

складовою захисту демократичної цілісності та безпеки всієї європейської спільноти.

1. European Council Conclusions on external relations (19 March 2015). <https://www.consilium.europa.eu/en/press/press-releases/2015/03/19/conclusions-russia-ukraine-european-council-march-2015/>
2. FIMI and disinformation as global threats. EUvsDisinfo, 30.01.2026. <https://euvsdisinfo.eu/fimi-and-disinformation-as-global-threats/>
3. 4th EEAS Report on Foreign Information Manipulation and Interference Threats : [Annual report] / European External Action Service (EEAS). Brussels, 2026. 12 March. URL: https://www.eeas.europa.eu/sites/default/files/2026/documents/EEAS%204th%20Threat%20Report_web%20version_1.pdf
4. European Commission. Communication from the Commission on ProtectEU: a European Internal Security Strategy. COM(2025) 148 final. April 2025.
5. European Commission. Joint Communication on the European Democracy Shield. November 2025.
6. Annual Report 2025 / The Soufan Center. New York, 2026. 24 p. URL: thesoufancenter.org (дата звернення: 10.05.2026).
7. EUvsDisinfo. (2020, April 22). “To Challenge Russia’s Ongoing Disinformation Campaigns”: The Story of EUvsDisinfo. <https://euvsdisinfo.eu/to-challenge-russias-ongoing-disinformation-campaigns-the-story-of-euvsdisinfo/>

Методологія збагачення подій SIEM результатами аналізу мережевого трафіку засобами машинного навчання

УДК 004.056.5

Юрій Коровайченко¹, Євгеній Педченко²,
Сергій Гахов³

*Державний університет інформаційно-комунікаційних технологій,
¹y.korovaichenko@duikt.edu.ua, ²e.pedchenko@duikt.edu.ua,
³s.gakhov@duikt.edu.ua*

Операційні центри безпеки в реальній корпоративній інфраструктурі стикаються з протиріччям: платформи моніторингу подій безпеки, такі як SIEM (IBM QRadar, Wazuh, Rapid7 InsightIDR, Palo Alto XSIAM) - консолідують сотні тисяч, або навіть мільйони подій на добу, але реальне покриття мережевого рівня при цьому залишається неповним. Більшість кореляційних правил спирається на стандартні джерела подій, а саме: журнали кінцевих точок та застосунків. Мережевий потік або передається як сирий NetFlow без подальшої обробки, або взагалі залишається поза полем кореляційного аналізу. На практиці це означає, що аналітик першої лінії вимушений працювати з

інтерфейсами мінімум двох рішень: SIEM та NDR і вручну зіставляти те, що система SIEM могла б зробити автоматично.

Метою даної роботи є обґрунтування методології збагачення подій структурованими метаданими машинного навчання в кореляційну логіку SIEM-систем для підвищення ефективності виявлення кіберзагроз в операційних центрах безпеки.

Наявний досвід інтеграції аналізу машинним навчанням мережевого трафіку з платформами моніторингу вже накопичений: відомі підходи на базі Zeek у поєднанні з машинним навчанням в Elastic Security, а також передача алертів від NDR-рішень до SIEM через Syslog або API. Проте в більшості реалізацій зовнішній компонент машинного навчання залишається ізольованим: до платформи надходить готовий алерт без структурованих метаданих класифікації: без оцінки впевненості, без інформації про ознаки, що обумовили рішення про подію. Наявні інтеграції, наприклад, у рішеннях Flowmon і Corelight з тими самими SIEM-платформами працюють за схожою логікою. Це суттєво обмежує можливості побудови складних кореляційних правил, які враховували б не лише факт виявленої аномалії, а й її характеристики.

На відміну від існуючих підходів, запропонована методологія будується на ідеї перенесення контексту машинного навчання всередину події SIEM. Джерелом даних слугують структуровані метадані мережевих сесій, зокрема JSON-логи Corelight на базі Zeek-аналізатора, або записи потоків Flowmon з поведінковими мітками. Міжмережеві екрани, наприклад, Palo Alto Networks NGFW з розширеним логуванням сесій, також можуть бути задіяні як додаткове джерело контексту. З цих даних формується ознаковий вектор: тривалість з'єднання, співвідношення вхідного/вихідного трафіку, кількість пакетів, розподіл їхніх розмірів, значення TTL, ентропія корисного навантаження. Перелік ознак уточнюється за результатами feature importance у процесі навчання.

Класифікацію в запропонованій методології виконує модель Random Forest. Цю модель обрано, насамперед, через стійкість до незбалансованих класів, що типово для трафіку в SOC, та можливість відстежити, які саме ознаки вплинули на конкретне рішення моделі. Швидкодія в умовах потокової обробки також вкладається в допустимі межі. Модель повертає мітку класу та оцінку впевненості, яка може бути безпосередньо використана як динамічний поріг у кореляційних правилах.

Вибір ознак ґрунтується на їхній здатності відображати поведінкові характеристики сесії незалежно від корисного навантаження: тривалість, байтове співвідношення та розподіл розмірів пакетів дозволяють виявляти аномалії на рівні патерну з'єднання, тоді як ентропія навантаження та значення TTL чутливі до тунелювання і спуфінгу відповідно. Остаточний склад вектора визначається за результатами аналізу feature importance на навчальній вибірці.

Етап збагачення події в SIEM реалізується по-різному залежно від платформи: для IBM QRadar - через Custom Properties або DSM-розширення, для Wazuh - через декодери з полями типу data.ml.*, для InsightIDR та XSIAM - через API-інтеграції та custom log parsers. Підсумок: кореляційні правила отримують поля ml_class, ml_score і ml_features_summary безпосередньо в тілі

події. Жодного перемикання між консолями, аналітик бачить висновок системи машинного навчання там, де і всі інші події, а система аналізує збагачені дані.

Попередні лабораторні тести на відкритому датасеті CIC-IDS2017 підтвердили принципову працездатність підходу та прийнятні метрики класифікації для більшості представлених класів атак. Повна валідація методології в умовах реального виробничого середовища запланована на наступному етапі дослідження. Відповідно, кількісні оцінки впливу на середній час виявлення наразі не наводяться.

В роботі представлено методологію збагачення подій структурованими метаданими машинного навчання в кореляційну логіку SIEM-систем, що базується на класифікації мережевого трафіку моделлю Random Forest з поверненням мітки класу та оцінки впевненості. Ознаковий вектор формується з мережевих метаданих сесій, а результати класифікації передаються в тіло події через поля `ml_class`, `ml_score` та `ml_features_summary`. Це усуває необхідність ручного зіставлення даних між консолями та створює передумови для скорочення часу виявлення загроз.

1. Sharafaldin I., Habibi Lashkari A., Ghorbani A.A. Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization. Proceedings of the 4th International Conference on Information Systems Security and Privacy (ICISSP). 2018. P. 108–116. URL: <https://www.scitepress.org/Papers/2018/66398/>
2. Breiman L. Random Forests. Machine Learning. 2001. Vol. 45. P. 5–32. URL: <https://link.springer.com/article/10.1023/A:1010933404324>
3. Palo Alto Networks. XSIAM Platform Architecture. Technical Whitepaper, Palo Alto Networks. 2024. URL: <https://www.paloaltonetworks.com/resources/techbriefs/cortex-xsiam>

The Eastin-Knill Theorem: Fundamental Limitations of Quantum Fault Tolerance

UDK 621.395.7 (043.2)

Yevgen Kotukh¹

Yevhenii Bereznyak Military Academy, ¹yevgenkotukh@gmail.com

The construction of a large-scale universal quantum computer remains constrained by the inherent fragility of quantum states, which are continuously subject to decoherence, dephasing and energy dissipation. Reliable quantum information processing therefore relies on quantum error correction (QEC) codes, in which a single logical qubit is encoded into an entangled state of many physical qubits. Within the framework of fault-tolerant quantum computing (FTQC), logical operations on encoded data must be implemented in a manner that prevents the uncontrolled propagation of physical errors. The principal mechanism that ensures this property is transversality: a transversal logical gate acts as a tensor product of local unitaries on the constituent subsystems, guaranteeing that a single physical fault propagates to at most one qubit per code block [1].