

ШТУЧНИЙ ІНТЕЛЕКТ ЯК ІНСТРУМЕНТ ПРОТИДІЇ ПОЛІТИЧНІЙ ДЕЗІНФОРМАЦІЇ ТА ІНФОРМАЦІЙНИМ ОПЕРАЦІЯМ

Проаналізовано роль штучного інтелекту в протидії політичній дезінформації та інформаційним операціям в умовах цифрової трансформації суспільства. Доведено, що сучасні інформаційні загрози мають системний і динамічний характер та потребують інноваційних технологічних рішень, оскільки традиційні підходи до інформаційної безпеки є недостатніми. Розглянуто еволюцію політичної дезінформації, вплив соціальних платформ і алгоритмічного поширення контенту, а також потенціал штучного інтелекту для виявлення, аналізу й прогнозування дезінформаційних кампаній. Окреслено ризики та обмеження використання штучного інтелекту, зокрема алгоритмічну упередженість, непрозорість рішень, етичні дилеми та адаптацію дезінформаційних акторів до нових технологій, що зумовлює потребу в правових та інституційних механізмах контролю. Зроблено висновок про доцільність інтеграції штучного інтелекту в державну політику протидії дезінформації як складову комплексної стратегії інформаційної безпеки для зміцнення демократичної стійкості та суспільної довіри.

Ключові слова: штучний інтелект (ШІ); політична дезінформація; інформаційні операції; інформаційна безпека; гібридні загрози; алгоритмічні системи; соціальні медіа; інформаційна державна політика; цифрові технології; демократична стійкість.

But Sergii. Artificial intelligence as a tool to counter political disinformation and information operations

The article analyzes the role of artificial intelligence in countering political disinformation and information operations in the context of the digital transformation of society. It argues that modern information threats are systemic and dynamic and therefore require innovative technological solutions, as traditional approaches to information security are insufficient. The study

© БУТ Сергій Сергійович – аспірант Інституту держави і права імені В. М. Корецького НАН України; ORCID: 0009-0004-9964-8110; e-mail: sergey.buth@gmail.com

examines the evolution of political disinformation, the influence of social platforms and algorithmic content distribution, and the potential of artificial intelligence for detecting, analyzing, and forecasting disinformation campaigns. It also outlines the risks and limitations of AI use, including algorithmic bias, lack of transparency in decision-making, ethical dilemmas, and the adaptation of disinformation actors to new technologies, highlighting the need for legal and institutional oversight mechanisms. The article concludes that integrating artificial intelligence into state policy for countering disinformation is advisable as part of a comprehensive information security strategy aimed at strengthening democratic resilience and public trust.

Key words: artificial intelligence (AI); political disinformation; information operations; information security; hybrid threats; algorithmic systems; social media; information state policy; digital technologies; democratic sustainability.

Вступ. Сучасний інформаційний простір дедалі частіше стає полем цілеспрямованих політичних маніпуляцій, дезінформаційних кампаній та комплексних інформаційних операцій, які використовуються як інструмент впливу на громадську думку, виборчі процеси та політичну стабільність держав. Розвиток цифрових платформ, соціальних мереж і миттєвих каналів комунікації суттєво знизив бар'єри для поширення неправдивої або маніпулятивної інформації, водночас ускладнивши її виявлення та нейтралізацію традиційними методами [1].

Особливу загрозу становить поєднання дезінформації з автоматизованими технологіями – бот-мережами, синтетичним контентом, deepfake-відео та алгоритмічним таргетуванням аудиторій. Такі інструменти дозволяють масштабувати інформаційні операції, адаптувати повідомлення під конкретні соціальні групи та приховувати джерела впливу. В результаті зростає ризик підриву довіри до демократичних інститутів, державних органів і засобів масової інформації.

В умовах повномасштабної гібридної війни проти України проблема політичної дезінформації набула екзистенційного характеру, ставши складовою ширших інформаційно-психологічних операцій. Це актуалізує пошук нових технологічних рішень для раннього виявлення, аналізу та протидії інформаційним загрозам [2].

Штучний інтелект дедалі частіше розглядається як перспективний інструмент у цій сфері, проте його використання супроводжується низкою етичних, правових і методологічних викликів. Саме ці суперечності формують наукову проблему, що потребує системного осмислення.

Метою статті є аналіз можливостей та обмежень використання технологій штучного інтелекту в протидії політичній дезінформації та інформаційним операціям, а також визначення ключових викликів і перспектив їх застосування в умовах сучасного інформаційного середовища та гібридних загроз.

Аналіз останніх досліджень і публікацій. Проблематика політичної дезінформації та інформаційних операцій активно досліджується в межах політичної науки, комунікаційних студій, безпекових досліджень і суміжних міждисциплінарних напрямів. У сучасних наукових публікаціях дезінформація розглядається як інструмент політичного впливу, що використовується для маніпуляції громадською думкою, підриву легітимності демократичних інститутів і дестабілізації суспільно-політичних процесів. Значна увага приділяється аналізу механізмів інформаційно-психологічного впливу, ролі соціальних мереж у поширенні маніпулятивного контенту та трансформації політичної комунікації в цифровому середовищі.

Окремий блок досліджень зосереджений на феномені інформаційних операцій у контексті гібридних конфліктів. У працях українських і зарубіжних авторів наголошується, що дезінформаційні кампанії дедалі частіше інтегруються у ширші стратегії політичного та військово-політичного тиску, поєднуючись із кібератаками, дипломатичними інструментами та економічним впливом [3, 4]. В цьому контексті досвід України розглядається як показовий приклад системного застосування інформаційної зброї.

Паралельно зростає корпус досліджень, присвячених застосуванню технологій штучного інтелекту для аналізу інформаційних потоків, виявлення бот-мереж, класифікації дезінформаційного контенту та моніторингу онлайн-дискусій. У звітах міжнародних організацій, аналітичних центрів і технологічних компаній ШІ позиціонується як перспективний інструмент під-

вищення ефективності протидії інформаційним загрозам [5]. Водночас науковці звертають увагу на обмеження таких підходів, зокрема на ризики помилкової класифікації, алгоритмічної упередженості, непрозорості моделей та можливого зловживання технологіями [6].

Попри зростання кількості публікацій наявні дослідження часто мають фрагментарний характер і зосереджуються або на технологічних аспектах штучного інтелекту, або на політичній природі дезінформації. Недостатньо опрацьованим залишається комплексний політологічний аналіз використання ІІІ саме як інструмента протидії політичній дезінформації та інформаційним операціям, з урахуванням безпекового, інституційного та етичного вимірів, що й зумовлює наукову новизну даного дослідження.

Виклад основного матеріалу. Політична дезінформація як інструмент політичної боротьби має тривалу історію, однак у цифрову епоху вона зазнала якісної трансформації, що змінила її характер, масштаби та наслідки. Якщо в традиційних медіасистемах поширення маніпулятивної інформації було обмежене редакційним контролем, часовими рамками та відносно стабільною структурою публічного простору, то сучасне цифрове середовище створило принципово нові умови для інформаційного впливу. Швидкість циркуляції повідомлень, низький поріг входу для створення контенту та глобальний характер онлайн-платформ сприяли перетворенню дезінформації на системний фактор політичного процесу.

Еволюція політичної дезінформації тісно пов'язана з трансформацією моделей політичної комунікації. У цифровому середовищі зникає чітке розмежування між виробниками, посередниками та споживачами інформації, що призводить до розмивання відповідальності за її достовірність. Дезінформаційні повідомлення дедалі частіше подаються не у формі очевидної неправди, а як фрагментарні, контекстуально викривлені або емоційно забарвлені наративи, які апелюють до наявних соціальних страхів, ідентичностей та політичних упереджень. Такий підхід ускладнює їхнє розпізнавання та підвищує ефективність маніпулятивного впливу [1].

Особливу роль у поширенні політичної дезінформації відіграють соціальні платформи та алгоритмічні системи управління контентом. Алгоритми рекомендацій, орієнтовані на максимізацію залученості користувачів, фактично стимулюють поширення поляризуючих і конфліктних повідомлень, незалежно від їхньої фактичної достовірності. У такому середовищі дезінформація отримує структурні переваги, оскільки часто є емоційно насиченою, простою для сприйняття та здатною швидко викликати реакцію аудиторії. Додатково цьому сприяють механізми мікротаргетингу, які дозволяють адаптувати політичні повідомлення до конкретних соціальних, регіональних або ідеологічних груп, фрагментуючи публічний простір і ускладнюючи формування спільного інформаційного порядку денного [7].

У сучасних умовах політична дезінформація дедалі частіше використовується не ізольовано, а в межах цілеспрямованих інформаційних операцій. Такі операції являють собою комплексні кампанії, що поєднують поширення маніпулятивних наративів, використання автоматизованих акаунтів, координацію поведінки онлайн-спільнот і синхронізацію інформаційного впливу з іншими інструментами тиску. Інформаційні операції дедалі частіше інтегруються у ширші гібридні стратегії, поєднуючись із кібератаками, дипломатичними маневрами, економічними заходами та військово-політичними сигналами.

У межах гібридних загроз інформаційні операції спрямовані не лише на введення аудиторії в оману, а й на тривалий підрив довіри до державних інститутів, демократичних процедур та незалежних медіа. Їхнім кінцевим результатом стає зростання соціальної поляризації, демобілізація суспільства та ослаблення спроможності держави до ефективного реагування на кризи. Для України, яка перебуває в умовах тривалої гібридної агресії, інформаційні операції стали постійним і структурним елементом безпекового середовища, що істотно підвищує вимоги до систем протидії інформаційним загрозам.

У цьому контексті традиційні підходи до боротьби з дезінформацією – фактчекінг, спростування та регуляторні обмеження – виявляються недостатніми. Це зумовлює необхідність пошуку інновацій-

них технологічних рішень, здатних діяти в умовах високої динаміки, масштабності та складності сучасного інформаційного простору, що безпосередньо підводить до аналізу потенціалу штучного інтелекту як інструмента протидії інформаційним загрозам.

Стрімке ускладнення інформаційного середовища та зростання масштабів політичної дезінформації зумовили обмеженість традиційних підходів до її виявлення та нейтралізації. У цьому контексті штучний інтелект (ШІ) розглядається як один із ключових технологічних інструментів, здатних забезпечити системну, масштабовану та оперативну протидію інформаційним загрозам. Його застосування дає змогу перейти від реактивних моделей реагування до проактивного аналізу інформаційного простору та раннього виявлення маніпулятивних кампаній.

Основу сучасних рішень у сфері протидії дезінформації становлять алгоритми машинного навчання та обробки природної мови (NLP), які здатні автоматично аналізувати великі обсяги текстового, візуального та аудіоконтенту. Такі системи використовуються для ідентифікації характерних ознак дезінформаційних повідомлень, зокрема повторюваних наративів, емоційної поляризації, атипових мовних конструкцій та нетипових моделей поширення інформації [8]. На відміну від ручного аналізу ШІ-інструменти забезпечують безперервний моніторинг інформаційних потоків у режимі реального часу, що є критично важливим у динамічному цифровому середовищі.

Окремим напрямом застосування штучного інтелекту є виявлення координованої неавтентичної поведінки, включно з діяльністю бот-мереж та скоординованих груп акаунтів. Аналіз мережових зв'язків, часових патернів публікацій і синхронізованих реакцій дає змогу ідентифікувати штучно створену активність, яка імітує громадську думку. Такі інструменти особливо ефективні у виявленні прихованих інформаційних операцій, де маніпуляція здійснюється не через окремі повідомлення, а через системне формування інформаційного фону.

Важливим аспектом використання ШІ є можливість прогнозування розвитку дезінформаційних кампаній. На основі аналізу попередніх кейсів та накопичених даних алгоритми можуть

моделювати ймовірні сценарії поширення наративів, оцінювати їхній потенційний вплив на різні аудиторії та визначати найбільш уразливі сегменти інформаційного простору. Це створює підґрунтя для превентивних заходів, спрямованих на зниження ефективності інформаційних атак ще до досягнення ними критичної фази.

Водночас використання штучного інтелекту в протидії інформаційним загрозам не обмежується виключно технічним аналізом контенту. Перспективним напрямом є інтеграція ШІ-рішень у системи стратегічних комунікацій держави, де алгоритми можуть допомагати у формуванні контрнاراتивів, оцінці їхньої сприйнятливості аудиторіями та оптимізації каналів поширення офіційної інформації. За такого підходу ШІ виступає не лише як інструмент фільтрації, а як елемент ширшої комунікаційної екосистеми [9].

Водночас ефективність застосування штучного інтелекту значною мірою залежить від якості навчальних даних, інституційної спроможності та міжсекторальної взаємодії. Алгоритми, навчені на упереджених або обмежених наборах даних, можуть відтворювати помилки або демонструвати низьку адаптивність до нових форм дезінформації. Це зумовлює необхідність постійного оновлення моделей, залучення експертного контролю та поєднання автоматизованих рішень із аналітичною роботою фахівців.

Таким чином, штучний інтелект формує технологічну основу сучасних систем протидії інформаційним загрозам, однак його застосування потребує комплексного підходу, що враховує як технічні можливості, так і політичні, етичні та інституційні обмеження. Саме ці аспекти визначають умови ефективного використання ШІ у сфері інформаційної безпеки та зумовлюють подальший аналіз ризиків і викликів, пов'язаних з його впровадженням.

Попри значний потенціал штучного інтелекту в протидії політичній дезінформації та інформаційним операціям його впровадження супроводжується низкою суттєвих ризиків, обмежень та етичних викликів. Вони зумовлені як технічними характеристиками самих алгоритмів, так і специфікою політичного та соціаль-

ного контексту, в якому ці технології застосовуються. Ігнорування таких аспектів може не лише знизити ефективність ШІ-рішень, а й створити додаткові загрози для демократичних інститутів і прав людини.

Одним із ключових обмежень є проблема алгоритмічної упереженості. Системи штучного інтелекту навчаються на історичних даних, які часто відображають наявні політичні, культурні або соціальні дисбаланси. Внаслідок цього алгоритми можуть некоректно класифікувати контент, помилково маркуючи легітимні політичні висловлювання як дезінформацію або, навпаки, не виявляти маніпулятивні матеріали, замасковані під нейтральні повідомлення. У політичному середовищі такі помилки мають особливо чутливий характер, оскільки можуть впливати на свободу слова та рівність доступу до публічної дискусії.

Іншим суттєвим викликом є обмежена прозорість і пояснюваність алгоритмічних рішень. Багато сучасних моделей, зокрема глибокі нейронні мережі, функціонують як «чорні скриньки», що ускладнює розуміння логіки прийняття рішень. У контексті протидії дезінформації це створює проблеми підзвітності: складно визначити, чому певний контент було заблоковано або обмежено, і хто несе відповідальність за такі дії. Відсутність пояснюваності підриває довіру суспільства до технологічних рішень і може стати підґрунтям політичних спекуляцій.

Використання ШІ також породжує ризики надмірної автоматизації процесів управління інформаційним простором. Делегування ключових рішень алгоритмам може призвести до зниження ролі людського судження, контекстуального аналізу та професійної етики. У ситуаціях політичної поляризації або кризових подій автоматизовані системи можуть реагувати надто жорстко або, навпаки, чутливо до специфіки ситуації, що підвищує ймовірність помилкових рішень.

Окремої уваги потребують етичні дилеми, пов'язані з балансом між забезпеченням інформаційної безпеки та захистом прав людини. Масовий моніторинг цифрових комунікацій, навіть із метою протидії дезінформації, може розглядатися як втручання у приватність і свободу вираження поглядів [4]. За відсутності чітких

правових обмежень і механізмів контролю ШІ-інструменти потенційно можуть бути використані для політичного тиску, цензури або переслідування опозиційних голосів.

Не менш важливим є ризик адаптації дезінформаційних акторів до використання штучного інтелекту. Інформаційні операції дедалі частіше використовують ті самі технології машинного навчання для створення більш переконливого фейкового контенту, зокрема *deepfake*-матеріалів, автоматизованих текстів та синтетичних акаунтів. Це формує своєрідну «технологічну гонку», в якій перевага ШІ-рішень є тимчасовою та потребує постійного оновлення й удосконалення [5].

Таким чином, використання штучного інтелекту в сфері протидії політичній дезінформації є одночасно необхідним і проблематичним. Ефективність таких рішень визначається не лише технічними параметрами, а й наявністю етичних принципів, прозорих процедур та демократичного контролю. Усвідомлення й системне врахування зазначених ризиків створює передумови для формування збалансованих стратегій застосування ШІ, орієнтованих на захист як інформаційної безпеки, так і фундаментальних демократичних цінностей.

Інтеграція штучного інтелекту в державну політику протидії політичній дезінформації та інформаційним операціям відкриває нові можливості для підвищення ефективності управління інформаційною безпекою в умовах цифрової трансформації. Водночас цей процес потребує системного підходу, що поєднує технологічні інновації з інституційними, правовими та етичними механізмами регулювання. ШІ у даному контексті має розглядатися не як автономне рішення, а як складова комплексної політики, зорієнтованої на зміцнення демократичної стійкості.

Одним із ключових напрямів є інституціоналізація ШІ-рішень у структурі державного управління. Це передбачає створення або адаптацію спеціалізованих підрозділів, відповідальних за моніторинг інформаційного простору, аналіз дезінформаційних загроз та координацію міжвідомчих дій. Використання штучного інтелекту в таких структурах дає змогу об'єднати розрізнені джерела даних,

забезпечити їхню оперативну обробку та підтримати ухвалення управлінських рішень на основі доказів, а не інтуїтивних оцінок.

Важливою передумовою ефективної інтеграції ІІІ є формування чіткої нормативно-правової бази. Державна політика має визначати межі допустимого використання алгоритмічних інструментів, критерії їхньої прозорості та механізми підзвітності. Особливу увагу слід приділяти регулюванню питань доступу до даних, захисту персональної інформації та процедур оскарження рішень, ухвалених за участі ІІІ. У цьому сенсі перспективним є запровадження принципів «пояснюваного штучного інтелекту», які дають змогу забезпечити баланс між ефективністю та демократичним контролем.

Не менш значущим напрямом є розвиток людського капіталу в сфері інформаційної безпеки. Інтеграція ІІІ у державну політику вимагає підготовки фахівців, здатних не лише користуватися відповідними технологіями, а й критично оцінювати їхні результати, розуміти обмеження алгоритмів і враховувати соціально-політичний контекст. Це зумовлює потребу в міждисциплінарних освітніх програмах, що поєднують знання з політології, інформаційних технологій, права та етики [9].

Окрему роль у впровадженні ІІІ відіграє партнерство між державою, приватним сектором та громадянським суспільством. Значна частина інновацій у сфері штучного інтелекту зосереджена в технологічних компаніях та дослідницьких центрах, тоді як громадські організації часто мають унікальну експертизу в сфері медіаграмотності та аналізу дезінформації. Налагодження стійкої взаємодії між цими акторами дає змогу підвищити адаптивність державної політики та зменшити ризики монополізації контролю за інформаційним простором.

У перспективі важливим елементом державної політики має стати інтеграція ІІІ-інструментів у ширші стратегії стратегічних комунікацій [10]. Йдеться не лише про виявлення та нейтралізацію дезінформації, а й про активне формування довіри до офіційних джерел інформації, підтримку якісного публічного діалогу та підвищення рівня медіаграмотності населення. За такого підходу штучний інтелект може використовуватися для аналізу суспільних

настроїв, оцінки ефективності комунікаційних кампаній та адаптації повідомлень до різних аудиторій без маніпулятивного впливу.

Водночас перспективи інтеграції ШІ значною мірою залежать від здатності держави забезпечити тривалу стратегічну послідовність. Фрагментарне або ситуативне використання технологій не дозволяє досягти стійких результатів і може посилити недовіру суспільства. Тому застосування штучного інтелекту в протидії дезінформації має ґрунтуватися на чітко сформульованих цілях, системі індикаторів ефективності та механізмах постійного перегляду політики з урахуванням динаміки інформаційних загроз.

Отже, перспективи інтеграції штучного інтелекту в державну політику протидії дезінформації полягають у переході від ізольованих технологічних рішень до комплексної моделі управління інформаційною безпекою. Така модель поєднає інноваційні інструменти з демократичними принципами, забезпечуючи не лише захист інформаційного простору, а й зміцнення інституційної стійкості та суспільної довіри.

Висновки та перспективи подальших досліджень. Політична дезінформація та інформаційні операції в умовах цифрової епохи набули системного, багаторівневого та динамічного характеру, перетворившись на один із ключових чинників впливу на політичні процеси, громадську думку та функціонування демократичних інститутів. Аналіз еволюції дезінформаційних практик свідчить, що сучасні інформаційні загрози дедалі частіше ґрунтуються на поєднанні технологічних інструментів, алгоритмічного поширення контенту та цілеспрямованих інформаційних операцій, інтегрованих у ширший контекст гібридних конфліктів.

Штучний інтелект виступає одним із найбільш перспективних інструментів протидії інформаційним загрозам завдяки здатності аналізувати великі масиви даних, виявляти приховані патерни та забезпечувати раннє виявлення дезінформаційних кампаній. Використання алгоритмів машинного навчання, аналізу соціальних мереж і прогнозування інформаційних впливів створює підґрунтя для переходу від реактивних до превентивних моделей управління інформаційною безпекою. Водночас ефективність таких рішень значною мірою залежить від якості даних, інституційної спромож-

ності та поєднання автоматизованого аналізу з експертною оцінкою.

Водночас існує низка ризиків і обмежень, пов'язаних із використанням штучного інтелекту в політичній сфері. Алгоритмічна упередженість, обмежена прозорість рішень, етичні дилеми щодо свободи слова та приватності, а також технологічна адаптація дезінформаційних акторів формують складне поле викликів. Це зумовлює необхідність критичного та відповідального підходу до впровадження ШІ, який враховуватиме не лише технічну ефективність, а й демократичні цінності та принципи верховенства права.

Щодо перспектив інтеграції штучного інтелекту в державну політику протидії дезінформації то успішне застосування ШІ можливе лише за умови його інституціоналізації в межах цілісної політики інформаційної безпеки, розвитку нормативно-правового регулювання, інвестицій у людський капітал та налагодження партнерства між державою, технологічним сектором і громадянським суспільством.

Оте, по-перше, політична дезінформація в цифрову епоху є комплексною загрозою, яка не може бути ефективно нейтралізована виключно технологічними або суто регуляторними засобами; по-друге, штучний інтелект має значний потенціал у сфері протидії інформаційним операціям, однак його використання потребує чітких етичних і правових обмежень та постійного людського контролю; по-третє, інтеграція ШІ у державну політику має відбуватися в межах тривалої стратегії щодо зміцнення демократичної стійкості та суспільної довіри.

Перспективи подальших наукових розвідок у даному напрямі пов'язані з поглибленим аналізом ефективності конкретних ШІ-інструментів у різних політичних контекстах, порівняльними дослідженнями національних моделей протидії дезінформації, а також вивченням впливу алгоритмічного регулювання на політичну участь і свободу вираження поглядів. Окремим перспективним напрямом є дослідження взаємодії між штучним інтелектом і стратегічними комунікаціями держави, а також розробка методологій

оцінювання соціальних і політичних наслідків використання ШІ у сфері інформаційної безпеки.

1. Лорян Р. Штучний інтелект як суперінструмент для дезінформації та пропаганди. *ОПОРА*. 30.01.2023 р. URL: https://www.oporaua.org/polit_ad/shtuchnii-intelekt-iaksuperinstrument-dlia-dezinformatsiyi-ta-propagandi-24507.
2. Штучний інтелект і дезінформація: можливості та ризики в умовах війни. *Центр стратегічних комунікацій та інформаційної безпеки SPRAVDI*. 05.04.2023 р. URL: <https://spravdi.gov.ua/shtuchnyi-intelekt-i-dezinformacziya-mozhlyvosti-ta-ryzyky-v-umovah-vijny/>.
3. Бучко У. Виявлення та протидія дезінформації з допомогою штучного інтелекту. Природничі та гуманітарні науки. *Актуальні питання: матер. VIII Міжнародної студентської наук.-техн. конференції (м. Тернопіль, 24–25 квітня 2025 р.)*. Тернопіль: ТНТУ, 2025. С. 136–138. URL: <http://elartu.tntu.edu.ua/handle/lib/49349>.
4. Драбюк С. Штучний інтелект і пропаганда та дезінформація: основні виклики. *Вісник Ужгородського національного університету*. Сер. Право. 2025. Вип. 90. Ч. 5. С. 348–356. URL: <https://visnyk-juris-uzhnu.com/wp-content/uploads/2025/09/47-4.pdf>.
5. Shoaib M. R., Wang Z., Ahvanooe M. T. et al. Deepfakes, Misinformation, and Disinformation in the Era of Frontier AI, Generative AI, and Large AI Models. Cornell University. 29 Nov 2023. URL: <https://arxiv.org/abs/2311.17394>.
6. Walker C. P., Schiff D. S., Jackson Schiff K. Merging AI Incidents Research with Political Misinformation Research: Introducing the Political Deepfakes Incidents Database, 2024. URL: <https://arxiv.org/abs/2409.15319>.
7. Doshi J., Novacic I., Fletcher C. et al. Sleeper Social Bots: A New Generation of AI Disinformation Bots Are Already a Political Threat. Cornell University. 7 Aug 2024. URL: <https://arxiv.org/abs/2408.12603>.
8. Дзюбановська Н.В. Штучний інтелект для виявлення фейкової інформації: міжнародні тренди і можливості імплементації. *Інноваційна економіка*. 2024. № 3. С. 154–160. URL: <https://inneco.org/index.php/innecoua/uk/article/view/1288>.
9. Solopova V. From Trust to Truth: Actionable Policies for the Use of AI in Fact-Checking in Germany and Ukraine. *Cornell University*. 24 Mar 2025. URL: <https://arxiv.org/abs/2503.18724>.
10. The Digital Services Act. European Union. Brussels, 2023. URL: <https://digital-strategy.ec.europa.eu/en/policies/digital-services-act-package>.

References

1. Lorian R. Shtuchnyi intelekt yak superinstrument dlia dezinformatsii ta propahandy. *OPORA*. 30.01.2023 . URL: https://www.oporaua.org/polit_ad/shtuchnii-intelekt-iaksuperinstrument-dlia-dezinformatsiyi-ta-propagandi-24507.
2. Shtuchnyi intelekt i dezinformatsiia: mozhlyvosti ta ryzyky v umovakh viiny. *Tsentr stratehichnykh komunikatsii ta informatsiinoi*

bezpeky SPRAVDI. 05.04.2023. URL: <https://spravdi.gov.ua/shtuchnyj-intelekt-i-dezinformacziya-mozhlyvosti-ta-ryzyky-v-umovah-vijny/>. 3. Buchko U. Vyavlennia ta protydiia dezinformatsii z dopomohoiu shtuchnoho intelektu. Pryrodnychi ta humanitarni nauky. *Aktualni pytannia: Proceedings of the VIII International Student Scientific and Technical Conference (Ternopil, 24–25 April 2025)*. Ternopil: TNTU, 2025. P. 136–138. URL: <http://elartu.tntu.edu.ua/handle/lib/49349>. 4. Drabiuk S. Shtuchnyi intelekt i propahanda ta dezinformatsiia: osnovni vyklyky. *Visnyk Uzhhorodskoho natsionalnoho universytetu*: Ser. Law. 2025. Issue 90. P. 5. P. 348–356. URL: <https://visnyk-juris-uzhnu.com/wp-content/uploads/2025/09/47-4.pdf>. 5. Shoaib M. R., Wang Z., Ahvanooy M. T. et al. Deepfakes, Misinformation, and Disinformation in the Era of Frontier AI, Generative AI, and Large AI Models. Cornell University. 29 Nov 2023. URL: <https://arxiv.org/abs/2311.17394>. 6. Walker C. P., Schiff D. S., Jackson Schiff K. Merging AI Incidents Research with Political Misinformation Research: Introducing the Political Deepfakes Incidents Database. Cornell University. 5 Sep 2024. URL: <https://arxiv.org/abs/2409.15319>. 7. Doshi J., Novacic I., Fletcher C. et al. Sleeper Social Bots: A New Generation of AI Disinformation Bots Are Already a Political Threat. Cornell University. 7 Aug 2024. URL: <https://arxiv.org/abs/2408.12603>. 8. Dziubanovska N. V. Shtuchnyi intelekt dlia vyavlennia feikovoï informatsii: mizhnarodni trendy i mozhlyvosti implementatsii. *Innovatsiina ekonomika*. 2024. No. 3. P. 154–160. URL: <https://inneco.org/index.php/innecoua/uk/article/view/1288>. 9. Solopova V. From Trust to Truth: Actionable Policies for the Use of AI in Fact-Checking in Germany and Ukraine, *Cornell University*. 24 Mar 2025. URL: <https://arxiv.org/abs/2503.18724>. 10. The Digital Services Act. European Union. Brussels, 2023. URL: <https://digital-strategy.ec.europa.eu/en/policies/digital-services-act-package>.

But Sergii. Artificial intelligence as a tool to counter political disinformation and information operations

The article is devoted to the analysis of the role of artificial intelligence in countering political disinformation and information operations in the context of the digital transformation of society. It is argued that modern information threats are systemic and dynamic in nature, combine technological, political and social dimensions and are increasingly used as a tool to influence political processes and public opinion. In this context, traditional approaches to ensuring information security are insufficient, which actualizes the need for the application of innovative technological solutions.

The article considers the evolution of political disinformation in the digital age, the role of social platforms and algorithmic content distribution, as well as the place of information operations in the structure of hybrid threats. The potential of artificial intelligence as a tool for detecting, analyzing and predicting

disinformation campaigns is analyzed, in particular through the use of machine learning methods, natural language processing and network behavior analysis.

Particular attention is paid to the risks and limitations of the use of artificial intelligence in the political sphere, including problems of algorithmic bias, opacity of decisions, ethical dilemmas regarding freedom of speech and privacy, as well as the adaptation of disinformation actors to new technologies. It is substantiated that the effective use of artificial intelligence requires a combination of technological capabilities with legal, institutional and ethical control mechanisms.

As a result, the prospects for integrating artificial intelligence into the state policy of countering disinformation as a component of a comprehensive information security strategy focused on strengthening democratic stability, institutional capacity and public trust are formulated.

Key words: artificial intelligence (AI); political disinformation; information operations; information security; hybrid threats; algorithmic systems; social media; information state policy; digital technologies; democratic sustainability.

Дата першого надходження рукопису до редакції: 30.01.2026 р.

Дата прийнятого до друку рукопису після рецензування: 15.02.2026 р.

Дата публікації: 03.03.2026 р.