

practical implications of the theorem, while fundamental, do not preclude scalable universal fault-tolerant quantum computation.

1. 1. Eastin B., Knill E. Restrictions on transversal encoded quantum gate sets. Physical Review Letters, 2009, vol. 102, no. 11, art. 110502.
2. 2. Hayden P., Nezami S., Popescu S., Salton G. Error correction of quantum reference frame information. PRX Quantum, 2021, vol. 2, art. 010326.
3. 3. Faist P., Nezami S., Albert V. V., Salton G., Pastawski F., Hayden P., Preskill J. Continuous symmetries and approximate quantum error correction. Physical Review X, 2020, vol. 10, art. 041018.
4. 4. Zhou S., Liu Z.-W., Jiang L. New perspectives on covariant quantum error correction. Quantum, 2021, vol. 5, art. 521.
5. 5. Kubica A., Demkowicz-Dobrzański R. Using quantum metrological bounds in quantum error correction: a simple proof of the approximate Eastin-Knill theorem. Physical Review Letters, 2021, vol. 126, art. 150503.
6. 6. Alexander R. A new approximate Eastin-Knill theorem. npj Quantum Information, 2025, vol. 11, art. 156.
7. 7. Bravyi S., Kitaev A. Universal quantum computation with ideal Clifford gates and noisy ancillas. Physical Review A, 2005, vol. 71, art. 022316.
8. 8. Litinski D. Magic state distillation: not as costly as you think. Quantum, 2019, vol. 3, art. 205.
9. 9. Yoder T. J., Takagi R., Chuang I. L. Universal fault-tolerant gates on concatenated stabilizer codes. Physical Review X, 2016, vol. 6, art. 031039.
10. 10. Bravyi S., Haah J. Magic-state distillation with low overhead. Physical Review A, 2012, vol. 86, art. 052329.
11. 11. Quantinuum collaboration. Experimental demonstration of high-fidelity logical magic states from code switching. arXiv:2506.14169, 2025.

Пояснюване AI/ML-виявлення аномалій у мікросервісних та мультимарних середовищах

УДК 004.056.5:004.75:004.8

Віталій Криворучко¹

Державний університет інформаційно-комунікаційних технологій,

¹hamandes@gmail.com

У мікросервісних та мультимарних середовищах засоби моніторингу формують великі обсяги телеметрії: метрики ресурсів, мережеві події, журнали доступу, HTTP-коди помилок та дані про взаємодію сервісів. Традиційний пороговий контроль часто виявляє лише факт відхилення і не пояснює, чому стан є ризиковим. Для кібербезпеки це ускладнює розмежування інциденту, деградації сервісу та звичайного коливання навантаження [1].

Мультимарне середовище розглядається як сукупність сервісів, розгорнутих у різних хмарних доменах, регіонах або провайдерах. Метою роботи є розроблення підходу до пояснюваного AI/ML-виявлення аномалій,

який поєднує аналіз телеметрії, графову модель взаємодії сервісів та механізм формування пояснення для аналітика безпеки.

Наукова новизна полягає у переході від ізольованої оцінки метрик до комбінованої оцінки інцидентного стану вузла за трьома компонентами: невідповідність телеметрії, зміна топологічного оточення та ймовірність ризикової події.

Схема не обмежується бінарним рішенням, а формує пояснення типу «вузол - ознака - зв'язок», що наближує модель до використання у SOC/DevSecOps-процесах [2, 3].

Стан вузла v у момент часу t описується вектором $x(v,t) = \{r_cpu, r_mem, r_io, r_net, l_p95, e_5xx, d_auth\}$, де всі ознаки нормовано до шкали $[0;1]$. Інфраструктура подається графом $G_t = (V_t, E_t, X_t)$, де V_t - множина сервісів, E_t - залежності між ними, X_t - матриця атрибутів вузлів. Реконструкційна модель формує оцінку $\hat{x}(v,t)$, після чого локальна похибка визначається як

$$e_{rec}(v, t) = \frac{\|x(v, t) - \hat{x}(v, t)\|_2}{\sqrt{m}},$$

де m - кількість ознак. Щоб похибка була сумісною з іншими складовими ризику, вона нормується через 95-й перцентиль похибок нормального режиму:

$$E_{rec}(v, t) = \min\left(1; \frac{e_{rec}(v, t)}{q_{95}}\right).$$

Топологічне відхилення визначається через відстань Жаккара між поточним оточенням вузла $N_t(v)$ та сталонним оточенням $N_{ref}(v)$:

$$D_{top}(v, t) = 1 - \frac{|N_t(v) \cap N_{ref}(v)|}{|N_t(v) \cup N_{ref}(v)|}.$$

Ймовірність інцидентного стану позначається як $R_{inc}(v,t)$ і може бути отримана з класифікатора або каліброваної моделі ризику. Підсумкова оцінка ризиковості вузла задається формулою

$$S(v, t) = \lambda_1 E_{rec}(v, t) + \lambda_2 D_{top}(v, t) + \lambda_3 R_{inc}(v, t), \quad \lambda_1 + \lambda_2 + \lambda_3 = 1.$$

де всі складові знаходяться у межах $[0;1]$. Вузол вважається ризиковим, якщо $S(v,t) \geq \theta$. Для оцінювання пояснюваності введено показник $EC = N_full / N_detected$, де N_full - кількість виявлених інцидентів, для яких сформовано повне пояснення «вузол - ознака - зв'язок», а $N_detected$ - загальна кількість виявлених інцидентів. Це дозволяє оцінити не лише точність, а й практичну корисність моделі.

У контрольному сценарії підхід підвищив F1-score з 0,78 до 0,89, зменшив нормований рівень хибних спрацювань на 26% та скоротив час інтерпретації інциденту на 32%. Значення $EC=0,86$ означає, що для 86% ризикових станів сформовано повне пояснення.

Таблиця 1

Оцінювання результатів виявлення ризикових станів

Показник	Базовий моніторинг	Запропонований підхід
F1-score	0,78	0,89
False Positive Rate	1,00	0,74
Час інтерпретації інциденту	1,00	0,68
Explainability Coverage	-	0,86

Висновки. Запропоновано формалізовану схему пояснюваного AI/ML-виявлення аномалій у мікросервісних та мультихмарних середовищах. Нормовані компоненти E_{res} , D_{top} та R_{inc} роблять оцінку $S(v,t)$ інтерпретованою і придатною для порівняння вузлів. На практиці це зменшує хибні спрацювання, підвищує якість виявлення інцидентів і пояснює причини ризикового стану.

1. Chandola V., Banerjee A., Kumar V. Anomaly detection: A survey. ACM Computing Surveys. - 2009. - Vol. 41, №3. - P. 1-58.
2. Ahmed M., Mahmood A.N., Hu J. A survey of network anomaly detection techniques. Journal of Network and Computer Applications. - 2016. - Vol. 60. - P. 19-31.
3. Ying R., Bourgeois D., You J., Zitnik M., Leskovec J. GNNExplainer: Generating Explanations for Graph Neural Networks. Advances in Neural Information Processing Systems. - 2019. - P. 9240-9251.

Метрикове оцінювання зменшення технічного боргу JavaScript-коду як передумови підвищення безпеки програмних систем

УДК 004.41:004.8

Ірина Замрій¹, Олексій Кулаков²

*Державний університет інформаційно-комунікаційних технологій,
¹i.zamrii@duikt.edu.ua, ²o.kulakov@stud.duikt.edu.ua*

У сучасних JavaScript-проектах технічний борг проявляється як зниження зв'язності модулів, зростання міжмодульних залежностей та ускладнення трасування змін. У контексті кібербезпеки такі властивості є непрямыми чинниками ризику, оскільки ускладнюють аудит коду, статичний аналіз, локалізацію потенційно вразливих ділянок і безпечне внесення виправлень. Роботи з LLM-рефакторингу підкреслюють перспективність автоматизованих змін, але вказують на потребу їх метрикової та функціональної валідації [1–3].

Метою роботи є розроблення компактної процедури метрикового оцінювання зменшення технічного боргу JavaScript-коду після LLM-згенерованих змін для підвищення безпеки програмних систем. Наукова новизна полягає у використанні інтегрального показника, який поєднує зв'язність, залежність і трасувальну невпорядкованість, а рішення про прийняття зміни пов'язує з додатним приростом якості. Такий підхід