

токени чи бали за використання безпечних налаштувань приватності. Ці рішення є реалістичними завдяки сучасним технологіям, таким як Ethereum для блокчейну та доступним бібліотекам машинного навчання, і можуть бути впроваджені платформами чи сторонніми розробниками.

Висновки. Інтеграція машинного навчання та блокчейн-технологій пропонує ефективний підхід до підвищення приватності та безпеки в соціальних медіа. Запропонована система аномального виявлення та АВАС забезпечує реальний захист від загроз, тоді як децентралізовані платформи та інноваційні інструменти, такі як ШІ-помічники та біометрична автентифікація, відкривають нові можливості для безпечного цифрового середовища. Впровадження цих рішень може значно знизити ризики та підвищити довіру користувачів до соціальних медіа.

1. Randa Aljably, "Preserving Privacy in Multimedia Social Networks Using Machine Learning Anomaly Detection." Security and Communication Networks, 2020. URL: <https://www.hindawi.com/journals/scn/2020/5874935/> (дата звернення 29.04.2025)

2. Мельник К.С. Обробка та захист персональних даних в соціальних мережах. Київ: Інститут інформації, безпеки і права Національної академії правових наук України, 2014. URL: <https://ippi.org.ua/sites/default/files/14mksdsm.pdf> (дата звернення 29.04.2025)

3. Top 7 Blockchain Social Media Platforms In 2024. Flying V Group, 2024. URL: <https://www.flyingvgroup.com/blockchain-social-media/> (дата звернення 29.04.2025)

4. Витік даних Facebook: що сталося і як захистити себе. *Українська правда*, 2021. URL: <https://www.pravda.com.ua/news/2021/04/4/7288929/> (дата звернення 29.04.2025)

5. Blockchain Social Media - Towards User-Controlled Data. *LeewayHertz*, 2019. URL: <https://www.leewayhertz.com/blockchain-social-media-platforms/> (дата звернення 29.04.2025)

6. How Blockchain Can Solve Social Media Privacy. *Investopedia*. 2018. URL: <https://www.investopedia.com/news/ethereum-blockchain-social-media-privacy-problem-linkedin-indorse/> (дата звернення 29.04.2025)

7. Дослідження медіаспоживання українців: третій рік повномасштабної війни. Київ: ОПОРА, 2024. URL: <https://www.oporaua.org/viyna/doslidzhennya-mediaspozhyvannya-ukrayinciv-tretyi-rik-povnomasshtabnoyi-viyni-25292> (дата звернення 29.04.2025)

Застосування ML-моделі для виявлення SQL-ін'єкцій у веб-додатках на Flask та Node.js

УДК 004.056.5:004.75:004.738.5

Іванна Мелько¹, Ігор Ігнатєв²

^{1,2}Західноукраїнський національний університет

¹ivannamelko59@gmail.com, ²iiv@wunu.edu.ua

У роботі розглянуто можливість виявлення SQL-ін'єкцій за допомогою алгоритмів машинного навчання. Подано приклад реалізації простої моделі, яка

класифікує запити як легітимні або потенційно шкідливі. Система може бути інтегрована у веб-додатки, створені з використанням Flask або Node.js[4]. Запропоноване рішення забезпечує високу точність та швидкість реагування при мінімальних ресурсних витратах, що дозволяє використовувати його навіть у невеликих проєктах.

SQL-ін'єкції залишаються однією з найпоширеніших загроз у веб-додатках. Вони дозволяють зловмиснику маніпулювати запитами до бази даних, обходити автентифікацію або зчитувати конфіденційну інформацію. Звичайні методи фільтрації не завжди встигають адаптуватись до нових видів атак, особливо zero-day[1]. Саме тому машинне навчання стає перспективним напрямом для побудови більш гнучких та інтелектуальних захисних механізмів.

Мета дослідження — розробити та продемонструвати підхід до виявлення SQL-ін'єкцій у HTTP-запитах на основі машинного навчання, з можливістю інтеграції в веб-додатки на Flask і Node.js.



Рис.1. Схема обробки HTTP-запиту з ML-детектором SQL-ін'єкцій

Для формалізації процесу прийняття рішення можна скористатись наступною формулою класифікації:

$$P(\text{SQLi} \mid x) = (1 / N) \sum f_i(x) \tag{1}$$

де: x — векторизований HTTP-запит, $f_i(x)$ — рішення i -го дерева у випадковому лісі, N — загальна кількість дерев, $P(\text{SQLi} \mid x)$ — ймовірність, що запит є SQL-ін'єкцією.

Приклад обчислення:

Тобто модель визначає запит як SQLi з ймовірністю 60%, згідно вимог я правильно написати 1. Підписи до рисунків і таблиць

$$P(\text{SQLi} \mid x) = 1/5(1 + 1 + 1 + 0 + 0) = 0.6 \tag{2}$$

Таблиця 1

Ефективність моделей		
Модель	Точність	Час реакції
Logistic Regression	92.4%	2.1 мс
Decision Tree	94.3%	3.0 мс
Random Forest	96.8%	3.4 мс

Запропонована модель на основі машинного навчання демонструє високу ефективність у виявленні SQL-ін'єкцій у веб-запитах. Вона працює швидко, не потребує великих обчислювальних ресурсів і легко інтегрується у популярні фреймворки, зокрема Flask та Node.js[3].

Модель стабільно розпізнає типові форми шкідливих запитів, а її адаптивність дозволяє враховувати нові сценарії атак. Такий підхід можна розширити й на інші типи загроз, зокрема на бот-активність, аномальні дії користувачів або zero-day атаки.

Нижче наведено приклад реалізації простої ML-моделі для виявлення SQL-ін'єкцій з використанням Python, бібліотек scikit-learn та Flask. Модель аналізує HTTP-запити та класифікує їх як легітимні або SQL-ін'єкційні[1].

- Збір навчальних даних з прикладами запитів (SQLi та OK).
- Побудова моделі з використанням TF-IDF та Random Forest.
- Тестування моделі на нових запитах.
- Збереження моделі у форматі .pkl (joblib)[2].
- Створення Flask-сервера, що приймає запит і повертає результат (SQLi або OK).

Такий підхід дозволяє легко вбудувати інтелектуальний захист у будь-який REST API та проводити перевірку запитів у реальному часі. Рішення підходить як для Flask, так і для Node.js через HTTP-запити.

1. OWASP Foundation. SQL Injection. URL: https://owasp.org/www-community/attacks/SQL_Injection
2. Scikit-learn documentation. URL: <https://scikit-learn.org/stable/>
3. Géron A. Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow. O'Reilly, 2019.
4. RFC 7231. URL: <https://datatracker.ietf.org/doc/html/rfc7231>

Розробка застосунку для підвищення рівня кібербезпеки від атак методами соціальної інженерії

УДК 62

Маргарита Мельник¹, Денис Завадський²

¹SET University, ²Національний університет, "Одеська Політехніка"

¹ m.melnyk@setuniversity.edu.ua ²9560447@stud.op.edu.ua

Соціальна інженерія є особливо небезпечною ще через те, що вона не потребує високотехнологічних інструментів. Атаки в соціальній інженерії часто спрямовані на виявлення слабких місць у людському мисленні або поведінці, що робить їх особливо важкими для виявлення і нейтралізації. Тому навіть найсучасніші технології захисту не можуть гарантувати повну безпеку, якщо користувачі не навчені реагувати на ці загрози.

Успіх соціальної інженерії значною мірою базується на нестачі обізнаності серед користувачів. Люди не завжди розпізнають потенційну загрозу в листі, дзвінку чи особистій розмові. Захист від такого виду атак передбачає впровадження різноманітних методів навчання користувачів у кіберпросторі.

Своєчасне виявлення та попередження спроб соціальної інженерії значною мірою знижує ризики витоку даних. Захист починається з людини - найслабшого, але й найважливого елементу інформаційної безпеки. У час стрімкого розвитку цифрових технологій для протидії таким методам атак важливе значення має систематичне навчання користувачів. Необхідно впроваджувати комплексні програми з підвищення кіберграмотності. Існують різні програми навчання, які орієнтовані на формування у користувачів навичок критичного мислення та обережності у взаємодії з цифровими технологіями. Одним з найбільш ефективних підходів є створення навчальних курсів та