

МОЖЛИВОСТІ МАНІПУЛЯЦІЇ ТА МЕТОДИ ВИЯВЛЕННЯ ФАЛЬШИВИХ ЗОБРАЖЕНЬ І ВІДЕО

UDC 004.932

R. R. Mytulinskyi

POSSIBILITIES OF MANIPULATION AND METHODS OF DETECTION FAKE IMAGES AND VIDEOS

Анотація. У доповіді висвітлюються можливості маніпуляцій зображеннями та відео за допомогою сучасних алгоритмів, таких як DeepFake. Описано ключові принципи створення фальшивок, їхній вплив на суспільство, а також актуальні методи виявлення фальшивих матеріалів. Проаналізовано приклади використання таких технологій у різних сферах – від розважальної індустрії до політичних маніпуляцій. Окрема увага приділяється сучасним методам боротьби з фальшивими зображеннями, які включають аналіз фізичних закономірностей, використання метаданих та алгоритмів машинного навчання.

Ключові слова: DeepFake, маніпуляція зображеннями, фальшиві відео, цифрова безпека, машинне навчання, штучний інтелект, виявлення фальшивок.

Вступ. DeepFake – технологія, що використовує генеративно-змагальні нейронні мережі (GAN) для створення реалістичних фальшивих відео та зображень. Її поширення впливає на політику, бізнес і приватність, підкреслюючи важливість методів виявлення маніпуляцій.

Матеріали та методи. Дослідження базувалося на аналізі наукових статей, звітів компаній і матеріалів міжнародних конференцій. Розглянуто три основні підходи:

- Аналіз фізичних закономірностей. Наприклад, розбіжності в освітленні, неправильне відображення тіней або аномальні рухи очей і губ.
- Аналіз метаданих. Вивчення інформації, що міститься у файлах зображень і відео, наприклад, змінені параметри зйомки.
- Методи машинного навчання. Використання алгоритмів, які аналізують тисячі зображень і відео для виявлення специфічних артефактів, властивих фальшивим матеріалам.

Також вивчено автоматизовані рішення для виявлення фальшивок у реальному часі, які можуть бути інтегровані у платформи соціальних мереж.

Результати. Технології створення: GAN дозволяють генерувати високоточні фальшивки. Сфери застосування: Розваги, освіта, кінематограф; маніпуляції, дискредитація.

Методи виявлення: Аналіз моргання, рухів губ, піксельних аномалій; моделі машинного навчання із точністю понад 90%.

Висновки. DeepFake є серйозною загрозою для інформаційного середовища. Для протидії необхідно розвивати адаптивні моделі виявлення, посилювати співпрацю між експертами та підвищувати цифрову грамотність.

Література

1. Deepfake detection using deep learning methods: A systematic and comprehensive review <https://wires.onlinelibrary.wiley.com/doi/10.1002/widm.1520> (дата звернення 02.12.2024).
2. Діпфейки: як розпізнати та захиститися? <https://netfreedom.org.ua/article/dipfejki-yak-rozpiznati-ta-zahistitisya> (дата звернення 02.12.2024).
3. Synthetic Media & Deepfakes <https://innovating.news/article/synthetic-media-deepfakes/> (дата звернення 02.12.2024).