

НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ «ОДЕСЬКА ЮРИДИЧНА АКАДЕМІЯ»

ФАКУЛЬТЕТ КІБЕРБЕЗПЕКИ ТА ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ

28 листопада 2025 року



КІБЕРБЕЗПЕКА В СУЧАСНОМУ СВІТІ: АКТУАЛЬНІ ВИКЛИКИ

**МАТЕРІАЛИ
VI МІЖНАРОДНОЇ НАУКОВО-ПРАКТИЧНОЇ
КОНФЕРЕНЦІЇ**

УДК 004.056

Відповідальний редактор:

О. В. Дикий, декан факультету кібербезпеки та інформаційних технологій,
кандидат юридичних наук, доцент

Матеріали видано в авторській редакції.
Повну відповідальність за достовірність та якість поданого матеріалу
несуть учасники конференції, їхні наукові керівники,
які рекомендували ці матеріали до друку.

Рекомендовано до друку
Вченою радою Національного університету
«Одеська юридична академія»
(протокол №3 від 29.11.2025 р.)

Матеріали Міжнародної науково-практичної конференції «Кібербезпека в сучасному світі: актуальні виклики» (м. Одеса, 28 листопада 2025 р.). Одеса, 2025. 287 с.

У конференції взяли участь студенти, аспіранти, молоді вчені, викладачі та науковці. Конференція проводиться на базі Національного університету «Одеська юридична академія».

Редакційна колегія:

ГОРБАЧЕНКО Станіслав, д.е.н., професор

СОКОЛОВ Артем, д.т.н., професор

АХМАМЕТЬЄВА Ганна, к.т.н., доцент

РАЗІНКІН Нікіта, асистент кафедри

СУЧАСНІ МЕТОДИ АНАЛІЗУ ДАНИХ ДЛЯ ВИЯВЛЕННЯ ДЕЗІНФОРМАЦІЇ

Гура Володимир

*доцент кафедри інформаційних технологій Національного університету
«Одеська юридична академія», кандидат технічних наук, доцент*

Гандзій Ілля

*студент 2-го курсу магістратури факультету кібербезпеки та інформаційних
технологій Національного університету «Одеська юридична академія»*

У сучасному суспільстві засоби масової інформації відіграють визначальну роль у формуванні громадської думки та соціальних процесів. Власники багатьох інформаційних каналів, керуючись інтересами бенефіціарів і спонсорів, часто подають новини під певним кутом, який може суттєво відрізнятися від реальних подій [1, 3]. Така практика призводить до спотворення фактів і поширення дезінформації, що швидко набуває вірусного характеру завдяки цифровим платформам [4]. У світі, де доступ до достовірної інформації стає стратегічним ресурсом, той, хто контролює її потік, здатен впливати на рішення людей, політичні процеси та суспільну стабільність [3, 5].

Критичне мислення, навички аналізу джерел і вміння розрізняти правду від маніпуляції стають необхідними для свідомого сприйняття процесів в інформаційному суспільстві [2,4]. Особливо гостро ця проблема проявляється в умовах цілеспрямованого впливу на окремі соціальні групи через медіа, де інтереси власників каналів часто переважають над об'єктивністю висвітлення подій і процесів у суспільстві. Як наслідок, виникає суперечливість подання новин у медійних потоках, що ускладнює сприйняття реальності та сприяє поляризації суспільства.

Активне використання комунікаційного інструменту Telegram, є наочним прикладом різноманітного інформаційного середовища. Відсутність ефективних інструментів автоматичного виявлення дезінформації в україномовному сегменті Telegram створює критичну вразливість державі, особливо в умовах гібридної інформаційної боротьби. Таким чином, розробка адаптованих технологій аналізу контенту, що розміщується в інформаційному середовищі є не лише науковим, а й суспільно-політичним пріоритетом [7].

Дослідження присвячено вивченню способів обробки медійних даних для виявлення дезінформаційного контенту у цифровому просторі. На основі аналізу наукових публікацій створено комбіновану систему, яка об'єднує текстові, візуальні, аудіо- та графові ознаки в єдине рішення. Ця система спеціально налаштована на особливості україномовного контенту в Telegram і доповнена інструментами, що пояснюють, чому саме прийнято те чи інше рішення. Перевірка на реальних даних довела кращі результати порівняно з відомими аналогами за точністю, швидкістю та зрозумілістю роботи. Новизна полягає в тому, що система враховує мовні особливості української аудиторії та робить процес перевірки прозорим для користувачів.

Сучасний інформаційний простір швидко трансформується завдяки цифровим технологіям, які впливають на те, як створюється, поширюється та сприймається контент. Великі обсяги медійних даних стають доступними всім учасникам цифрового простору, і це створює умови для широкого використання недостовірної інформації. Дезінформація – це навмисне поширення хибних відомостей з метою зміни думки чи поведінки людей. Такі дії часто поєднуються з використанням штучно створених зображень, відео чи аудіо, а також систем автоматичних акаунтів [7, 8]. Їх створення і розповсюдження може суттєво впливати на суспільну думку, виборчі процеси та інформаційну безпеку держави. Саме у цьому сенсі в українському суспільстві особливу роль відіграє комунікаційна платформа Telegram, де повідомлення поширюються дуже швидко, а контроль за її вмістом не врегульовано на законодавчому рівні. Без застосування спеціальних інструментів для перевірки це створює серйозні ризики. Тому створення систем, які працюють саме з українським сегментом контенту в Telegram, є важливим завданням, як для суспільства і для держави [1].

Метою роботи є створення мультимодальної системи автоматичного виявлення дезінформації в україномовному сегменті Telegram, яка одночасно аналізує текст, зображення, аудіо/відео та мережу поширення контенту, а також містить вбудовані інструменти пояснюваності рішень. Ці інструменти забезпечують прозорість роботи системи, надаючи користувачу чітке обґрунтування кожного вердикту, які саме слова чи мовні конструкції, деталі зображення, ознаки синтезованого аудіо або аномалії в схемі поширення (наприклад, скоординована активність ботів) стали підставою для класифікації повідомлення як дезінформації, разом із кількісною оцінкою внеску кожної ознаки у фінальне рішення.

Для реалізації цієї мети проводиться огляд сучасних методів виявлення дезінформації (текстових, мультимедійних, графових та гібридних підходів), розробляється модульна архітектура системи з об'єднанням різних типів контенту (текст, зображення, аудіо/відео та мережа поширення), виконується адаптація й тонке налаштування моделей під особливості української мови, інтегруються інструменти пояснення рішень на основі оцінки внеску кожної ознаки, а також здійснюється комплексна перевірка ефективності системи на реальних прикладах україномовного контенту.

Об'єктом дослідження є інформаційні потоки контенту у Telegram, що містять текст, зображення, аудіо, інформацію про взаємодію та спосіб поширення. Предметом дослідження виступають математичні моделі, алгоритми та способи оцінки, які дозволяють автоматично визначати недостовірну інформацію з високою точністю, швидкістю та адекватністю.

Новизна роботи полягає в тому, що система вперше об'єднує аналіз кількох типів даних з урахуванням особливостей української мови та реалізує повноцінну пояснення рішень на основі оцінки внеску кожної ознаки. Це відрізняється від більшості існуючих рішень, які або не адаптовані до української мови, або працюють як «чорні скриньки» без надання зрозумілих обґрунтувань класифікації [2, 3].

Сучасні дослідження надають поступовий перехід від простих методів перевірки тексту до складних систем, що враховують кілька типів контенту одночасно. Спочатку увага зосереджувалася на словах, стилі викладання контенту, та його емоційних відтінків. Але такі способи втрачають ефективність, коли з'являлася мультимедійна складова контенту, а саме зображення чи звук. Сьогодні сучасні системи (зокрема VERA, FakesCope, Reality Defender, Meta AI Multimodal Fact-Check та Microsoft Video Authenticator) здатні одночасно аналізувати текстову, фото-, відео- та аудіоскладові контенту й виявляти невідповідності між ними, а також ознаки підробки медіафайлів.

Окремий напрямок – вивчення поширення інформаційного контенту. Дослідження демонструють, що дезінформація характеризується аномально

швидким вірусним розповсюдженням, концентрацією активності навколо обмеженої кількості ключових акаунтів та наявністю скоординованих мереж ботів і тролів [8, 11].

Це допомагає передбачити, чи стане повідомлення вірусним. Додатково аналізується поведінка акаунтів, щоб виявити ботів і фейкові/скоординовані профілі, які працюють за однаковими сценаріями та шаблонами [8].

Найперспективнішим вважається об'єднання всіх типів даних. Сучасні моделі дозволяють порівнювати текст із зображенням чи аудіо фрагментом і знаходити невідповідності [9]. Проте більшість таких рішень створено для англійської мови, тому для українського мовного сегменту потрібні спеціальні налаштування. Сучасні способи виявлення недостовірного контенту різняться за цілями, методами та інформаційними потоками, де вони розповсюджуються. Кожен спосіб має свої сильні сторони, але й обмеження, що ускладнює повне вирішення проблеми. У таблиці 1 наведено порівняльний аналіз сучасних методів виявлення дезінформації та недостовірного контенту за такими критеріями, тип даних, рівень автоматизації, підтримка української мови, швидкодія та наявність механізмів пояснення.

Таблиця 1 – Порівняння способів перевірки недостовірної інформації

Мета	Опис	Канали поширення	Переваги	Недоліки
Знайти помилки в тексті	Аналізує слова, стиль, емоції	Telegram, X, Facebook	Швидко працює з текстом	Не бачить фото чи звук
Виявити підроблені зображення	Перевіряє фото, обличчя, підписи	Instagram, TikTok, Telegram	Добре ловить deepfake	Потребує текст для порівняння
Зупинити вірусне поширення	Стежить за репостами, ланцюжками	Telegram-канали, групи	Передбачає спалах фейків	Багато обчислень
Розпізнати фальшивий звук	Аналізує голос, шум, синхронізацію з відео	YouTube, TikTok, Telegram-аудіо	Знаходить підроблені записи	Вимагає якісного аудіо
Визначити бото-мережу	Шукає шаблонні дії, частоту постів	Telegram-чати, X, ботоферми	Виявляє штучне поширення	Потребує доступу до метаданих
Перевірити метадані	Аналізує час, місце, пристрій створення	Telegram, WhatsApp, email	Доводить автентичність файлів	Не працює з видаленими даними
Виявити емоційний маніпулятив	Шукає страх, гнів, тривогу в тексті	Telegram-канали, соцмережі	Блокує психологічний тиск	Складно відрізнити від емоцій
Повне виявлення фейків	Об'єднує текст, фото, звук, мережу	Усі платформи (Telegram, X, Meta)	Найточніший результат	Потребує налаштування під мову

Важливим стало питання пояснення прийнятих рішень – тобто чому система класифікувала конкретне повідомлення саме як дезінформацію. Системи аналізу контенту повинні не просто визначити недостовірну інформацію, а довести, чому саме вона класифікована, як недостовірна. Для цього використовують методи, які вказують, які частини контенту найбільше вплинули на визначення їх класифікації.

Для аналізу використано набір реального контенту з платформ Telegram, X та Meta (текст, зображення, аудіо та відео), доповнений спеціально зібраною україномовною базою текстових, візуальних і мультимедійних матеріалів. Увесь контент пройшов ретельну анонімізацію та видалення персональних метаданих для забезпечення конфіденційності джерел.

Спочатку контент звільнюється від метаданих та розбивається на частини. Текстова складова обробляється спеціальною моделлю для української мови, аудіо- та відеофрагменти обробляються моделями розпізнавання та верифікації (Whisper + спеціалізовані детектори deepfake-аудіо), що порівнює картинки з текстом. Мережа поширення аналізується за швидкістю розповсюдження та зв'язками між акаунтами. Усі зроблені класифікації зводяться разом, і система вирішує, чи є контент достовірним.

Після класифікації система чітко пояснює своє рішення та показує користувачу, які саме слова, деталі фото, фрагменти аудіо чи аномалії поширення «переважили» і чому повідомлення визнано дезінформацією. Пояснення формується за допомогою стандартних методів інтерпретації, які кількісно оцінюють внесок кожної ознаки у кінцевий результат [5].

Система побудована за принципом поетапної обробки. Спочатку контент надходить у кілька паралельних блоків, один відповідає за текст, другий за зображення, третій за аудіо складову, четвертий за мережу поширення. Кожен блок видає свій класифікатор, а потім усі результати об'єднуються з урахуванням того, наскільки важлива кожна складова у кожному контенті.

Для ефективного виявлення недостовірного контенту використовується комбінація традиційних і сучасних інструментів [3, 5, 9]. Кожен з них має свою сферу застосування, але разом вони створюють надійну основу для аналізу. Перевірка фактів, аналіз джерел і змісту, моніторинг соцмереж та залучення експертів – це перевірені методи, які доповнюють автоматизовану класифікацію контенту. Їхня адаптація до україномовного сегменту платформи Telegram вимагає урахування локальних особливостей, мовних нюансів, культурного контексту та лексичної специфіки. Порівняння цих інструментів наведено в таблиці 2.

Таблиця 2 – Інструменти виявлення недостовірних новин

Інструмент	Опис	Як працює на практиці	Переваги	Недоліки	Адаптація до української мови
Перевірка фактів	Перевірка правдивості фактів у новинах через аналіз джерел, пошук доказів та порівняння з надійними базами	Журналісти чи аналітики беруть заяву, шукають першоджерело, перевіряють документи, фото, свідчення	Висока точність при наявності доказів	Потребує часу та доступу до джерел	Потрібні відкриті бази перевірених фактів українською
Дослідження джерел	Оцінка надійності видання: репутація, акредитація, історія, авторство	Перевіряється, хто написав, хто володіє каналом, чи є ліцензія, чи є спростування в минулому	Допомагає відсікати сумнівні канали	Не завжди є публічна інформація	Адаптація через реєстр українських ЗМІ та Telegram-каналів
Аналіз змісту	Вивчення тексту, фото, відео на логіку, послідовність, підробки	Шукають невідповідності: старе фото в новій події, нелогічні аргументи, підроблені цитати	Виявляє маніпуляції в деталях	Складно з сарказмом чи гумором	Потребує україномовних моделей для аналізу стилю
Моніторинг соцмереж	Відстеження поширення, трендів, активності акаунтів	Аналіз репостів, реакцій, хто і як швидко ділиться	Передбачає вірусність фейків	Багато шуму, потрібні фільтри	Фокус на Telegram-каналах і групах
Експертні оцінки	Залучення фахівців для аналізу теми	Експерт оцінює за своїм досвідом: лікар — медичні фейки, історик — історичні	Глибоче розуміння теми	Суб'єктивність, час	Потрібні українські експерти з різних галузей

Як видно з таблиці 2, жоден інструмент не є універсальним і вони доповнюють один одного. У запропонованій системі ці методи інтегровано в автоматизовану платформу, перевірка достовірності фактів – через аналіз джерел і змісту, моніторинг – через граф поширення, тобто схему зв'язків між акаунтами, де видно, хто, коли і як швидко поширює повідомлення (репости, коментарі та їх ланцюжки), експертні оцінки – через пояснення, тобто автоматичне визначення, які саме слова, деталі зображення чи зв'язки

найбільше вплинули на рішення системи, ніби «експерт» пояснює свої висновки. Адаптація до української мови досягається за допомогою спеціалізованих моделей, таких як ukrBERT – це українська версія моделі BERT (Bidirectional Encoder Representations from Transformers), яка є одним із найпоширеніших інструментів для обробки природної мови [8]. Модель навчається на великому наборі україномовних текстів, включаючи новини, Вікіпедію, соціальні мережі та літературу – загалом понад 100 мільйонів слів та їх комбінацій. UkrBERT вміє розуміти контекст речень, розпізнавати емоції, стиль письма та маніпулятивні конструкції, такі як перебільшення чи приховані упередження [7, 11, 15]. Наприклад, якщо в повідомленні є фраза з емоційним забарвленням, як "шокуюча правда про вакцини", модель може виявити, що це не відповідає існуючим фактам, базуючись на набутих шаблонах україномовного підробленого медіаконтенту. Адаптація до української мови робить її ефективною для платформи Telegram, де змішується літературна мова та її діалекти, підвищуючи точність на 5–10% порівняно з загальними моделями [1, 2]. Доступна на платформі Hugging Face (відкритий репозиторій моделей машинного навчання) для подальшої адаптації та інтеграції під конкретні завдання, як дезінформація, та UkrCLIP – це українська адаптація моделі CLIP (Contrastive Language-Image Pretraining), розроблена для зв'язку тексту та зображень. Вона навчається на парах "текст-зображення" з україномовних джерел, таких як новина з фото чи пости в Instagram, – понад 100 тисяч прикладів. Модель перевіряє, чи відповідає зображення опису в тексті: наприклад, якщо новина про "протести в місті" містить фото з міста іншої країни, UkrCLIP виявить невідповідність [14]. Вона також розпізнає підробки зображень, такі як підроблений медіаконтент, за абстрактами пікселів і стилем. Адаптація до української мови враховує локальні візуальні шаблони, як українські прапори чи символи, що робить її корисною для перевірки контенту в Telegram. Точність зростає на 7–12 % для україномовних пар текст-зображення порівняно з базовою CLIP, яка орієнтована на англійську [3, 4]. Модель доступна на GitHub і Hugging Face для інтеграції в систему перевірки достовірності фактів.

Ці моделі роблять систему гнучкою для україномовного сегменту, де недостовірні інформація часто опановує певні теми. Навчання на контенті з українських джерел, таких як UA-News чи UA-Wiki, дозволяє уникнути помилок через переклад. Наприклад, ukrBERT може розпізнати маніпуляцію в фразі "секретна правда про подію", а UkrCLIP – перевірити супровідне фото. Загалом, така адаптація підвищує загальну ефективність системи на 10–15%, роблячи її придатною для реального використання в Telegram [5, 7]. Це дозволяє системі працювати швидше та точніше, ніж ручні методи. Такий підхід робить процес перевірки контенту прозорим і доступним для широкого кола користувачів. UkrCLIP, які навчені на україномовному контенті. Ці моделі

дозволяють системі точно аналізувати текст і зображення, враховуючи особливості української мови, такі як морфологія, синтаксис і культурний контекст. Без такої адаптації моделі, створені для англійської чи інших мов, часто помиляються через відмінності у словнику, граматиці чи ідіомах, що знижує точність виявлення недостовірного контенту [10, 11, 16].

Запропонована система для виявлення дезінформації в україномовному сегменті Telegram побудована за модульним принципом. Вона складається з кількох блоків, які працюють одночасно, що дозволяє швидко обробляти великий потік даних і давати точний результат.

На вході система приймає контент з Telegram, а саме текстові повідомлення, зображення, аудіо-, відео-, а також інформацію про те, хто і коли його опублікував, як його поширюють – через репости, коментарі чи реакції.

Перший блок відповідає за аналіз текстового контенту. Він перевіряє, чи є в повідомленні маніпуляції, перебільшення, емоційне забарвлення, відсутність джерел або логічні невідповідності. Для цього використовується спеціальна модель, навчена на україномовних текстах.

Другий блок аналізує зображення. Він порівнює, чи відповідає зображення опису в тексті. Наприклад, якщо повідомлення говорить про одну подію, а фото – з іншої, система це помітить. Також перевіряються ознаки підробки зображень.

Третій блок аналізує аудіо. Він розпізнає мову, визначає, чи є голос справжнім, чи синтезованим, і перевіряє, чи збігається звук із рухами губ у відео.

Четвертий блок стежить за поширенням. Він дивиться, як швидко і ким поширюється повідомлення, чи є шаблонні дії, що вказують на ботоферми або скоординовану кампанію.

Усі висновки з чотирьох блоків збираються в одному місці. Система автоматично визначає, наскільки важлива кожна частина даних у конкретному випадку – наприклад, якщо немає аудіо, його вплив не враховується.

Далі працює блок пояснення, він показує, які саме слова, деталі зображення чи зв'язки між акаунтами вплинули на рішення. Це робить процес прозорим – аналітик бачить, чому система вважає повідомлення підозрілим.

На останньому етапі система приймає рішення, це правда чи фейк. Результат виводиться на зручний інтерфейс перегляду – з графіками, прикладами, поясненнями та можливістю експорту даних.

Система працює швидко, адаптована до української мови і може бути встановлена навіть на простих пристроях. Вона перевершує окремі текстові чи візуальні моделі, бо враховує всі типи даних і пояснює свої дії.

Робота системи перевіряється за кількома показниками, наскільки точно вона знаходить недостовірну інформацію, як швидко реагує та наскільки зрозуміло пояснює свої дії. Усе це об'єднується в одну оцінку.

Система була протестована на наборі з реальних повідомлень з платформи Telegram. Дані включали текст, зображення, аудіо, відео та метадані про поширення – усе очищене від особистих відомостей. Тестування проводилося в режимі імітації живої роботи, повідомлення надходили по одному, як у реальному часі, з випадковими інтервалами від 1 до 30 секунд. Систему було реалізовано на звичайному сервері без спеціалізованого обладнання. Кожне повідомлення оброблялося повністю, від аналізу тексту до присвоєння класифікації та пояснення.

Порівняння проводилося з трьома відомими рішеннями – ukrBERT (тільки текст), UkrCLIP (текст + зображення) та Whisper + GraphSAGE (аудіо + поширення). Усі моделі тестувалися за однакових умов, той самий набір даних, той самий сервер, той самий порядок повідомлень. Жодна модель не мала доступу до додаткових ресурсів чи попередньої інформації.

Результати показали, що розроблена система перевершує аналоги за всіма трьома критеріями, точністю, швидкістю та зрозумілістю. Вона досягла балансу між точністю та повнотою моделі класифікації 0,94, тоді як ukrBERT – 0,89, UkrCLIP – 0,91. Час реакції склав менш 3 секунд на повідомлення, що лише на 0,5 секунди повільніше за найшвидший окремий блок, але з повним поясненням.

Тестування показало, що об'єднання в одній системі всіх типів даних (текст + зображення + аудіо + граф поширення) дає суттєво кращий результат, ніж окремі моделі. Перевірка в реальних умовах (імітація живого потоку з Telegram) підтвердила готовність системи до розгортання.

В результаті створено комбіновану систему для перевірки недостовірного контенту в Telegram з урахуванням особливостей української мови. Вона об'єднує різні типи контенту, швидко реагує та пояснює свої дії. Перевірка на реальних прикладах довела її перевагу.

Список використаних джерел

1. Alenezi F. From Misinformation to Insight: A Systematic Review of AI-Driven Solutions for Combating Fake News. *Information*, 2025, vol. 16, no. 3, p. 189. DOI: 10.3390/info16030189. PDF: <https://www.mdpi.com/2078-2489/16/3/189/pdf>.
2. Almarashy H., et al. Artificial Intelligence in the Battle Against Disinformation and Misinformation: A Systematic Review of Challenges and Approaches. *Journal of Artificial Intelligence and Capsule Networks*, 2025, vol. 7, no. 1, pp. 69–90. DOI: 10.36548/jain.2025.1.006. PDF: <https://www.researchgate.net/publication/388421309>.
3. Buchanan, J., et al. Fake News, Disinformation and Misinformation in Social Media: A Review. *Social Network Analysis and Mining*, 2023, vol. 13, no. 1, p. 69. DOI:

10.1007/s13278-023-01028-1.

PDF:

<https://link.springer.com/content/pdf/10.1007/s13278-023-01028-1.pdf>.

4. Choras M., et al. Machine Learning and Deep Learning Applications in Disinformation Detection: A Bibliometric Assessment. *Electronics*, 2024, vol. 13, no. 22, p. 4352. DOI: 10.3390/electronics13224352. PDF: <https://www.mdpi.com/2079-9292/13/22/4352/pdf>.
5. De Fraga T., et al. AI in Disinformation Detection: A Review of Techniques and Challenges. *ACIG Journal of Computer Science*, 2024, vol. 2, no. 1, pp. 1–20. DOI: 10.47861/acigs.2024.2.1.1. PDF: <https://www.acigjournal.com/AI-in-Disinformation-Detection,200200,0,2.html>.
6. X. Men, V. Y. Mariano, "Explainable Fake News Detection Based on BERT and SHAP Applied to COVID-19", *International Journal of Modern Education and Computer Science*, Vol.16, No.1, pp. 11-22, 2024.
7. Choras, M., et al. Machine Learning and Deep Learning Applications in Disinformation Detection. *Electronics*, 2024, vol. 13, no. 22, p. 4352. URL: <https://www.mdpi.com/2079-9292/13/22/4352>.
9. Freeze, M., Baumgartner, M., et al. Fake news, fast and slow: Deliberation reduces belief in false (but not true) news headlines. *Journal of Experimental Psychology: General*, 2020, vol. 149, no. 8, pp. 1608–1613. URL: <https://doi.org/10.1037/xge0000729>.
10. Schnellenbach, J. The economics of fake news. *CESifo Economic Studies*, 2020, vol. 66, no. 3, pp. 207–227. URL: <https://doi.org/10.1093/cesifo/ifaa011>.
- We Are Social & Hootsuite. *Digital 2025: Global Report*. 2025. URL: <https://datareportal.com/reports/digital-2025-global-overview-report>.
11. Alenezi, F. From Misinformation to Insight: A Systematic Review of AI-Driven Solutions for Combating Fake News. *Information*, 2025, vol. 16, no. 3, p. 189. URL: <https://www.mdpi.com/2078-2489/16/3/189>.
12. Kaliyar R., Goswami A., Narang P. (2021). DeepFakE: improving fake news detection using tensor decomposition-based deep neural network. *The Journal of Supercomputing*. V. 77. P. 1015-1037. <https://doi.org/10.1007/s11227-020-03294-y>.
13. Martel C., Pennycook G., Rand D. (2020). Reliance on emotion promotes belief in fake news. *Cognitive Research: Principles and Implications*. V. 5 (1). P. 47. <https://doi.org/10.1186/s41235-020-00252-3>.
14. Palani B., Elango S., Viswanathan V. (2022). A multimodal deep learning framework for automatic fake news detection using capsule neural network and BERT. *Multimedia Tools and Applications*. V. 81. P. 5587-5620.
15. Pehlivanoglu D., Lin T. et al. (2021). The role of analytical reasoning and source credibility on the evaluation of real and fake full-length news articles. *Cognitive Research: Principles and Implications*. V. 6: 24. <https://doi.org/10.1186/s41235-021-00292-3>.
16. Verstraete M., Bambauer D.E., and Bambauer J.R. (2017). Identifying and Countering Fake News. *Arizona Legal Studies Discussion Paper*. V. 73 (3). Available at: https://repository.uchastings.edu/hastings_law_journal/vol73/iss3/6.

КІБЕРБЕЗПЕКА В СУЧАСНОМУ СВІТІ: АКТУАЛЬНІ ВИКЛИКИ ТА ПРАКТИКА

**МАТЕРІАЛИ
VI МІЖНАРОДНОЇ
НАУКОВО-ПРАКТИЧНОЇ
КОНФЕРЕНЦІЇ**

*28 листопада 2025 року
м. Одеса*

Відповідальний редактор: О.В. Дикий

Дизайн та верстка:
М.Ю. Овчинников