

1. Alimoradi, R. (2016). A Study of Hyperelliptic Curves in Cryptography. IJCNIS, 8(8), 67–72. <https://doi.org/10.5815/ijcnis.2016.08.08>
2. Yevseiev, S., Kots, H., Liekariiev, Y. (2016). Developing of multi-factor authentication method based on niederreiter-mceliece modified crypto-code system. Eastern-European Journal of Enterprise Technologies, 6(4(84)), 11–23. <https://doi.org/10.15587/1729-4061.2016.86175>.
3. Hubrechts, H. (2011). MEMORY EFFICIENT HYPERELLIPTIC CURVE POINT COUNTING. Int. J. Number Theory, 07(01), 203–214. <https://doi.org/10.1142/S1793042111004034>.
4. Conceição, R. (2020). ON INTEGRAL POINTS ON ISOTRIVIAL ELLIPTIC CURVES OVER FUNCTION FIELDS. Bull. Aust. Math. Soc., 102(2), 177–185. <https://doi.org/10.1017/S0004972720000155>

Дослідження методів та засобів ідентифікації дезінформативних новин у соціальних мережах

УДК 004.8:004.7

Тарас Труш¹, Марія Стадник²

*Тернопільський національний технічний університет імені Івана Пулюя,
¹larastrush.dev@gmail.com, ²stadnyk_m@tntu.edu.ua*

У період стрімкої цифровізації суспільства соціальні мережі стали ключовим джерелом поширення інформації, але водночас є інструментом для навмисного поширення дезінформації. Досвід України, яка з 2014 року опинилася серед перших держав, що зіткнулися з масштабними дезінформаційними кампаніями в соціальних мережах, демонструє, що ця проблема переросла у фактор загрози національній безпеці. Фейкові профілі, ботоферми, маніпулятивні матеріали та синтетичний медіаконтент активно використовуються як інструменти гібридної війни. Додатковим викликом є швидкий розвиток генеративного штучного інтелекту, який значно ускладнює відокремлення правдивої інформації від неправдивої. У зв'язку з цим зростає потреба у створенні дієвих автоматизованих систем виявлення дезінформації.

Метою дослідження є аналіз сучасних методів автоматизованої детекції дезінформації в соціальних мережах та оцінка ефективності гібридних і мультимодальних підходів для виявлення маніпулятивного контенту. Поширення дезінформації в соціальних мережах та розвиток генеративного штучного інтелекту створюють серйозні загрози інформаційній безпеці. Сучасні методи виявлення фейкової інформації потребують поєднання текстового, візуального та мережевого аналізу, однак україномовний інформаційний простір залишається недостатньо дослідженим.

На основі аналізу публікацій з баз даних Scopus, Web of Science, Google Scholar та arXiv за 2017–2025 роки встановлено, що методи детекції дезінформації поділяються на три основні категорії. Контент-орієнтовані підходи базуються на аналізі текстових та візуальних ознак новин із використанням NLP-технік і трансформерних моделей (BERT, RoBERTa, DeBERTa) [1]. Мультимодальні методи поєднують обробку тексту та зображень

через механізми early fusion, late fusion і cross-modal attention, що дозволяє виявляти deepfake-контент і складні семантичні невідповідності між текстом і зображенням [2]. Користувач-орієнтовані підходи базуються на поведінкових характеристиках акаунтів: зокрема, ECS (Ensemble of Specialized Classifiers) застосовує правило максимуму для виявлення нових типів ботів, невідомих під час навчання [3]. Мережеві підходи аналізують структуру поширення інформації в соціальних мережах через графові нейронні мережі (GCN, GAT, BiGCN), моделюючи часові та структурні закономірності поширення [4]. Гібридні підходи, що поєднують контентний і мережевий аналіз (dEFEND, BERT+LightGBM), демонструють стабільно вищу ефективність порівняно з унімодальними системами [5,6].

Окремо досліджено україномовний контекст: датасет EUvsDisinfo фіксує значне зростання проросійської дезінформації напередодні повномасштабного вторгнення у 2022 році [7], а Shared Task UNLP 2025 представив перший публічний бенчмарк для виявлення маніпуляцій у Telegram-дописах українською мовою (9 557 записів, 10 технік маніпуляції) [8].

За результатами дослідження встановлено, що гібридні підходи стабільно перевершують унімодальні: dEFEND досягає 90.4% точності на PolitiFact [5], BERT+LightGBM - 82.9% на LIAR [6]. Водночас жодна з розглянутих систем не є універсальною через проблему зміщення домену: моделі з точністю ~99% на WELFake демонструють падіння на 30% і більше при тестуванні на нових доменах. Україномовний інформаційний простір залишається критично недослідженим - існуючі бенчмарки (EUvsDisinfo, UNLP 2025) є лише першими кроками. Перспективами подальших досліджень є розробка крос-лінгвальних моделей для україномовного середовища, мультимодальних систем детекції deepfake-контенту, а також методів виявлення дезінформації, згенерованої великими мовними моделями.

1. Sayyadiharikandeh M., Varol O., Yang K.-C. et al. Detection of novel social bots by ensembles of specialized classifiers // Proceedings of the 29th ACM International Conference on Information and Knowledge Management. New York : Association for Computing Machinery, 2020. P. 2725–2732.
2. Lv J., Gao Y., Li L., Shi L., Li S. Multi-modal fake news detection: a comprehensive survey on deep learning technology, advances, and challenges // Journal of King Saud University – Computer and Information Sciences. 2025. Vol. 37. P. 306.
3. Sayyadiharikandeh M., Varol O., Yang K.-C. et al. Detection of novel social bots by ensembles of specialized classifiers // Proceedings of the 29th ACM International Conference on Information and Knowledge Management. New York : Association for Computing Machinery, 2020. P. 2725–2732.
4. Sayyadiharikandeh M., Varol O., Yang K.-C. et al. Detection of novel social bots by ensembles of specialized classifiers // Proceedings of the 29th ACM International Conference on Information and Knowledge Management. New York : Association for Computing Machinery, 2020. P. 2725–2732.
5. Sayyadiharikandeh M., Varol O., Yang K.-C. et al. Detection of novel social bots by ensembles of specialized classifiers // Proceedings of the 29th ACM

- International Conference on Information and Knowledge Management. New York : Association for Computing Machinery, 2020. P. 2725–2732.
6. Essa E., Omar K., Alqahtani A. Fake news detection based on a hybrid BERT and LightGBM models // Complex & Intelligent Systems. 2023. Vol. 9. P. 6581–6592.
 7. Leite, J. A., Razuvaevskaya, O., Bontcheva, K., Scarton, C. EUvsDisinfo: A dataset for multilingual detection of pro-Kremlin disinformation in news articles. Proceedings of the 33rd ACM International Conference on Information and Knowledge Management. New York: Association for Computing Machinery, 2024, pp. 5380–5384.
 8. Kyslyi, R., Romanyshyn, N., Sydorskyi, V. The UNLP 2025 Shared Task on Detecting Social Media Manipulation. Proceedings of the Fourth Ukrainian Natural Language Processing Workshop, 2025, pp. 105–111. URL: <https://aclanthology.org/2025.unlp-1.12.pdf>.

Середовище для аналізу атак на SDN-орієнтовані системи

УДК 004.056:004.7

Юрій Кльоц¹, Сергій Мостовий²

*Хмельницький національний університет,
¹klots@khmnu.edu.ua, ²serhii.mostovyi@khmnu.edu.ua*

Сучасні мережеві інфраструктури дедалі частіше використовують централізовані засоби керування, моніторингу та конфігурування обладнання, що наближає їх до принципів програмно-конфігурованих мереж. Одним із прикладів такої системи є Omada Controller, який забезпечує централізоване керування маршрутизаторами, комутаторами, точками доступу та бездротовими клієнтами. Концентрація керуючих функцій в одному компоненті підвищує зручність адміністрування, однак одночасно формує критичну точку впливу для потенційних атак.

Атаки на SDN-системи можуть бути спрямовані на порушення доступності контролера, зміну або блокування службової взаємодії між контролером і мережевими пристроями, імітацію легітимної активності, перевантаження каналів керування або виявлення вразливих сервісів. Особливу складність становить те, що частина атакувального трафіку за окремими ознаками може бути подібною до нормальних службових процесів, зокрема реєстрації пристроїв, оновлення станів, обміну телеметрією або підключення клієнтів. Тому дослідження таких атак потребує не лише фіксації факту аномальної активності, а й аналізу варіативності нормального трафіку в умовах реальної або наближеної до реальної мережевої інфраструктури [2].

Використання готових наборів пакетів для розроблення методів виявлення та протидії атакам не завжди забезпечує достатню достовірність результатів. Такі набори часто не враховують особливостей конкретного контролера, моделі мережевого обладнання, структури службового обміну, кількості клієнтів, типів підключення та характеру адміністративних дій. Крім того, у відкритих датасетах зазвичай складно встановити точний контекст формування трафіку, умови проведення атаки та відповідність отриманих пакетів реальним