

Міжнародні підходи до розвитку ШІ та принципів відповідальності за використання: висновки для України

Ольга Петрів, Олена Спесивцева



Видавець

Представництво Фонду Конрада Аденауера в Україні (Київ)

Авторки

Ольга Петрів, юристка у сфері штучного інтелекту напряму «Незалежні медіа» Центру демократії та верховенства права;

Олена Спесивцева, юристка з авторського права та медійного права напряму «Незалежні медіа» Центру демократії та верховенства права.

Усі права захищені. Запити щодо рецензійних примірників та інші звернення стосовно цієї публікації слід надсилати видавцю. Відповідальність за погляди, висновки та рекомендації, висловлені в даній публікації, повністю покладається на автора (авторів), і їх інтерпретації не обов'язково відображають погляди чи політику Фонду Конрада Аденауера.

Публікація підготовлена в рамках проекту «Посилення аналітичних спроможностей прийняття рішень у сфері зовнішньої політики за допомогою громадянського суспільства» Центру міжнародної безпеки за підтримки Представництва Фонду Конрада Аденауера в Україні (Київ).



© 2025 Представництво Фонду Конрада Аденауера в Україні (Київ)

вул. Богомольця Академіка, 5, кв. 1, 01024 Київ, Україна

Тел.: +380444927443

<https://www.kas.de/en/web/ukraine>



© 2025 Центр міжнародної безпеки

вул. Бородіна Інженера, 5-А, 02092 Київ, Україна

Тел.: +380999833140

<https://intsecurity.org/>



ЦЕНТР ДЕМОКРАТІЇ ТА
ВЕРХОВЕНСТВА ПРАВА

© 2025 Центр демократії та верховенства права

вул. М. Заньковецької, 3/1, оф. 12, м. Київ, 01001

Тел.: +380678282074

info@cedem.org.ua

<https://cedem.org.ua>

<https://www.facebook.com/CEDEMUA/>

Зміст

Вступ	4
Підхід ЄС до розвитку ШІ та принципів відповідальності за використання	6
Імплементація положень AI Act в українське законодавство	8
Рамкова конвенція Ради Європи зі штучного інтелекту, прав людини, демократії та верховенства права як орієнтир для України	11
Підхід США до регулювання штучного інтелекту	13
Законодавча «мозаїка»: Конгрес, експерименти штатів та інституційний поділ	14
Використання AI у сфері національної безпеки США	16
Структурні виклики США	17
Висновки та рекомендації для України	18

Вступ

Стрімкий розвиток штучного інтелекту (далі – ШІ) сформував суперечливе ставлення до нього. ШІ визнається потужним інструментом для економічного зростання держав, розвитку інновацій і підвищення результативності та ефективності праці в різних сферах – від охорони здоров'я та культури до національної безпеки та правосуддя. З іншого боку, очевидно, що застосування інструментів ШІ супроводжується ризиками, пов'язаними з дискримінацією, втручанням у приватність, загрозою персональним даним, відсутністю прозорості рішень. Це подвійне сприйняття вимагає збалансованих підходів до регулювання ШІ, яке буде впроваджувати ефективні механізми контролю та юридичної відповідальності та одночасно сприятиме, а не обмежуватиме розвиток технологій.

Розвиток ШІ з самого свого початку був обумовлений потребами національної безпеки. І якщо сьогодні державна політика України у сферах національної безпеки і оборони спрямовується комплексно на забезпечення воєнної, зовнішньополітичної, державної, економічної, інформаційної, екологічної безпеки, безпеки критичної інфраструктури, кібербезпеки України та на інші її напрями, то розвиток систем штучного інтелекту безумовно може сприяти виконанню зазначеного завдання.

Сьогодні постають принципово важливі питання, що стосуються етичних, правових і стратегічних аспектів використання ШІ:

- ▶ Яким чином забезпечити надійність і безпечність функціонування систем ШІ в умовах воєнного конфлікту?
- ▶ Як гарантувати дотримання прав людини при впровадженні ШІ у сфері безпеки?
- ▶ Хто нести відповідальність за рішення ШІ?
- ▶ Яким чином запобігти дискримінаційним практикам і упередженості в алгоритмах?
- ▶ Як досягти прозорості й підзвітності систем ШІ, зокрема при використанні їх для обробки великих обсягів персональних даних?
- ▶ Як зберегти баланс між інноваційністю та регулюванням, щоб не гальмувати розвиток, але й не допустити зловживань?

Ці питання потребують не лише фахового аналізу, а й формування цілісної національної політики щодо використання ШІ в контексті безпеки, яка відповідала б міжнародним стандартам, національним інтересам і принципам верховенства

права. Тож сьогодні на порядку денному провідних держав та регіонів світу стоїть формування підходу до регулювання ШІ, який поєднує етичні принципи, правові та безпекові гарантії. Наприклад, підхід Європейського Союзу (далі – ЄС) є превентивним та орієнтованим на права людини та національну безпеку. США ж віддають перевагу гнучкості, стимулюючи розвиток ШІ через партнерство з приватним сектором та поступове формування рекомендацій. Втім, досвід обох регіонів є важливим для України, де модель ЄС є орієнтиром правової гармонізації, а запозичення досвіду США може сприяти розвитку інновацій.

Варто згадати про те, що ще у грудні 2020 року Кабінет Міністрів України затвердив [Концепцію розвитку штучного інтелекту в Україні](#), яка визначає мету, принципи та завдання розвитку технологій ШІ в Україні як одного з пріоритетних напрямів у сфері науково-технологічних досліджень. У ній зазначено про проблеми, які потребують розв'язання, серед яких, зокрема, низький рівень цифрової грамотності населення щодо загальних аспектів, можливостей, ризиків та безпеки використання ШІ, відсутність або недосконалість правового регулювання ШІ, низький рівень інвестицій у розроблення технологій ШІ, відсутність єдиних підходів, що застосовуються при визначенні критеріїв етичності під час розроблення та використання технологій ШІ для різних галузей, видів діяльності та сфер національної економіки тощо. Вирішення цих проблем, безумовно, потребує комплексного підходу: аналізу практики регулювання ШІ в інших державах і регіонах, розробки законодавчих ініціатив і комплексних принципів використання систем ШІ, консультацій з зацікавленими сторонами (серед яких як розробники та користувачі систем ШІ, так і інвестори та представники влади).

[Концепція розвитку штучного інтелекту в Україні](#) вже визначила принципи розвитку та використання технологій ШІ, які хоч і є рекомендаційними, проте можуть слугувати орієнтиром. Серед принципів зокрема:

- ▶ Розроблення та використання систем ШІ лише за умови дотримання верховенства права, основоположних прав і свобод людини і громадянина, демократичних цінностей, а також забезпечення відповідних гарантій під час використання таких технологій;
- ▶ Відповідність діяльності та алгоритму рішень систі вимогам законодавства про захист персональних даних, а також додержання конституційного права кожного на невтручання в особисте і сімейне життя у зв'язку з обробкою персональних даних;
- ▶ Забезпечення прозорості та відповідального розкриття інформації про системи ШІ, надійне та безпечне функціонування систем ШІ протягом усього їх життєвого циклу та здійснення на постійній основі їх оцінки та управління потенційними ризиками;
- ▶ Покладення на організації та осіб, які розробляють, впроваджують або використовують системи ШІ, відповідальності за їх належне функціонування відповідно до зазначених принципів тощо.

Підхід ЄС до розвитку ШІ та принципів відповідальності за використання

Результати багаторічних обговорень можливостей, викликів та ризиків, пов'язаних із розвитком систем ШІ, поступово знаходять закріплення в регіональних і національних джерелах права. Відбувається перехід від декларацій, рекомендацій щодо відповідального використання систем ШІ чи внутрішніх правил (кодексів) поведінки окремо взятих організацій до ухвалення нормативно-правових актів.

Відповіддю на потребу в регулюванні технологій ШІ та їхнього впливу на суспільство через стрімкі темпи розвитку ШІ в ЄС став, зокрема, [Artificial Intelligence Act \(Regulation \(EU\) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence](#), далі – AI Act) – регламент ЄС, який набрав чинності 1 серпня 2024 року, мета якого – створити безпечне середовище для використання та розвитку ШІ.

Для України вивчення досвіду ЄС має принципове значення з кількох причин. У ситуації, коли технології ШІ вже активно застосовуються, а спеціалізоване законодавство ще не сформоване, виникає нагальна потреба у впровадженні чітких стандартів щодо розробки та використання ШІ. Формуючи власну нормативну модель, Україні варто враховувати підхід ЄС, що відповідає стратегічному курсу на євроінтеграцію та зобов'язанням, які випливають із Угоди про асоціацію з ЄС. Крім того, ЄС сьогодні виступає глобальним лідером у сфері цифрового регулювання, прикладом чого є прийняття у 2016 році Загального регламенту про захист даних ([General Data Protection Regulation](#), далі – GDPR), який [спрямований](#) на захист фізичних осіб стосовно обробки їх персональних даних та вільний рух таких даних.

AI Act поширюється не на всі системи ШІ. Системи, які випускаються на ринок, вводяться в експлуатацію або використовуються для військових, оборонних або національних цілей безпеки, виключені зі сфери його застосування незалежно від того, який тип суб'єкта здійснює таку діяльність (державний чи приватний суб'єкт). Національна безпека залишається виключною відповідальністю держав-членів відповідно до статті 4(2) Договору про ЄС.

Проте, якщо така система ШІ (що розроблена та використовується для військових, оборонних або національних цілей безпеки) використовується одночасно і поза цими цілями (наприклад, для цивільних або гуманітарних цілей, цілей правоохоронної діяльності або громадської безпеки), така система підпадає під дію AI Act. У такому випадку, суб'єкт, який використовує систему ШІ для інших цілей, ніж

військові, оборонні або національної безпеки, повинен забезпечити відповідність систем ШІ вимогам AI Act.

AI Act складається з 12-ти розділів, кожен з яких регулює окрему сферу застосування та розробки ШІ. Основні положення цього законодавчого акта передбачають (детальніше за [посиланням](#)):

1. Визначення рівнів ризику систем ШІ (заборонені системи ШІ, системи ШІ з високим ступенем ризику, системи ШІ з обмеженим рівнем ризику, системи ШІ з мінімальним ризиком або не ризикові).
2. Введення обов'язкової сертифікації для певних систем ШІ (систем біометричної ідентифікації, критичної інфраструктури, освітні або професійні оцінювальні системи).
3. Встановлення вимог для певних систем ШІ щодо необхідності інформування користувачів про те, що вони взаємодіють із системою ШІ, а не людиною.
4. Встановлення правил прозорості для систем ШІ, призначених для взаємодії з фізичними особами, систем розпізнавання емоцій, а також систем ШІ, які використовуються для створення або оброблення зображень, аудіо- або відеоконтенту.
5. Створення єдиного європейського ринку ШІ – введення єдиних правил для роботи з технологіями ШІ на території ЄС, що сприятиме вільному руху систем ШІ та послуг на європейському ринку.
6. Заборону на використання певних методів ШІ. Наприклад, заборонено розміщення на ринку та використання систем ШІ, які застосовують техніки, що виходять за межі свідомості людини, змушують прийняти певне рішення.

Щодо відповідальності AI Act наголошує на тому, що доцільно, щоб конкретна фізична або юридична особа, визначена як постачальник систем ШІ, несла відповідальність за розміщення на ринку або введення в експлуатацію системи ШІ з високим рівнем ризику, незалежно від того, чи є ця фізична або юридична особа особою, яка спроектувала або розробила систему. Водночас у науковій спільноті критикують те, що відповідальність у AI Act не має імперативного характеру, а заборона дискримінації є лише імпліцитною. Це свідчить про подальшу потребу удосконалення законодавства з урахуванням нових викликів.

Фундаментом цієї моделі регулювання систем ШІ є поєднання правових і етичних принципів: прозорість, справедливість, повага до прав людини, уникнення дискримінації, людський контроль та запобігання шкоді. ЄС заклав основи системи регулювання, яка в контексті національної безпеки базується на забороні систем із неприйнятним рівнем ризику, посиленому регулюванні високоризикових систем, зобов'язанні постачальників систем ШІ дотримуватися процедур управління даними, прозорості та кібербезпеки, покладенні відповідальності за відповідність систем ШІ на розробників і постачальників.

Імплементация положень AI Act в українське законодавство

AI Act сьогодні є інструментом, який запроваджує комплексний підхід до врегулювання розробки та безпечного застосування інструментів ШІ у різних сферах діяльності людини та держав. Але кожна окрема сфера вимагає імплементції та адаптації до її специфіки.

Наприклад, сьогодні в Україні регулювання використання результатів роботи ШІ відбувається відповідно до Закону України «Про авторське право і суміжні права». AI Act наголошує на тому, що моделі ШІ загального призначення, здатні генерувати текст, зображення та інший контент, відкривають унікальні можливості для інновацій, але також становлять виклики для художників, авторів та інших творців, а також для способу створення, розповсюдження, використання та споживання їхнього творчого контенту. Розробка та навчання таких моделей вимагають доступу до величезних обсягів тексту, зображень, відео та інших даних. Будь-яке використання контенту, захищеного авторським правом, вимагає дозволу відповідного правовласника, якщо не застосовуються відповідні винятки та обмеження авторського права.

Регулювання ШІ також може, ба більше, має, відбуватися відповідно до Закону України «Про захист персональних даних». Наприклад, AI Act визначає, що будь-яка обробка персональних даних, пов'язана з використанням систем ШІ для біометричної ідентифікації, заборонена, крім випадків, пов'язаних з використанням систем біометричної ідентифікації з метою забезпечення правопорядку. Оскільки за допомогою ШІ можуть здійснюватись автоматичні процеси збору та обробки персональних даних, наприклад, для аналізу поведінки користувачів на сайті або для розробки персоналізованих рекламних кампаній, то Закон України «Про захист персональних даних» має містити вимоги щодо обробки персональних даних під час функціонування систем ШІ. Зокрема, це стосується їх збору, зберігання, використання, передачі та захисту, вимог до автоматизованої обробки даних, заборони на біометричну ідентифікацію без спеціального дозволу, гарантій безпеки при використанні ШІ в державному управлінні.

Наразі Україна бере участь в імплементції AI Act шляхом запровадження регуляторних «пісочниць» (sandbox), про створення який йдеться у розділі 5 AI Act. Регуляторні «пісочниці» – це безпечний і контрольований простір для експериментів, який мають створити національні компетентні органи на національному рівні для сприяння розробці та випробуванню інноваційних систем ШІ під суворим регуляторним наглядом, перш ніж ці системи будуть випущені на ринок або введені в експлуатацію.

Зокрема, можливість обробки персональних даних для розробки певних систем ШІ в інтересах суспільства в регуляторних «пісочницях» окремо регулюється AI Act. Так, вони можуть оброблятися з метою розробки, навчання та тестування певних систем ШІ в пісочниці, якщо виконуються наступні умови, серед яких:

1. Системи ШІ розробляються для захисту істотних суспільних інтересів державним органом або іншою фізичною чи юридичною особою;
2. Існують ефективні механізми моніторингу для виявлення будь-яких високих ризиків для прав і свобод суб'єктів даних, які можуть виникнути під час експерименту;
3. Персональні дані, що підлягають обробці в контексті пісочниці, знаходяться у функціонально захищеному середовищі обробки даних під контролем постачальника, і доступ до цих даних мають лише уповноважені особи;
4. Персональні дані не можуть бути передані за межі пісочниці, захищаються за допомогою відповідних технічних та організаційних заходів і видаляються після закінчення участі;
5. Короткий виклад проекту ШІ, розробленого в пісочниці, його цілей та очікуваних результатів публікується на веб-сайті компетентних органів та інші умови.

AI Act визначає, що держави-члени ЄС повинні забезпечити, щоб їхні компетентні органи створили принаймні одну регуляторну «пісочницю» для штучного інтелекту на національному рівні, яка має запрацювати до 2 серпня 2026 року. Враховуючи стрімкий євроінтеграційний рух України, національні компетентні органи мають зважати на те, що, хоча визначений строк не застосовний до України, дія AI Act автоматично пошириться на Україну одразу після набуття членства, а тому робота над створенням регуляторних пісочниць триває і має перші результати.

Міністерство цифрової трансформації та Український фонд стартапів оголосили про запуск [Sandbox](#) – «пісочниці» для українських компаній, що пропонують високотехнологічні рішення для сфер цифрової економіки, загальної інфраструктури, публічних послуг, оптимізації роботи органів державної влади, охорони здоров'я, біотехнологій, освіти і науки, агросектору, оборони тощо.

Тенденція до подальшого впровадження регуляторних «пісочниць» демонструє необхідність встановлення чітких вимог національного законодавства до захисту персональних даних, етичної перевірки, контролю за навчанням та тестуванням систем у таких «пісочницях». У грудні 2020 року Кабінет Міністрів України затвердив [Концепцію розвитку штучного інтелекту в Україні](#), у жовтні 2023 році було представлено [Дорожню карту з регулювання штучного інтелекту в Україні](#), а у червні 2024 року було презентовано [Білу книгу з регулювання ШІ в Україні](#). У Білій книзі з регулювання ШІ окрему увагу приділено саморегулюванню як інструменту

формування етичних стандартів та доповнення державної політики. У червні 2025 року 14 провідних українських компаній підписали [Меморандум щодо саморегулювання у сфері ШІ](#), взявши на себе добровільні зобов'язання щодо безпечної та прозорої розробки технологій.

Рекомендаційні документи та формування національної стратегії є орієнтиром у сфері ШІ та невід'ємною складовою формування культури розробки, введення на ринок та використання систем ШІ. Втім, першочерговим завданням держави сьогодні є розробка нормативно-правової бази у різних сферах, зокрема, щодо відповідальності за рішення, ухвалені або виконані ШІ-системами, із розмежуванням обов'язків розробника, користувача, постачальника та держави.

Рамкова конвенція Ради Європи зі штучного інтелекту, прав людини, демократії та верховенства права як орієнтир для України

Україна має розробляти нормативно-правову базу для регулювання ШІ, оскільки у 2025 році вона приєдналася до [Рамкової конвенції Ради Європи зі штучного інтелекту, прав людини, демократії та верховенства права](#) (далі – Конвенція з ШІ). Це [перший обов'язковий договір](#) у сфері ШІ, що охоплює весь життєвий цикл систем ШІ. Кожна держава вживає заходи для забезпечення підзвітності та відповідальності за негативний вплив на права людини, демократію та верховенство права, що виникає в результаті діяльності протягом життєвого циклу систем штучного інтелекту.

У [пояснювальній записці до Конвенції з ШІ](#) зазначається, що принцип підзвітності та відповідальності стосується необхідності створення механізмів, за допомогою яких особи, організації або суб'єкти, відповідальні за діяльність у рамках життєвого циклу систем ШІ, нестимуть відповідальність за негативний вплив на права людини, демократію або верховенство права. Від держав вимагається створити механізми, які можуть бути застосовані до діяльності протягом життєвого циклу систем ШІ, щоб забезпечити виконання цієї вимоги. Це може включати судові та адміністративні заходи, цивільні, кримінальні та інші режими відповідальності, а в державному секторі – адміністративні та інші процедури, щоб рішення могли бути оскаржені.

Необхідно чітко розподіляти відповідальність та можливості відстежувати дії та рішення конкретних осіб або суб'єктів з урахуванням різноманітності відповідних учасників та їхніх ролей і обов'язків. Всі суб'єкти, відповідальні за діяльність у рамках життєвого циклу систем ШІ, незалежно від того, чи є вони державними чи приватними організаціями, повинні підлягати існуючій системі правил, правових норм та інших відповідних механізмів держави.

Підхід України для виявлення, аналізу, оцінки ризиків та впливу систем ШІ на права людини, демократію та верховенство права сьогодні може розвиватися на основі методології HUDERIA (Human Rights, Democracy, and the Rule of Law Impact Assessment), яка розроблена [Комітетом зі штучного інтелекту](#) (CAI) Ради Європи (детальніше за [посиланням](#)), яка складається з декількох етапів: ідентифікація ризику системи ШІ, оцінка впливу, оцінка управління (відповіді на соціотехнічні питання та питання, що стосуються прав людини, демократії та верховенства права), пом'якшення негативного впливу та оцінка системи ШІ.

Методологія може бути використана як державними інституціями, так і приватними суб'єктами. Вона не має обов'язкової юридичної сили, водночас держави – сторони Конвенції з ШІ – можуть гнучко використовувати чи адаптовувати її повністю або частково під власні національні інтереси. Методологія може сприяти тому, щоб розробки у сфері ШІ не порушували права людини, дотримувалися демократичних принципів та верховенства права.

Підхід США до регулювання штучного інтелекту

Американська політика регулювання штучного інтелекту від початку будується на внутрішній напрузі: Вашингтон прагне утримати технологічне лідерство, але водночас зіштовхується з дедалі гучнішими вимогами щодо безпеки, прав людини та конкурентної рівноваги. Від [жовтневого указу Байдена – ЕО 14110](#) (2023) – до «дерегуляційного» [ЕО 14179](#) (січень 2025) Білий дім пройшов шлях від спроби встановити обов'язкові «запобіжники» до сигналу про лібералізацію для бізнесу. Така швидка зміна курсу створила ситуацію «політики маятника»: ринок отримує суперечливі імпульси, а законодавець все ще не спромігся дати всеосяжну відповідь ([Federal Register](#), [The White House](#)).

[Указ ЕО 14110](#) запровадив вісім загальних принципів розвитку і впровадження ШІ: безпечність та ефективність, захист прав споживачів, приватність та громадянські свободи, справедливість, прозорість та пояснюваність, управління ризиками, інновації та конкуренція, глобальне лідерство США.

[Національному інституту стандартів і технологій \(NIST\)](#) та [Офісу менеджменту і бюджету \(OMB\)](#) було доручено виробити стандарти оцінки ризиків. Також Указ вперше зобов'язав федеральні агентства публікувати AI Impact Assessment для критичних систем. На практиці документ залишався переважно «soft law», мав рекомендаційний характер – санкцій за невиконання він не передбачав.

Проте вже у січні 2025 року інший указ – [ЕО 14179](#) – скасував низку попередніх обмежень. Його підписав президент Дональд Трамп, аргументуючи це прагненням зберегти світове лідерство США в технологіях.

Щоб не залишити ШІ поза контролем, через два місяці Офіс менеджменту і бюджету США видав [меморандум М-24-10](#). У ньому йшлося про мінімальні вимоги для найбільш потужних моделей ШІ, які включають тестування на безпеку, зовнішні аудити та ведення журналів роботи алгоритмів.

Законодавча «мозаїка»: Конгрес, експерименти штатів та інституційний поділ

У 119-му Конгресі на розгляді перебуває одразу декілька важливих законопроектів. Усі ці ініціативи демонструють наступну тенденцію: складні питання ШІ Конгрес розв'язує «точковими» законами ([congress.gov](https://www.congress.gov), [congress.gov](https://www.congress.gov), [WIRED](https://www.wired.com)).

Прикладом точкового підходу є [TAKE IT DOWN Act](#), поданий до Сенату у січні 2023 року. Законопроект криміналізує створення та розповсюдження сексуалізованих deepfake-зображень без згоди постраждалих. Його територія дії – усі штати США. Це перший випадок, коли Конгрес адресно зачіпає конкретну технологію, а не весь ШІ загалом.

На розгляді в Конгресі перебуває [CREATE AI Act \(H.R. 2385\)](#), представлений у березні 2023 року. Він передбачає створення державних обчислювальних кластерів для науковців, щоб забезпечити рівний доступ до обчислювальних потужностей, які нині зосереджені у приватних хмарних гігантів. Закон не містить вимог до безпеки моделей.

Ще один важливий проєкт – [American Privacy Rights Act \(APRA\)](#), представлений у березні 2024 року. Він має на меті уніфікувати підхід до захисту персональних даних на федеральному рівні, проте наразі застряг на етапі обговорення в енергетичному комітеті Палати представників.

Усі ці ініціативи демонструють: замість створення єдиного закону про ШІ, США реалізують точкові втручання через окремі галузеві чи тематичні акти. Якщо порівнювати підходи США та ЄС можна сказати про те, що ЄС обрав широкий підхід і запроваджує загальні принципи регулювання ШІ, тоді як США має галузевий підхід. А Україні ж слід мати золоту середину.

Експерименти з регулювання ШІ проводяться на рівні штатів. У штаті Техас, підписаний [SB 1621](#) у червні 2023 року, криміналізує створення deepfake-зображень за участю неповнолітніх. Закон діє лише в межах штату. [SB 20](#), ухвалений у травні 2025, підписаний 20 червня 2025 року, забороняє володіння, розповсюдження та створення AI-генерованої дитячої порнографії. Він передбачає до 2 років ув'язнення та штраф до \$10 000. У штаті Вірджинія [HB 697 \(2024\)](#) вводить кримінальну відповідальність за використання «synthetic media» з метою шахрайства. Закон надає нову категорію злочинів. У штаті Теннесі [ELVIS Act](#) захищає артистів від несанкціонованого використання голосу або образу за допомогою ШІ. Закон набув чинності 1 липня 2024 року. У штаті Юта [SB 149](#) закон Utah SB 149 (Artificial Intelligence Policy Act) набрав чинності 1 травня 2024 року і

передбачає створення Офісу AI-політики, запуск пілотної програми «AI Learning Lab» для тестування генеративних AI-систем, а також встановлює вимоги щодо розкриття факту використання AI під час взаємодії з користувачами. Спочатку він діяв лише до 1 травня 2025 року, але був продовжений до 1 липня 2027 року через ухвалення поправок (SB 226 і SB 332). Закон залишається чинним і вимагає прозорості та обережності від усіх, хто використовує генеративний AI у штаті Юта. [Local Law 144](#) у Нью-Йорку набув чинності 5 липня 2023 та забороняє автоматизовані інструменти найму без попереднього аудиту на предмет упередженості.

У США немає єдиного регулятора штучного інтелекту – замість цього діє мережева модель координації. Окремі агентства відповідають за різні напрями:

1. [NIST](#) (Національний інститут стандартів і технологій) – розробив AI Risk Management Framework 1.0 (AI RMF), а в липні 2024 року представив спеціалізований профіль для generative AI. Цей документ описує 12 специфічних ризиків генеративного ШІ (наприклад, фальшиві факти, витoki даних, екологічне навантаження) та пропонує понад 200 конкретних дій для їхнього управління. Хоча профіль є добровільним, він фактично став стандартом, якого дотримуються компанії.
2. [OSTP](#) (Офіс науково-технологічної політики при Білому домі) – задає напрям політичної дискусії. У 2022 році він представив Blueprint for an AI Bill of Rights, а після указу президента EO 14179 координує міжвідомчий AI Action Plan.
3. [NSF](#) (Національний науковий фонд) – адмініструє пілотну програму NAIRR (National AI Research Resource). Її мета – надати науковцям доступ до обчислювальних ресурсів і відкритих датасетів, щоб зрівняти шанси з Big Tech (логіка – «демократизувати обчислення»).
4. [DARPA](#) (Агентство передових оборонних дослідницьких проєктів) – працює над інтерпретованими AI-системами (XAI 2.0) та інструментами безпеки (Guardrails). Агентство часто використовує закриті дані, що ускладнює незалежну оцінку.
5. [DoD / CDAO](#) (Міністерство оборони / Головне управління цифрових та AI-операцій) – після ініціативи Task Force Lima (2023–2024) створив AI Rapid Capabilities Cell з бюджетом \$100 млн (на 2024–2025 роки). Осередок працює над 15 сценаріями бойового використання генеративного AI.

Отже, у США функції регулювання ІШ розподілені між багатьма установами, і хоча централізованого нагляду немає, інституції формують потужну систему координованого управління ризиками, схожу на модель у сфері кібербезпеки.

Використання AI у сфері національної безпеки США

У сфері національної безпеки США застосовується відокремлений і значно менш публічний режим управління штучним інтелектом. Міністерство оборони (DoD) розробляє власні протоколи та політики, що суттєво відрізняються від цивільного регулювання. Формально зберігається принцип human-in-the-loop, тобто участі людини у процесі прийняття рішень. Водночас цей принцип застосовується гнучко: людський контроль реалізується лише там, де це не суперечить оперативним чи бойовим вимогам. У межах чітко визначених і технічно контрольованих сценаріїв системи штучного інтелекту можуть діяти з частковою або повною автономією.

Використання генеративного ШІ в оборонному секторі інтегрується у цифрову бойову інфраструктуру через JWCC (Joint Warfighting Cloud Capability) – хмарну платформу, що об'єднує сервіси Amazon, Microsoft, Google і Oracle. Вона забезпечує безперервний доступ до даних, обчислювальних ресурсів та аналітичних інструментів у реальному часі для підтримки місій. Хоч JWCC не виконує функцій прямого моніторингу роботи ШІ, вона є основним середовищем для їхнього розгортання. Рівень прозорості у цій сфері значно нижчий, ніж у цивільному використанні: більшість характеристик моделей, методики тестування й результати аудитів не розголошуються, що обмежує незалежну оцінку технологій.

У серпні 2023 року Міністерство оборони ініціювало Task Force Lima – міжвідомчу робочу групу з вивчення потенціалу генеративного ШІ в оборонному контексті. У грудні 2024 року її діяльність була інституціоналізована у вигляді постійного підрозділу – AI Rapid Capabilities Cell (AI RCC) – з бюджетом 100 мільйонів доларів на 2024–2025 роки. AI RCC займається впровадженням генеративного ШІ в низку пріоритетних напрямів, таких як аналіз розвідувальної інформації, планування операцій, інформаційні впливи, логістика та підтримка процесів командування і управління.

Структурні виклики США

США досі не мають всеосяжного федерального закону на кшталт європейського AI Act. Замість цього формується фрагментована екосистема – Конгрес діє точково, тоді як штати ухвалюють власні регуляції, що призводить до появи «комплаєнс-пазлу» з 50 різних режимів. Компанії змушені орієнтуватися одночасно на місцеве, федеральне й навіть міжнародне право, не маючи чіткої стратегічної рамки.

Регуляторна траєкторія також залишається непослідовною. З одного боку, указ президента EO 14179 стимулює розвиток інновацій, спрощуючи доступ до обчислювальних ресурсів і даних. З іншого – меморандум OMB M-24-10 встановлює жорсткі вимоги до федеральних агентств щодо тестування, аудиту та прозорості ШІ-систем. Це створює ефект «регуляторного маятника», коли ринок не має чіткої відповіді: дерегуляція чи контроль.

На етичному рівні гостро постає питання deepfake-маніпуляцій. У відповідь на зростання таких ризиків штати самостійно впроваджують регулювання: Техас (SB 20) криміналізує створення AI-генерованої дитячої порнографії, а Вірджинія (HB 697) забороняє поширення змінених зображень із сексуальним підтекстом без згоди зображеної особи.

На перетині етики та безпеки знаходиться питання прозорості у сфері національної безпеки. Більшість ініціатив Міністерства оборони, зокрема у межах Task Force Lima чи AI RCC, проводяться за зачиненими дверима. З одного боку, це забезпечує оперативну перевагу й швидке впровадження технологій. З іншого – відсутність незалежного аудиту підриває довіру до рішень, які потенційно можуть впливати на безпеку, права людини й міжнародну стабільність.

США демонструють парадокс: найбільший у світі драйвер ШІ досі лишається без «єдиного» закону про ШІ. Натомість вони покладаються на soft-law (NIST RMF), вузькі кримінальні та приватні ініціативи (TAKE IT DOWN, deepfake-акти штатів) і швидкі оборонні ініціативи.

Висновки та рекомендації для України

Україна має шанс адаптувати міжнародні підходи до ШІ під власні потреби, особливо у сфері національної безпеки, яка включає кіберзахист, боротьбу з інформаційними загрозами, охорону критичної інфраструктури та прикордонний контроль.

Подальші кроки України в контексті розвитку та запровадження систем ШІ мають стосуватися наступного:

1. Формування національної стратегії, яка стане орієнтиром у сфері ШІ та невід'ємною складовою формування культури розробки, введення на ринок та використання систем ШІ.
2. Формування цілісної нормативно-правової бази для регулювання ШІ. Важливим кроком до впровадження регулювання систем ШІ та їх ефективного функціонування є інформаційно-просвітницька діяльність серед громадян та українського бізнесу, що допоможе сформувати стале розуміння того, як підвищити продуктивність, результативність, виробництво за допомогою ШІ, як убезпечити себе від ризиків та негативних наслідків використання ШІ.
3. Гармонізація з AI Act зокрема щодо класифікації ШІ-систем за рівнем ризику, вимог прозорості, інформування користувачів, безпеки, закріплення принципу відповідальності постачальників та розробників систем ШІ. Сьогодні Україна має обрати шлях добровільного та ініціативного наближення до права ЄС, яке наразі не є обов'язковим для імплементації.
4. Прийняти рекомендації, що стосуються системи управління ризиками штучного інтелекту (аналог RMF на базі NIST; фреймворк [NIST AI RMF 1.0](#) пропонує практичні орієнтири з управління ризиками ШІ). Це важливо для єдиного підходу для оцінки ризиків у нацбезпеці (наприклад, під час оцінки розробок для Міністерства внутрішніх справ, Служби безпеки України чи Державної прикордонної служби України), дозволяє використовувати його як вимогу в державних закупівлях ШІ-систем, знижує бюрократичне навантаження. Імплементація має передбачати, зокрема, такі кроки: переклад RMF, адаптація як добровільна ДСТУ-рекомендація, рекомендація до застосування в публічному секторі, насамперед в органах національної безпеки.
5. Гарантія людського контролю у критичних рішеннях ШІ та забезпечення громадського контролю, залучення експертної спільноти до оцінки ризиків при впровадженні нових систем ШІ. В межах нацбезпеки ШІ має

застосовуватися з гарантією людського контролю, особливо в сферах автоматизованого виявлення загроз, прикордонної перевірки осіб, аналізу кіберризиків. Нормативно закріпити принцип «людина приймає остаточне рішення», заборонити повну автономність у критичних сценаріях, впровадити симуляційне тестування та логування рішень ШІ в системах безпеки.

6. Створення спеціального незалежного органу з етики ШІ або підрозділу у межах існуючих, який контролюватиме дотримання прав людини у застосуванні ШІ. Зокрема, створити Координаційну раду з ШІ при Кабінеті Міністрів України (КМУ) з безпековим мандатом. Американський підхід OSTP демонструє перевагу мережевого врядування без надмірної централізації. Таким чином, Координаційна рада з ШІ при КМУ має включати представників Міністерства цифрової трансформації, Міністерства внутрішніх справ, Служби безпеки України, Національної академії наук України, Національного координаційного центру кібербезпеки, Головного управління розвідки та ухвалювати методичні рекомендації для впровадження безпечного ШІ в публічних структурах.
7. Сприяння залучення інвестицій у сферу ШІ через державно-приватні партнерства, підтримку стартапів та наукових установ.
8. Інтеграція в міжнародні ініціативи з розроблення стандартів ШІ, зокрема, в межах Ради Європи та ЄС.
9. Запуск пілотного проєкту «U-AIRR» (Ukrainian Artificial Intelligence Research Resource) – національний AI-ресурс для безпеки та науки. [Аналог NAIRR](#) у США – це публічний доступ до GPU-ресурсів для науковців і безпекових установ. В Україні проєкт можна реалізувати на базі спеціалізованої державної хмарної інфраструктури у партнерстві з приватними провайдерами за координації Міністерства цифрової трансформації України. Доступ до GPU-ресурсів може надаватися в межах грантової підтримки для команд, що працюють над ШІ у сферах безпеки, оборони, боротьби з дезінформацією та виявлення загроз.
10. Врегулювання deepfake у політичній рекламі. Штати США, зокрема, Техас і Вірджинія, вже ухвалили закони, що обмежують використання AI для маніпуляцій громадською думкою. В Україні рекомендується внести зміни до Виборчого кодексу, які зобов'язують розкривати використання ШІ або deepfake у політичній агітації.

Для України цінним є саме комбінований підхід: добровільні стандарти, точкове законодавство та міжвідомча координація, що дозволяє не блокувати інновації, але й не залишати прогалин у критичних секторах.