

УДК 347.77:316.472.4:004.8

DOI <https://doi.org/10.32782/chc.v057.2025.17>**Омельчук Олександр Сергійович,**

кандидат юридичних наук, доцент,

доцент кафедри цивільного права

Національного університету «Одеська юридична академія»

ORCID ID: 0000-0002-0082-3619

ЦИВІЛЬНО-ПРАВОВИЙ АСПЕКТ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В СОЦІАЛЬНИХ МЕРЕЖАХ

Постановка проблеми. На сьогодні перспективи використання штучного інтелекту (ШІ) викликають інтерес не лише в експериментальних та прикладних технічних галузях, соціальній та економічній сфері. Одночасно з тим як наукова спільнота працює над удосконаленням технологій та вирішенням філософських питань стосовно ШІ, національні законодавці та міжнародні організації шукають підходи до визначення юридичної природи ШІ та надання певного правового режиму результатам діяльності ШІ. Як відзначає А. Юшина, на сьогоднішній день в наукових колах актуалізується дискусія навколо правомірності використання ШІ творів для навчання, які охороняються авторським правом, оскільки на основі їх аналізу ШІ створює нові схожі об'єкти, а автори творів не отримують винагороду за таке використання [16]. Крім того, авторка вбачає проблемним питання, чи порушується авторське право під час генерування нового контенту ШІ, адже авторські твори використовуються для машинного навчання.

Сфери та методи використання ШІ виглядають невичерпними на даний момент. Все ширше нейромережі використовуються і в звичайному житті, для спілкування, передачі інформації та створення непрофільного розважального аудіовізуального контенту. Вітчизняні дослідники проблем правового регулювання штучного інтелекту також вказують на проблему ризиків дезінформації, зумовлених поширенням ШІ (приклади останньої можуть з'являтися у текстах, згенерованих за допомогою ШІ), а, поряд з цим, і на проблему потенційного збільшення кількості випадків порушення авторського права. Таким чином, вважаємо актуальним для наукового осягнення використання штучного інтелекту в соціальних мережах в цивільно-правовому аспекті, а **метою даного дослідження є визначення перспектив державного регулювання правових відносин в сфері використання штучного інтелекту в соціальних мережах та мережі Інтернет загалом.**

Аналіз останніх досліджень і публікацій. Проблематиці надання юридичної оцінки діяльності з використання штучного інтелекту в мережі Інтернет присвячені роботи таких дослідників як Н. Авер'янова, О. Баранов, А. Бірюкова, Т. Воропаєва, Ю. Гайдай, Г. Колісникова, І. Колеснікова та багатьох інших. Однак, враховуючи дискусійність, певну фрагментарність та новаторський характер наукових розвідок з проблематики теми слід визнати недостатність наявних таких наукових розвідок, а також перспективність та нагальність наукових розробок чинних та можливих правових механізмів нівелювання або обмеження негативного впливу штучного інтелекту в соціальних мережах.

Виділення невирішених раніше частин загальної проблеми. Аналіз наявного масиву наукових праць дає змогу констатувати відсутність цільного систематизованого дослідження цивільно-правових відносин, що виникають в сфері використання штучного інтелекту в соціальних мережах. Крім того потребують доопрацювання питання відповідальності за шкідливе використання ШІ в соціальних мережах.

Виклад основного матеріалу. ШІ вже сьогодні здатен обробляти велику кількість даних та навчатися, щоб вирішувати задачі, генерувати текст, фото, відео, аудіо, код, інформацію та сукупність цих об'єктів. ШІ вже зараз здатний аналізувати дані, змішувати їх і генерувати на їхній основі новий та неповторний результат [12].

За прогнозами експертів найближчим часом ШІ перейде від узагальнених продуктів до рішень, орієнтованих на розв'язання конкретних завдань. Окрім того, ШІ проникатиме в різні сфери життя й бізнесу – від зміни клімату до охорони здоров'я [1]. Більшість людей не усвідомлюють, що вже зараз використовують технології штучного інтелекту у щоденному житті, в освіті та медицині, комунікаціях та соціальній сфері, економіці та юридичній сфері [9].

Ю. Гайдай визначає наступні ризики етичного характеру щодо використання ШІ в мережі Інтернет:

- ризики для особистих даних і приватності. Мовні моделі, по-перше, можуть робити дуже багато узагальнень, поєднань інформації і висновків з отриманих даних, які можуть бути недоступними для людини, що аналізує інформацію. Наприклад, з наборів даних про поведінку людини у соцмережах та інтернеті загалом, тих же запитів до чатботів великих мовних моделей, можна робити глибокі висновки щодо її психічного стану, звичок, розкладу дня тощо. Поєднання великих масивів даних та їх аналіз ШІ створюють додаткові загрози приватності, які не очевидні для користувачів, бо окремо ці дані не створюють загроз.

- ризики сегрегації – мовні моделі, за рахунок того, як саме і на основі яких даних навчаються, можуть надавати перевагу окремим соціальним групам. Вони будуть краще працювати з цими групами, а іншим приділяти менше уваги і маргіналізувати їх, вилучати певні меншини (етнічні або інші) та відтворювати людські упередження щодо цих меншин;

- ризики неточності і неправдивості інформації – мовні моделі не дуже добре відрізняють факти від вигадок, але при цьому створюють враження, що вони є достовірними співрозмовниками, що може створювати багато загроз;

- безпекові ризики – мовні моделі дуже зручно використовувати для масової дезінформації, причому персоналізованої дезінформації «з людським обличчям», заточеної під конкретні групи людей, наприклад, шляхом коментування у соцмережах та у медіа;

- загроза правам людини через використання штучного інтелекту державними органами – передбачається насамперед здатність опрацьовувати величезні масиви даних, роблячи узагальнення і висновки щодо окремих людей. Це можливості розпізнавання, а отже, відстежування людей як фізично, так і деанонімізація їхньої онлайн-поведінки. Це дає державним органам потужні інструменти втручання у приватність, а безконтрольне застосування ШІ може призвести до ухвалення неправомірних чи хибних рішень стосовно громадян – застосування санкцій, відмови у доступі до послуг, кредиту, працевлаштування тощо [3].

Одним з найбільш поширених і зрозумілих на сьогоднішній день видів ШІ, що на сьогоднішній день використовуються в мережі Інтернет є гене-

ративний штучний інтелект. Генеративний ШІ – це інноваційний напрямок штучного інтелекту, що дозволяє створювати унікальний контент: тексти, зображення та аудіо. Він відрізняється від традиційного ШІ тим, що не просто аналізує або класифікує дані, а генерує новий унікальний контент [15].

«Творчість» штучного інтелекту вже сьогодні широко використовується та непогано монетизується. Однак, слід однак зазначити, що авторські права на творчість алгоритмів – поки сіра зона світу юриспруденції [7].

Два роки тому набрав чинності Закон України «Про авторське право і суміжні права». Згідно положень даного закону результат роботи штучного інтелекту підпадає під правове регулювання як неоригінальний об'єкт, згенерований комп'ютерною програмою, що охороняється правом особливого роду (*sui generis*). Однак, як відмічають дослідники, питання авторських та суміжних прав, пов'язаних з штучним інтелектом, є складними і потребують уточнення у законодавстві. На сьогоднішній день в більшості країн законодавці продовжують вивчати цю проблему і шукати оптимальні рішення, які враховували б потреби розвитку технологій та гарантували захист прав власників інтелектуальної власності [2, с. 14].

Значна частина «побутового» використання ШІ припадає на соціальні мережі. Так, соціальні мережі, з одного боку, є платформою для вільного обміну інформацією та думками, проте, з іншого боку, вони можуть стати засобом поширення ненависті, дискримінації та маніпуляції [13].

Лідером із поширення фейків у соціальних мережах є Facebook. У Facebook найрозгалуженіша інструментальна мережа, що працює в двох напрямках: для збору даних і зміни поведінки людей – впливу чи нав'язування [8].

Наявні діджитал-інструменти дозволяють за допомогою ШІ розрізняти органічне та скоординоване розповсюдження контенту, виявляти автоматизовані системи розповсюдження спаму, оцінювати вплив на аудиторію різних облікових записів користувачів соцмереж, відрізняти ботів від реальних користувачів тощо [14].

Широкого застосування штучний інтелект набув і в основних месенджерах – Telegram, Whatsapp, Viber, Messenger. Серед основних напрямків його використання – створення чат-ботів з різноманітним функціоналом (генерування зображень та відео, збереження графічного та аудіовізуального контенту з окремих інтернет-ресурсів, пошук та підбір інформації тощо), використання можливо-

стей месенджера для генерації запитів до нейромереж та ШІ, використання останніх з метою введення в оману та шахрайських схем.

На сьогоднішній день можливості нейромереж генерувати голосові та відео-повідомлення, які є настільки реалістичними та неможливими для виявлення їх штучного походження, надають зловмисникам можливість використовувати дані технології для введення в оману користувача (або цілої групи користувачів) месенджера. Генеративний ШІ може навіть зробити застарілими найкращі на сьогодні заходи запобігання шахрайству, як-то голосова автентифікація чи «живі перевірки», коли зображення людини у реальному часі має збігатися із тим, що зберігається в системі [11]. На теперішній час значного поширення набрали шахрайські схеми, в яких знайома особа просить грошові кошти в борг або доступ до якогось додатку, але сама така особа не в курсі, що від її імені і, нерідко, її голосом відбувається спроба заволодіти чужими коштами. В умовах же інформаційної війни дані технології використовуються для отримання оперативних військових даних, дезінформації військовослужбовців, що однозначно створює загрозу державній безпеці.

Боротьба ж випадками шкідливого використання генеративного ШІ ускладнюється з наступних причин:

- подальша еволюція технології дозволяє обходити технологічні засоби протидії застосування ШІ – генеративні ШІ вже зараз здатні обманювати детектори ШІ;
- розвиток даної технології напругу пов'язаний з перевагами, які вона надає при створенні ефективного голосового помічника або в допомозі особам, що мають вади мови;
- ефективність боротьби напругу пов'язана з інформаційною гігієною користувачів, а отже й залежить від раніше означеної проблеми наявності значного масиву інформації про людину в соціальних мережах.

Окрім класичних соціальних мереж широке використання штучного інтелекту та нейромереж можемо споглядати й у випадку дейтинг-платформ та додатків. Під дейтингом зазвичай розуміють діяльність спеціалізованих онлайн-платформ, які використовуються користувачами для онлайн-знайомств. Основною метою користувачів дейтинг-платформ є пошук людини на основі своїх уподобань, для чого означені сервіси використовують різні механізми пошуку за певними критеріями.

Вже зараз з метою забезпечення ефективності надання комунікативних послуг дейтинг-сервіси застосовують сучасні технології, в тому числі ШІ. Так, наприкінці 2024 року компанія Match Group, яка володіє такими популярними додатками для знайомств як Tinder та Hinge, оголосила про плани інтегрувати генеративний ШІ у свої сервіси для покращення біографій та профілів користувачів [5]. Крім того ШІ та нейромережі використовуються дейтинг-платформами для оцінки сумісності пари, надання порад в спілкуванні та оформленні профілю, опосередкованого спілкування замість користувача, здійснення перекладу текстових та голосових повідомлень тощо.

Одночасно з цим штучний інтелект в дейтинг-сфері використовується також і з метою надання користувачу послуг персоналізованого спілкування з віртуальним партнером. Слід відмітити розвиток технологічних можливостей таких платформ, що зумовлюється стрімким збільшенням попиту на такий цифровий продукт. Так, сервіси на зразок Candy.ai та Kupid.ai надають можливість користувачу, вибравши статтю майбутнього партнера, створити вигаданого персонажа, який за допомогою штучного інтелекту дає можливість спілкуватися з ним голосовими повідомленнями [10].

Відкидаючи морально-етичні моменти відносин з цифровою особою слід однак звернути увагу на юридичний контекст, а саме на захист прав користувача подібних дейтинг-платформ. В першу чергу слід говорити про захист інформації, отриманої в межах комунікації з віртуальним партнером, власником платформи від розповсюдження та витоку, а також збереження режиму приватності щодо персональних даних та самого факту надання послуг, адже користувач може не бажати публічності факту надання таких послуг дейтинг-сервісом.

Сучасні додатки для знайомств вже збирають величезну кількість чутливої інформації, такої як геолокація, їхні вподобання, історію листування та навіть фото і відео [5]. Окрім традиційного неконтрольованого поширення таких даних загрозу становить також використання таких даних для створення небіологічних осіб, наділених штучним інтелектом, які, навчаючись на зазначених даних, можуть створювати цифрові клони реальної людини, тим самим користуючись чужою особистістю та вводячи в оману осіб, з якими спілкується в подальшому.

Використання таких цифрових копій осіб вже зараз зумовило проблему поширення випадків свідомого використання можливостей ШІ для

маніпулювання емоціями і шахрайських схем. На даний момент практично надзвичайно важко захистити свої права щодо такого «викрадення» особистості через складність доказування та відсутність законодавчих механізмів захисту від використання чужої особистості без дозволу особи.

Використання таких цифрових інструментів як нейромережі та інших видів штучного інтелекту в Інтернеті не обмежується означеними сферами та методами, адже сама сфера інтернет-відносин формується та змінюється ледь не щодня, а отже на тлі цього кількість напрямків застосування даних технологічних рішень збільшується в геометричній прогресії. Вважаємо за необхідне дослідити окремі приклади застосування ШІ, що актуалізувалися саме в умовах гібридної війни.

Функціональні можливості генеративного ШІ успішно використовуються при створенні так званих «емоційних фейків» в соціальних мережах. У соціальних мережах зараз поширюється багато дописів, які містять фото з ознаками генерації штучним інтелектом. На них зображені нібито українські військові з закличками привітати їх чи то з днем народження, чи то з весіллям. Також є схожі дописи про врятованих тварин, дітей або літніх людей [4].

Ледве не кожен в запропонованих механізмами соціальних мереж публікація стикається з зображеннями, на яких відображено надзвичайні результати творчості, що викликають захват, відображені непересічні ситуації людського горя, нестандартних ситуацій та несприятливих умов, а також з аудіовізуальним контентом, що відображає надзвичайно сентиментальні ситуації та закликає до проявлення емпатії (так звані пости «ніхто нас не вітає», «ви не повірите» або «ці брати»).

Завдання контенту таких публікацій викликати емоції та заохотити користувачів відреагувати, спіймати вразливу аудиторію в емоційну пастку і змусити поширити їх далі, тиснучи на почуття провини, викликаючи емпатію та співчуття.

Довгий час дані публікації сприймалися з гумором, адже, на перший погляд дані публікації, не загрожують шкідливими наслідками, а самі зображення, згенеровані ШІ, містять ряд кумедних для сприйняття артефактів, що очевидно свідчать про несправжність запропонованого фото та ситуації на ньому. Однак, дані публікації не просто не є нешкідливими, при подальшому вивченні можна визначити загрози, які несуть на собі такі емоційні фейки.

Так, докладний аналіз релевантних шляхів використання емоційних фейків свідчить про наступні мотиви, з яких публікують такі записи у Facebook, Instagram та TikTok:

- збір аудиторії для подальшого просування рекламних продуктів, пропагандистських наративів тощо. В даному випадку завданням публікації виступає заохочення користувачів соціальної мережі до взаємодії, тобто йде мова про використання емоцій для прикриття шкідливого змісту самої публікації. Розміщення в соціальній мережі емоційних фейків певного характеру (співчутливі, ушановуючі, схвальні) дають змогу виділити певну групу людей, схильних до проявлення емоцій та не здатних відрізнити справжні фотографії подій від згенерованих за допомогою ШІ. В подальшому відповідний користувацький ресурс може бути використаний для успішного просування сумнівних продуктів або пропагандистських тез, що особливо небезпечно в умовах війни. Емоційні фейки в даному випадку слугують маркером людської сентиментальності та можливості нав'язувати певним особам необхідної думки;

- ілюстрація підтримки певної позиції або провокація аудиторії – особливо актуально через можливість заміни первинного матеріалу (цифрового контенту) без анулювання реакцій користувачів у вигляді лайків та коментарів. Крім того дані публікації є відомим майданчиком для використання ботів при провокуванні дискусії та м'якого просування шкідливих наративів;

- витіснення дійсно важливого та актуального контенту шляхом заповнення рекомендацій емоційними фейками та подібним шкідливим контентом;

- просування негативних настроїв в суспільстві в межах ІПСО, формування негативного ставлення до керівництва Збройних сил, політичної системи держави; формування

- збір даних про користувачів з метою подальшого використання; шахрайство, шляхом спонукання до переходу за сторонніми посиланнями, фішинг тощо.

Висновки. Підсумовуючи вищезазначене, можна зробити висновок, що на сьогоднішній день використання штучного інтелекту окрім явних переваг несе значну загрозу порушення прав людини, може суперечити індивідуальним та суспільним інтересам, загрожує державній безпеці та є прямим інструментом ведення інформаційної війни.

Серед практичних способів захисту від шкідливого впливу ШІ в мережі Інтернет в умовах гібридної війни серед найбільш актуальних залишаються пропаганда інформаційної гігієни на державному рівні, боротьба з фейками, шкідливими емоційними згенерованими ШІ шляхом запровадження відповідальності за таку діяльність та формування адміні-

стративної та судової практики з даних питань. Крім того вважаємо слушним на законодавчому рівні встановити вимогу маркування зображень, створених ШІ, а також розробку концептуальних положень контролю за використанням ШІ розробниками додатків з метою уникнення випадків створення «бекдорів» для корисних та неправомірних цілей.

ЛІТЕРАТУРА:

1. Баловсяк Н. Прикладний штучний інтелект і заборона соцмереж для дітей. Чого чекати від 2025-го технологічного. Український тиждень. 2024. URL: <https://tyzhden.ua/prykładnyj-shtuchnyj-intelekt-i-zaborona-sotsmerezhi-dlia-ditej-choho-chekaty-vid-2025-ho-tekhnohichnoho/>
2. Вальчук Д. В. Авторське право штучного інтелекту: проблеми та шляхи вирішення. *Modern scientific journal* (Сучасний науковий журнал). 2024. №3 (1). С. 7–15.
3. Гайдай Ю. Тренди ШІ: які етичні загрози несе використання штучного інтелекту. URL: https://speka.media/trendy-si-yaki-etichni-zagrozi-nese-vikoristannya-shtuchnogo-intelektu-v4q3wp?utm_source=google&utm_medium=cpc&utm_campaign=20972733776&gad_source=1
4. Денисяк О. «У військового день народження, а його ніхто не привітав»: чим небезпечні ШІ-світлини у Facebook. URL: <https://hromadske.ua/suspilstvo/232022-u-viyskovoho-den-narodzennia-a-yoho-nikhto-ne-pryvitav-chym-nebezpechni-shi-svitlyny-u-facebook>
5. Довірити своє серце алгоритмам: як штучний інтелект змінює сучасний дейтинг і що з цим може бути не так? 2025. URL: <https://zaborona.com/doviryty-svoje-sercze-algorytmam-yak-shtuchnyj-intelekt-zminyuye-suchasnyj-dejtyng-i-shho-z-czym-mozhe-buty-ne-tak/>
6. Довірити своє серце алгоритмам: як штучний інтелект змінює сучасний дейтинг і що з цим може бути не так? 2025. URL: <https://zaborona.com/doviryty-svoje-sercze-algorytmam-yak-shtuchnyj-intelekt-zminyuye-suchasnyj-dejtyng-i-shho-z-czym-mozhe-buty-ne-tak/>
7. Дудко В. Як палає Кремль. Нейромережі генерують картини з тексту і допомагають збирати на ЗСУ. Чи може нейро-арт стати бізнесом. *Forbes Digital*. 2022. URL: <https://forbes.ua/inside/yak-palae-kreml-neyromerezh-generuyut-kartini-z-tekstu-i-dopomagayut-zbirati-na-zsu-chi-mozhe-neyroart-stati-biznesom-23082022-7861>
8. Ейсмунт В. Як створюють фейки у Facebook, YouTube, Telegram та Viber. URL: https://zaxid.net/statti_tag50974/
9. Живцова Л.І. Штучний інтелект: сутність та перспективи розвитку. Український журнал будівництва та архітектури. 2023. № 3. С. 67–71.
10. Коноплицький С. Майбутнє дейтингу – «Тіндер» для знайомств із ШІ-моделями. URL: <https://speka.media/maibutnje-dejtingu-tinder-dlya-znaiomstv-iz-si-modelyami-pk4lzd>
11. Кофлін Д. ШІ допомагає всім, навіть шахраям. Як штучний інтелект став новою зброєю кіберзлочинців та стимулював афери на \$8 млрд. *Forbes Digital*. URL: <https://forbes.ua/innovations/shi-dopomagaе-vsimi-navit-shakhrayam-yak-shtuchniy-intelekt-stimulyuvav-afery-na-8-mlrd-19092023-16102>
12. Спесивцева О. Право sui generis: як штучний інтелект адаптує під себе авторське право? *Юридична газета*. 2024. URL: <https://yur-gazeta.com/publications/practice/informaciyne-pravo-telekomunikaciyi/pravo-sui-generis-yak-shtuchniy-intelekt-adaptue-pid-sebe-avtorske-pravo.html>
13. Хоменко Ю.І. Мова ворожнечі в соціальних мережах: форми прояву та механізми протидії. *Запоріжжя : ЗНУ*, 2024. 49 с.
14. Штучний інтелект і дезінформація: можливості та ризики в умовах війни. *Укрінформ*. URL: <https://www.ukrinform.ua/rubric-technology/3691961-stuchnij-intelekt-i-dezinformacia-mozlivosti-ta-riziki-v-umovah-vijni.html>
15. Що таке генеративний штучний інтелект і де він застосовується? URL: <https://blog.colobridge.net/uk/2023/11/generative-artificial-intelligence-ua/>
16. Юшина А. Штучний інтелект: кому належать авторські й майнові права. URL: <https://mind.ua/openmind/20255019-shtuchnij-intelekt-komu-nalezhat-avtorski-j-majnovi-prava>

Омельчук Олександр Сергійович

Цивільно-правовий аспект використання штучного інтелекту в соціальних мережах

В межах дослідження розглянуто питання використання штучного інтелекту в соціальних мережах в цивільно-правовому аспекті та сформовано рекомендації практичних способів захисту від шкідливого впливу ШІ в мережі Інтернет в умовах гібридної війни.

З'ясовано, що одним з найбільш поширених і зрозумілих на сьогоднішній день видів ШІ, що на сьогоднішній день використовуються в мережі Інтернет є генеративний штучний інтелект.

Досліджено питання використання емоційних фейків в соціальних мережах. Відзначено що функціональні можливості генеративного ШІ успішно використовуються при створенні так званих «емоційних фейків» в соціальних мережах. У соціальних мережах зараз поширюється багато дописів, які містять фото з ознаками генерації штучним інтелектом.

В контексті використання штучного інтелекту в дейтинг-індустрії звернуто увагу на юридичний контекст, а саме на захист прав користувача подібних дейтинг-платформ. В першу чергу слід говорити про захист інформації, отриманої в межах комунікації з віртуальним партнером, власником платформи від розповсюдження та витоку, а також

збереження режиму приватності щодо персональних даних та самого факту надання послуг, адже користувач може не бажати публічності факту надання таких послуг дейтинг-сервісом.

Зроблено висновок, що на сьогоднішній день використання штучного інтелекту окрім явних переваг несе значну загрозу порушення прав людини, може суперечити індивідуальним та суспільним інтересам, загрожує державній безпеці та є прямим інструментом ведення інформаційної війни.

Підсумовано, що серед практичних способів захисту від шкідливого впливу ШІ в соціальних мережах та мережі Інтернет загалом в умовах гібридної війни серед найбільш актуальних залишаються пропаганда інформаційної гігієни на державному рівні, боротьба з фейками, шкідливими емоційними згенерованими ШІ шляхом запровадження відповідальності за таку діяльність та формування адміністративної та судової практики з даних питань. Запропоновано на законодавчому рівні встановити вимогу маркування зображень, створених ШІ, а також розробку концептуальних положень контролю за використанням ШІ розробниками додатків з метою уникнення випадків створення «бекдорів» для корисних та неправомірних цілей.

Ключові слова: штучний інтелект, нейромережі, соціальні мережі, месенджер, емоційні фейки, право інтелектуальної власності, інформаційна гігієна, ІПСО, маркування зображень, інформаційна війна, діпфейк, дейтинг-сервіси, викрадення особистості, синтетичний контент.

Omelchuk Oleksandr

Civil-legal aspect of the use of artificial intelligence in social networks

The study examined the issue of using artificial intelligence in social networks in the legal aspect and formulated recommendations for practical ways to protect against the harmful effects of AI on the Internet in the context of hybrid warfare.

It was found that one of the most common and understandable types of AI used on the Internet today is generative artificial intelligence.

The issue of using emotional fakes in social networks was studied. It was noted that the functional capabilities of generative AI are successfully used in creating so-called “emotional fakes” in social networks. Many posts are currently circulating on social networks that contain photos with signs of generation by artificial intelligence.

In the context of using artificial intelligence in the dating industry, attention was paid to the legal context, namely, the protection of the rights of users of such dating platforms. First of all, we should talk about protecting information received within the framework of communication with a virtual partner, the owner of the platform, from distribution and leakage, as well as maintaining the privacy regime regarding personal data and the very fact of providing services, because the user may not want the fact of providing such services by a dating service to be public.

It is concluded that today the use of artificial intelligence, in addition to obvious advantages, carries a significant threat of human rights violations, may contradict individual and public interests, threatens state security and is a direct tool for waging information warfare.

It is summarized that among the practical ways to protect against the harmful effects of AI in social networks and the Internet in general in the conditions of hybrid warfare, the most relevant remain the promotion of information hygiene at the state level, the fight against fakes, harmful emotional generated by AI by introducing liability for such activities and the formation of administrative and judicial practice on these issues. It is proposed to establish a requirement for labeling images created by AI at the legislative level, as well as the development of conceptual provisions for controlling the use of AI by application developers in order to avoid cases of creating “backdoors” for useful and illegal purposes.

Key words: artificial intelligence, social networks, messenger, emotional fakes, intellectual property law, information hygiene, IPSO, image labeling, information warfare, deepfake, dating services, identity theft, synthetic content.