

ЦИФРОВІ ДОКАЗИ У ПРОТИДІЇ ДЕЗІНФОРМАЦІЇ ТА ДІПФЕЙКАМ ЩОДО ВОЄННИХ ЗЛОЧИНІВ

DIGITAL EVIDENCE IN COUNTERING DISINFORMATION AND DEEPPAKES ABOUT WAR CRIMES

Плетньов О.В., к.ю.н., доцент,
завідувач спеціальної кафедри № 2

Національний юридичний університет імені Ярослава Мудрого
orcid.org/0000-0002-0033-3011

Коваленко Є.В., к.ю.н., доцент,
професор спеціальної кафедри № 2

Національний юридичний університет імені Ярослава Мудрого
orcid.org/0000-0001-8630-8951

Стаття присвячена дослідженню цифрових доказів у протидії дезінформації та дипфейкам щодо воєнних злочинів. Автори зазначають, що використання генеративного штучного інтелекту для створення дипфейків змінює використання таких вигадок у дезінформаційних кампаніях наступними трьома фундаментальними способами: кількість – комерційно доступні програми дозволяють масово, швидко та дешево створювати дипфейки; якість – якість дипфейків покращується, вони виглядають природніше, що ускладнює їх виявлення, а також підвищує їхню достовірність і переконливість; кваліфікація – хоча створення дипфейків майже не вимагає навичок, експертиза, необхідна для їх виявлення, стає дедалі масштабнішою; мобілізація – дипфейки також можна використовувати для мобілізації населення проти сил безпеки; підривна діяльність – дипфейки можуть бути використані для створення страху та невизначеності. Автори зазначають, що в перший рік російського вторгнення синтетичний аудіоконтент був відносно рідкістю в російських інформаційних операціях, спрямованих проти України, а зловмисники використовували інші медіаформати для цілей явної дезінформації. Ситуація почала змінюватися протягом другого року війни, включаючи деякі помітні приклади. Автори підкреслюють, що текст, згенерований штучним інтелектом, залишається проблемою для спільноти розвідки з відкритим кодом та боротьби з дезінформацією, оскільки загальнодоступні інструменти не можуть надійно виявляти його, особливо в текстах невеликого обсягу, таких як коментарі в соціальних мережах. Однак у деяких випадках текст, згенерований штучним інтелектом, можна ідентифікувати іншими способами. Автори рекомендують застосування двох підходів для боротьби з дипфейками: технічне виявлення – величезна різноманітність способів маніпулювання медіа робить малоімовірним появу автоматизованого виявлення; стратегія реагування – ефективна стратегія реагування охоплює: загальну обізнаність з проблемою, постійний рівень залучення до вирішення проблеми дипфейків у поєднанні з моніторингом ЗМІ та здатністю швидко виявляти та оцінювати потенційні фальсифікації з технічної точки зору.

Ключові слова: дипфейк, OSINT, інформаційна безпека, роль державних та недержавних суб'єктів, штучний інтелект, цифровий доказ, дезінформація.

The paper is devoted to the study of digital evidence in countering disinformation and deepfakes regarding war crimes. The authors note that the use of generative artificial intelligence to create deepfakes is changing the use of such fictions in disinformation campaigns in the following three fundamental ways: quantity – commercially available programs that allow the mass, rapid and cheap creation of deepfakes; quality – the quality of deepfakes is improving, they look more natural, which makes them more difficult to detect, and also increases their credibility and persuasiveness; qualification – while the creation of deepfakes requires almost no skill, the expertise required to detect them is becoming increasingly extensive; mobilization – deepfakes can also be used to mobilize the population against security forces; subversion – deepfakes can be used to create fear and uncertainty. The authors note that in the first year of the Russian invasion, synthetic audio content was relatively rare in Russian information operations against Ukraine, with attackers using other media formats for explicit disinformation purposes. The situation began to change during the second year of the war, including some notable examples. The authors emphasize that AI-generated text remains a challenge for the open-source intelligence and counter-disinformation community because publicly available tools cannot reliably detect it, especially in short-form text such as social media comments. However, in some cases, AI-generated text can be identified by other means. The authors recommend the use of two approaches to combat deepfakes: technical detection – the vast variety of ways media manipulation makes automated detection unlikely; response strategy – an effective response strategy encompasses: general awareness of the problem, a sustained level of engagement in addressing the deepfakes problem, combined with media monitoring and the ability to quickly detect and assess potential fraud from a technical perspective.

Key words: deepfake, OSINT, information security, the role of state and non-state actors, artificial intelligence, digital evidence, disinformation.

Постановка проблеми. З початку повномасштабного вторгнення в Україну в лютому 2022 року росія продовжує свої зусилля в багатьох сферах, спрямовані на підрив України, зокрема в інформаційному просторі. Арсенал тактики та методів росії залишається різноманітним, починаючи від підроблених відео до фальшивих сервісів перевірки фактів.

Паралельно з цим періодом також спостерігається величезне поширення інструментів на базі штучного інтелекту. Зокрема, сила генеративного штучного інтелекту (GAI) привернула увагу всього світу, оскільки платформи GAI знизили поріг входу, дозволяючи створювати витончений фальшивий або маніпульований контент у великих масштабах, чи то з метою дезінформації, чи то пропа-

ганди. Вищезазначене обумовлює актуальність обраної теми дослідження.

Аналіз останніх досліджень. Питання цифрових доказів проти дезінформації, дипфейків та ШІ-згенерованих матеріалів у висвітленні воєнних злочинів є об'єктом принципового інтересу для вітчизняних і зарубіжних вчених і практиків, зокрема, таких, як: Г. Авдєєва, О. Бугера, О. Дуфенюк, Г. Мамедов, В. Шевчук та ін. Проте надзвичайно актуальним залишається дослідження ролі цифрових доказів у викритті дипфейків, аудіо- та текстових фальсифікацій щодо воєнних злочинів.

Відповідно, **метою статті** є дослідження цифрових доказів у протидії дезінформації та дипфейкам щодо воєнних злочинів.

Виклад основного матеріалу. Діпфейк-технологія – це нещодавній технологічний розробок, який по суті дозволяє людям створювати відео про події, яких ніколи не було. Вона здається особливо добре пристосованою для поширення дезінформації та «фейкових новин» на платформах соціальних мереж та в інших місцях онлайн. Діпфейки також дуже добре підходять для використання в кібервійні [1; 5; 17].

Використання генеративного штучного інтелекту для створення діпфейків змінює використання таких вигадок у дезінформаційних кампаніях наступними фундаментальними способами.

1. Кількість – комерційно доступні програми дозволяють масово, швидко та дешево створювати діпфейки. Це дає змогу не лише державам, а й групам та окремим особам, які мають обмежені ресурси, проводити власні дезінформаційні кампанії у великих масштабах.

2. Якість – якість діпфейків покращується, вони виглядають природніше, що ускладнює їх виявлення, а також підвищує їхню достовірність і переконливість.

3. Кваліфікація – хоча створення діпфейків майже не вимагає навичок, експертиза, необхідна для їх виявлення, стає дедалі масштабнішою.

Ці події мають потенціал значно збільшити охоплення та ефективність дезінформації у XXI столітті.

Більше того, з розвитком технології діпфейків, обмеження на виробництво та ефективне використання діпфейків більше не будуть технічного характеру, а будуть пов'язані виключно з питанням креативності.

3. Параліч – діпфейки могли бути використані у вигляді сфабрикованих доказів для паралізування або розколу союзників. Такий підхід був запропонований напередодні вторгнення в Україну. Американські експерти з безпеки, наприклад, припустили, що росія планувала створити підроблені відеодокази українських воєнних злочинів проти російських громад, щоб виправдати напад на Україну. Такі відеодокази були б доречними для початку дискусії в європейських державах щодо законності перетину росією кордону для захисту російських меншин, що могло б запобігти негайній реакції на користь України. У цьому конкретному випадку підозрювалося не використання діпфейків, а радше традиційних фейків з використанням реквізиту та акторів. Діпфейки могли б сприяти фабрикуванню таких сцен у майбутньому. Створення або зміна свідчень очевидців та нібито автентичних записів наказів, виданих з порушенням міжнародного права, можуть бути згенеровані вже сьогодні.

4. Мобілізація – діпфейки також можна використовувати для мобілізації населення проти сил безпеки. Можна експлуатувати існуючі етнічні, культурні, соціальні чи релігійні розбіжності всередині суспільств та між ними. Наприклад, жорстоке поводження з релігійними символами можна імітувати, створюючи фотографії чи відео наруки або фальсифікуючи свідчення очевидців.

5. Підривна діяльність – діпфейки можуть бути використані для створення страху та невизначеності. Фальшиві відеозаписи політичних чи військових лідерів, які закликають до капітуляції, зневажливі зауваження щодо власних військовослужбовців, загиблих у бою або ставлять під сумнів сенс і мету військової операції, або ймовірно, деморалізують збройні сили. Так само, вигадкування звірств, скоєних власними силами проти цивільного населення, може бути використано для підризу народної підтримки збройних сил. Крім того, може бути створена величезна кількість графічних зображень та аудіозаписів, що ілюструють жахи війни, щоб запобігти мобілізації населення та заохотити до дезертирства [4; 16].

Відео формуються за допомогою технології штучного інтелекту та зазвичай містять поєднання реального та фальшивого контенту. Це робить їх більш реалістичними та переконливими, ніж відео, повністю створені за

допомогою штучного інтелекту. Наприклад, технологія діпфейків може взяти справжнє відео та поміняти місцями обличчя двох людей у цьому відео або змінити рухи губ так, що людина, здається, каже щось інше, ніж те, що вона сказала спочатку.

Науковці та практики стверджують, що фальшиві відео можна створювати набагато легше та швидше за допомогою технології діпфейків, ніж за допомогою попередніх методів.

Існує багато випадків, коли наявність діпфейків викликала сумніви або плутанину. Вражаюче, що інколи люди звинувачують справжні відео у підробці. Було виявлено навіть докази того, що люди втрачали віру в усі відео, причому деякі люди підтримували теорії про те, що світові лідери померли, а їх замінили діпфейки [1; 17].

Діпфейки з'явилися майже одразу після повномасштабного вторгнення, що дозволило росії використовувати штучний інтелект для маніпулювання вже існуючими відео. 16 березня 2022 року проросійські актори зламали український медіа-сайт, опублікувавши невдало зроблений діпфейк президента України Володимира Зеленського.

Так, увечері 18 лютого 2022 року відвідувачів українського новинного веб-сайту зустріло знайоме видовище, відео з промовою Президента. Хоча схожість і була, обличчя здавалося трохи несинхронізованим з головою українського президента.

У відео фейкова версія Зеленського стверджувала, що він пішов у відставку, втік з Києва та закликав Збройні Сили України здатися та врятувати свої життя. Таким чином, Президенти начебто оголосив про закінчення війни, хоча більшість українців знала, що це фейк. Це було відео з фейком. Поки це відбувалося в Інтернеті, тикер внизу екрана в прямому ефірі каналу читав те саме повідомлення. У ньому знову ж таки неправдиво стверджувалося, що Україна капітулює.

Відео отримало значне поширення через Telegram; наратив також поширився в ефірі після російського злому новинної системи телеканалу «Україна 24», яку було змінено, щоб повторити те саме повідомлення. Надзвичайно низька якість відео чітко дала зрозуміти, що це діпфейк, тому проросійські джерела в Telegram почали називати його фейком, одночасно вихваляючись, що йому все ж вдалося змусити деяких українських військових здатися [4; 16].

У листопаді 2023 року в межах проросійських джерел було створено три відео, на яких тодішній головнокомандувач Збройних сил України генерал Валерій Залужний нібито звинувачував Зеленського у вбивстві свого помічника та попереджав про власну можливу передчасну смерть. Хоча сфальсифіковані відео мали видимі ознаки, такі як незграбні жести та міміка, вони, тим не менш, демонстрували, як еволюціонували діпфейки з часів фейку Зеленського на початку війни.

Більш цілеспрямований випадок обману, спричиненого штучним інтелектом, мав місце у формі витоку дзвінка в Zoom між членами Міжнародного легіону в Україні та неавтентичною версією колишнього президента України Петра Порошенка. Під час розмови хакери використали підроблені відеозаписи Порошенка, щоб обманом змусити членів легіону погодитися з провокаційними коментарями про Зеленського. Kyiv Independent встановила, що за цією кампанією стоять російські дезінформаційні актори, і зазначила, що нібито «погане з'єднання» під час розмови в Zoom дозволило виправити артефакти на відеозаписах Порошенка. Кампанія продемонструвала ефективність технологій на основі штучного інтелекту в таргетуванні невеликої групи людей за допомогою цілеспрямованої кампанії [3; 18].

У подібному випадку стверджувалося, що на відео зображені бійці 117-ї окремої бригади територіальної оборони, які звертаються до українського генерала Валерія Залужного

з проханням усунути чинний уряд. Згодом бригада спростувала правдивість цього відео та опублікувала заяву, що люди, зображені на відео, ніколи не були частиною бригади. Український центр протидії дезінформації наголосив, що у відео використовувалися дипфейк-технології.

Хоча вищезгадані відео, зроблені штучним інтелектом, демонструють прогрес технології та її потенціал для поширення дезінформації, зазначені приклади були вузько орієнтовані на певну аудиторію. Тим часом залишається таким же легким переробити реальні кадри та представити їх у хибному контексті або навіть залучити акторів, щоб представити сценарій як факт. В одному з помітних випадків відеокліп нібито показує, як українські солдати стріляють у жінку та дитину. Дослідники OSINT визначили, що фільм насправді було знято на окупованій росією території з акторами в українській формі [5; 14; 16].

З багатьох прикладів використання дипфейків в Інтернеті під час російсько-української війни, приклад заяви Зеленського про закінчення війни був, мабуть, найлакучішим. Це пояснюється тим, що він підкреслив, як дипфейки можуть бути використані разом зі зламаними медіасервісами для поширення контрфактичних повідомлень. Кінцевим результатом цього інциденту стало поширення неправдивої інформації з надійного джерела.

У перший рік російського вторгнення синтетичний аудіоконтент був відносно рідкістю в російських інформаційних операціях, спрямованих проти України, а зловмисники використовували інші медіаформати для цілей явної дезінформації. Ситуація почала змінюватися протягом другого року війни, включаючи деякі помітні приклади.

В одному з випадків було використано підібраний аудіозапис виступу губернатора Техасу Грега Ебботта, взятого з інтерв'ю Fox News про імміграційну політику США. Російські джерела опублікували змінену версію інтерв'ю в Telegram, у якій використовується фальшивий закардовий голос губернатора, коли відеозапис обрізається, щоб показати додаткові кадри. Фальшиве аудіо заявляло, що президент США Джо Байден повинен навчитися у президента росії Володимира Путіна, «як працювати в національних інтересах». У початковому інтерв'ю губернатор не згадував ні Путіна, ні Байдена, оскільки це було підтверджено речниками Fox News та адміністрацією губернатора [1; 17].

Ще однією важливою аудіотактикою є ліпсинкінг, під час якого справжнє відео покращується, щоб змінити вираз обличчя та вимову людини. Один із таких випадків стався після теракту в Крокус-Сіті-Холі в Москві 22 березня 2024 року. Проросійські джерела, намагаючись знайти докази «причетності України» до нападу, представили численні неправдиві історії та докази. Одна з них містила нібито інтерв'ю з Олексієм Даніловим, колишнім секретарем Ради національної оборони та безпеки України. Російський канал НТВ, що належить державній енергетичній компанії «Газпром», трансливав нібито інтерв'ю з Даніловим. Однак Данілов не з'являвся на українському телебаченні під час телемарафону. Саме відео було поєднанням двох різних інтерв'ю з Даніловим, накладених на них фальшивим голосом; справді, Шаян Сардарізде з BBC Verify підтвердив, що його голос був згенерований штучним інтелектом.

В іншому випадку з'явилося відео, на якому заарештована особа зізнається, що українська розвідка наказала їй убити ультраправого американського медіа-діяча Такера Карлсона під час візиту до Москви в лютому 2024 року. Пізніше повідомлялося про те, що поліграф визначив, що особа, почута на відео, мала ознаки «цифрової маніпуляції».

В інших випадках синтетичний звук також використовується для озвучування відео, щоб приховати акцент чи національність відеопродюсера. Примітно, що ця тактика використовувалася в тому, що пізніше було визнано найбільшою операцією впливу, виявленою в TikTok, в якій

DFRLab та BBC Verify спільно досліджували тисячі пропагандистських відеороликів сімома мовами, що звинувачують українське керівництво в корупції [3; 16].

Текст, згенерований штучним інтелектом, залишається проблемою для спільноти розвідки з відкритим кодом та боротьби з дезінформацією, оскільки загальнодоступні інструменти не можуть надійно виявляти його, особливо в текстах невеликого обсягу, таких як коментарі в соціальних мережах. Однак у деяких випадках текст, згенерований штучним інтелектом, можна ідентифікувати іншими способами.

Одним із найочевидніших випадків використання є переклад. Хоча письмова російська дезінформація в Україні регулярно висміювалася на початку вторгнення через її граматичні помилки, якість тексту пізніше значно покращилася. Наприклад, негативні історії про Україну за допомогою недосконалих, але правдоподібних перекладів, які виглядали краще, ніж попередні спроби. За даними Recorded Future, генеративний ШІ було використано для створення статей для веб-сайтів, пов'язаних з Doppelganger, англійською мовою, таких як Electionwatch.info. Однак повний обсяг використання тексту, згенерованого штучним інтелектом, все ще не визначений, оскільки потрібні нові інструменти для вирішення цієї проблеми та кращого виявлення тексту, згенерованого штучним інтелектом [4; 18].

В епоху, коли інформація поширюється за лічені хвилини, а не за дні, здатність швидко виявляти дипфейки та реагувати на них є надзвичайно важливою. Проте у цьому зв'язку рекомендується застосування наступних підходів.

1. Технічне виявлення – величезна різноманітність способів маніпулювання медіа робить малоімовірним появу автоматизованого виявлення – у сенсі універсального рішення – у найближчому майбутньому. Більше того, економічний стимул створювати все кращі дипфейки наразі набагато вищий, ніж стимул працювати над методами їх виявлення. Держава повинна протидіяти цьому, цілеспрямовано сприяючи експертизі в галузі медіа. Спектр методів виявлення є величезним, починаючи від індивідуального запису міміки та ритмів мовлення високопоставлених лідерів, наприклад, до збору коливальних електромереж для перевірки часу та місця запису або ідентифікації використаного обладнання. Щоб потенційному зловмиснику було важко ефективно використовувати переконливий дипфейк, вкрай важливо, щоб методи виявлення були різноманітними, а в деяких випадках трималися в секреті. В іншому випадку дискримінація у GAI буде постійно адаптуватися до відомих методів виявлення, щоб обійти їх.

2. Стратегія реагування – ефективна стратегія реагування охоплює: загальну обізнаність з проблемою, постійний рівень залучення до вирішення проблеми дипфейків у поєднанні з моніторингом ЗМІ та здатністю швидко виявляти та оцінювати потенційні фальсифікації з технічної точки зору. Крім того, необхідні перевірені процедури: всередині уряду, між відомствами, а також з партнерами в ЄС та НАТО [5; 16].

Висновки. У цілому, технологія штучного інтелекту знизила бар'єри входу для зловмисників, що дає потенціал для швидшого, дешевшого та складнішого виробництва фейкового контенту. Однак наразі технологія штучного інтелекту не може перевершити більш традиційні тактики та методи, такі як вилучення існуючих медіаматеріалів з контексту або їх редагування таким чином, щоб ввести цільову аудиторію в оману.

Нові випадки генеративного ШІ в російських інформаційних операціях, ймовірно, стануть більш поширеними, оскільки з'являться додаткові генеративні моделі. Поки що вони залишаються недосконалими, але є потенційно потужним інструментом у ширшому інформаційному арсеналі росії.

ЛІТЕРАТУРА

1. Бугера О. І. Використання штучного інтелекту для запобігання злочинності. Вчені записки ТНУ імені В. І. Вернадського. Серія: юридичні науки. 2021. Т. 32 (71). № 6. С. 82–86.
2. Дуфенюк О. Розслідування воєнних злочинів: логістичні, криміналістичні та судово-медичні питання. *Юридичний науковий електронний журнал*. 2022. № 4. С. 369–374. DOI: 10.32782/2524-0374/2022-4/88 (дата звернення: 17.11.2022)
3. Латиш К. Цифрова криміналістика у період війни в Україні: можливості використання спеціальних знань у сфері інформаційних технологій. *Kriminalistika ir teismo ekspertologija : mokslas, studijos, praktika*. 2022. Т. 18. С. 32. URL: https://repository.mruni.eu/bitstream/handle/007/18524/XVIII%20Vilnius%202022_compressed-18-21.pdf?sequence=1&isAllowed=y
4. Мамедов Г. Цифрова криміналістика. Як це допомогло зібрати докази злочинів у Бучі? / *New Voice*. 08.06.2022. URL: <https://nv.ua/ukr/opinion/viyna-v-ukrajini-yak-cifrovakriminalistika-vikrivaye-zlochini-rf-v-ukrajini-novini-ukrajini-50248411.html>
5. Матуєлене С., Шевчук В., Балтрунене Ю. Штучний інтелект в діяльності органів правопорядку та юстиції: вітчизняний та європейський досвід. *Теорія та практика судової експертизи і криміналістики*. 2022. Вип. 4 (29). С. 12–46.
6. Про затвердження плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2021-2024 роки : Розпорядження Кабінету Міністрів України від 12 травня 2021 р. № 438-р. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/438-2021-%D1%80#Text>
7. Про затвердження Правил забезпечення захисту інформації в інформаційних, телекомунікаційних та інформаційно-телекомунікаційних системах: Постанова Кабінету міністрів України від 29.03.2006 № 373. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/373-2006-%D0%BF#Text>
8. Про інформацію: Закон України від 02.10.1992 № 2657-XII. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/2657-12#Text>
9. Про основні засади забезпечення кібербезпеки України: Закон України від 05.10.2017 № 2163-VIII. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text>
10. Про рішення Ради національної безпеки і оборони України від 15 жовтня 2021 року «Про Стратегію інформаційної безпеки»: Указ Президента України «від 28 грудня 2021 року № 685/2021. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/n0080525-21#Text>
11. Про схвалення Концепції розвитку штучного інтелекту в Україні : Розпорядження Кабінету Міністрів України від 2 грудня 2020 р. № 1556-р. *Верховна Рада України. Законодавство України*. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text>
12. Стратегія розвитку штучного інтелекту в Україні: монографія / А.І. Шевченко, С.В.Барановський, О.В.Білокобильський, Є.В.Бодяньський та ін. [За заг. ред. А.І.Шевченка]. Київ: ІПШІ, 2023. 305 с.
13. Токарева К. С., Савліва Н. О. Особливості правового регулювання штучного інтелекту в Україні. *Юридичний вісник*. 2021. Вип. 3 (60). С. 151.
14. Шевчук В.М., Авдеева Г.К. Використання технологій штучного інтелекту та спеціальних знань у розслідуванні воєнних злочинів. *Правнича наука та законодавство України: європейський вектор розвитку в умовах воєнного стану: монографія*; Нац. акад. прав. наук України. Харків: Право, 2023. С. 491–503.
15. Юртаєва К. В. Використання технологій штучного інтелекту в реалізації стратегій «predictive policing»: можливості, проблеми та перспективи для України. Використання технологій штучного інтелекту у протидії злочинності: наук.-практ. семінар (Харків, 05.11.2020). Харків, 2020. С. 99–104.
16. Hood, L. (2023), "Deepfakes in warfare: new concerns emerge from their use around the Russian invasion of Ukraine", October 26. URL: <https://theconversation.com/deepfakes-in-warfare-new-concerns-emerge-from-their-use-around-the-russian-invasion-of-ukraine-216393>
17. Kleemann, A. (2023), "Deepfakes – When We Can No Longer Believe Our Eyes and Ears", SWP Comment, vol. 52. URL: <https://www.swp-berlin.org/publikation/deepfakes-when-we-can-no-longer-believe-our-eyes-and-ears>
18. Osadchuk, R. (2024), "AI tools usage for disinformation in the war in Ukraine," Digital Forensic Research Lab (DFRLab), July 9. URL: <https://dfriab.org/2024/07/09/ai-tools-usage-for-disinformation-in-the-war-in-ukraine>

Дата першого надходження рукопису до видання: 24.10.2025
 Дата прийнятого до друку рукопису після рецензування: 14.11.2025
 Дата публікації: 28.11.2025