою та нестандартною зайнятістю. Потрібні інноваційні підходи до гарантування базового рівня соціальної безпеки для всіх учасників ринку праці [6; с.88].

Підхід до нової соціальної політики має спиратися на кілька ключових принципів:

- Інклюзивність - охоплення всіх категорій працівників, зокрема самозайнятих і тих, хто працює через цифрові платформи.
- Гнучкість - адаптація соціальних інструментів до нових форм роботи (наприклад, можливість накопичувального страхування для гіг-працівників).
- Цифрова грамотність і участь - кожен громадянин має отримати рівний доступ до цифрових послуг, перекваліфікації та соціальних ресурсів онлайн.
- Співвідповідальність бізнесу і держави - великі технологічні компанії повинні брати участь у фінансуванні програм соціального розвитку, особливо в секторах, де автоматизація витісняє людей [5; с. 35].

Як підкреслюється в аналітичній доповіді Центру внутрішньо політичних досліджень за 2024 р.: «Майбутня система соціального захисту має враховувати гібридні форми зайнятості, використання штучного інтелекту в управлінні працею та необхідність забезпечення рівного доступу громадян до цифрових сервісів соціальної підтримки» [3; с. 3].

Отже, потрібна нова модель соціальної політики, де цифровізація - не просто інструмент автоматизації послуг, а механізм соціальної інтеграції: від перекваліфікації до цифрової етики, від рівного доступу до алгоритмів - до участі громадян у формуванні політики даних. Це шлях до справедливого цифрового суспільства, де технології працюють на людину, а не навпаки. Штучний інтелект трансформує не лише економічні структури, а й саму логіку соціальної мобільності, змінюючи уявлення про працю, компетентність і справедливість. Якщо в індустріальному суспільстві соціальний успіх визначався рівнем освіти, кваліфікації та досвідом, то у цифровому світі вирішальними стають доступ до технологій, цифровий капітал і здатність взаємодіяти з алгоритмами. Це створює нову форму соціальної стратифікації, де межі між «тими, хто програмує», і «тими, кого програмують», стають дедалі чіткішими.

Ключовим завданням сучасної соціальної політики стає відновлення цінності людської праці в умовах алгоритмізованого суспільства. Це передбачає:

- підтримку професій, що забезпечують соціальну згуртованість і людський контакт (освіта, охорона здоров'я, соціальна робота);
- інвестиції у розвиток цифрової грамотності й етичної культури;
- формування політики «справедливого цифрового переходу», яка гарантує, що користь від автоматизації стане суспільним благом, а не привілеєм вузької групи технологічних еліт.

Майбутнє соціальних відносин у добу штучного інтелекту залежить від здатності поєднати інноваційність із гуманізмом, а технологічний прогрес - із соціальною відповідальністю. Людина має залишитися не лише користувачем, а й автором змін, які формують цифрову цивілізацію.

### Список використаних джерел

1. Ключ І.Ю., Гуменюк В.В. Впровадження штучного інтелекту в діяльність організацій *Управління економікою: теорія та практика*. Чумаченківські читання, 2024. С. 154–167. URL: <https://www.chumachenko-readings.org/download/2024/12-Klius.pdf>
2. Маковоз О.С., Лисенко С.М. Цифрова трансформація – складова сталого розвитку. *Економіка підприємства: сучасні проблеми теорії та практики*: Матеріали тринадцятої міжнар. наук.-практ. конф., 13 вересня 2024 р. Одеса: ОНЕУ, 2024. С. 388 – 389. URL: <http://dspace.oneu.edu.ua/jspui/handle/123456789/18613>
3. Пріоритети розвитку реального сектора в умовах війни та повоєнного відновлення економіки України : аналіт. доп. / [О. В. Собкевич, А. В. Шевченко, В. М. Русан та ін.] ; за загальн. ред. Я. А. Жаліла. Київ : НІСД, 2024. 104 с. <https://doi.org/10.53679/NISSanalytrep.2024.03>
4. Сидоренко О.В., Цішевський Б.В. Базові аспекти цифровізації ринку праці в Україні. *Проблеми сучасних трансформацій. Серія : економіка та управління*. 2024. вип. 12 (Березень). <https://doi.org/10.54929/2786-5738-2024-12-03-01>.
5. Філімонова Т.О. Проблеми цифрової нерівності в Україні. Шляхи подолання. *Інформація та соціум*. 2021. С. 34-35. URL: <https://jias.donnu.edu.ua/article/view/11481>
6. Шлапак А.В., Іващенко О.А. Управління соціальною безпекою в умовах цифрової економіки: локальний рівень. *БІЗНЕСІНФОРМ*. 2024. № 1. С. 87-94. [https://elibrary.kubg.edu.ua/id/eprint/48631/1/Shlapak\\_A\\_Ivashchenko\\_O\\_BI\\_2024\\_FE\\_U.pdf](https://elibrary.kubg.edu.ua/id/eprint/48631/1/Shlapak_A_Ivashchenko_O_BI_2024_FE_U.pdf).

НАУКОВЕ ВИДАННЯ

**УКРАЇНА В УМОВАХ СОЦІАЛЬНОЇ ТА ЦИФРОВОЇ  
ТРАНСФОРМАЦІЇ: ТЕОРЕТИЧНІ МОДЕЛІ ПРАВОВОГО  
РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ**

МАТЕРІАЛИ НАУКОВО-ПРАКТИЧНОЇ КОНФЕРЕНЦІЇ

*Київ, 5 листопада 2025 року*

Електронне видання

*Макет збірника та комп'ютерна верстка:*

*М. Дубняк, С. Дорогих*

*Упорядкування:*

*І. Корж, В. Фурашев*

Ум-друк. арк. 12.

Зам. № 2512-03.

Видавець ПП «Фенікс»

(Свідоцтво суб'єкта видавничої справи ДК № 1044 від 17.09.02).

Україна, м. Одеса, 65009, вул. Зоопаркова, 25.

e-mail: fenix-izd@ukr.net