

АНОТАЦІЯ

Терещенко О.І. Методи і моделі машинного навчання для виявлення та класифікації вразливостей смарт-контрактів. – Кваліфікаційна наукова праця на правах рукопису.

Дисертація на здобуття наукового ступеня доктора філософії за спеціальністю 122 – Комп'ютерні науки. – Національний університет «Одеська політехніка», МОН України, Одеса, 2025.

У **вступі** обґрунтовано актуальність проблеми безпеки смарт-контрактів у блокчейн-екосистемах (зокрема у децентралізованих фінансах), які акумулюють значні фінансові активи та залишаються вразливими до помилок коду. Сформульовано **мету дослідження** – підвищити ефективність за відповідними метриками виявлення, класифікації та локалізації вразливостей у смарт-контрактах за рахунок розробки моделей та методів машинного навчання, які становлять цілісний технологічний стек інтелектуального аналізу вихідного коду смарт-контрактів.

Визначено **об'єкт дослідження** – процеси виявлення, класифікації та локалізації вразливостей у смарт-контрактах. **Предметом дослідження** є методи та моделі машинного навчання для багатоміткового виявлення, класифікації та локалізації вразливостей у смарт-контрактах.

Сформульовано комплекс взаємопов'язаних завдань, серед яких: створення збалансованого та репрезентативного навчального датасету смарт-контрактів із анотацією вразливостей; розроблення моделей і методів багатоміткової класифікації вразливостей смарт-контрактів; підвищення надійності й стабільності моделей за рахунок калібрування впевненості та аналізу стійкості до збурень та змін коду; розроблення слабосупервізованих методів локалізації вразливих фрагментів коду без токенної розмітки; апробація запропонованих методів шляхом їх реалізації у програмному засобі та перевірки ефективності на реальних смарт-контрактах.

У **першому розділі** здійснено систематичний огляд підходів до виявлення вразливостей у смарт-контрактах, що охоплює формальні методи, статичний та динамічний аналіз, а також сучасні методи на основі машинного навчання. Показано,

що кожен клас методів має суттєві обмеження, які стримують їхнє використання у великих блокчейн-екосистемах. Формальні методи забезпечують максимальну строгість і дозволяють строго довести коректність коду, однак є надзвичайно дорогими, потребують значної експертної участі та масштабуються лише на обмежену кількість критичних контрактів. Статичний аналіз є найбільш продуктивним підходом, здатним швидко сканувати великі репозиторії, але страждає від значної кількості хибнопозитивних спрацювань і часто не охоплює нетипові або нові патерни вразливостей. Динамічний аналіз наближають перевірку до реальних сценаріїв атак і добре виявляють логічні помилки, але є повільними, не гарантують повного покриття й сильно залежать від заданих сценаріїв взаємодії. Методи машинного навчання демонструють хорошу масштабованість і здатність виявляти складні закономірності, проте у практичних сценаріях працюють як «чорна скринька», що ускладнює інтерпретацію рішень і локалізацію вразливих фрагментів.

Окремо проведено аналіз трансформерних моделей для коду (CodeBERT/GraphCodeBERT), графові нейронні мережі (Graph Neural Networks, GNN) та гібридних трансформерно-рекурентних архітектур. Виявлено ключові практичні обмеження їхнього застосування: нестачу якісної фрагментної розмітки для локалізації, низьку стійкість до обфускації та рефакторингу без застосування спеціальних методів підвищення стійкості до збурень та змін, а також схильність моделей до надмірної впевненості у прогнозах під час роботи в production-середовищах. Сформовано рамку оцінювання: для детекції використано micro-F1 (мікроусереднену F1-міру) та macro-F1 (макроусереднену F1-міру), а також ROC-AUC (площу під кривою робочої характеристики приймача, Receiver Operating Characteristic – Area Under the Curve); для локалізації застосовано mAP (mean Average Precision – середню точність), IoU@ τ (Intersection over Union при пороговому значенні τ – міру перекриття передбачених та істинних областей) і Hit@K (частоту потрапляння релевантного фрагмента серед перших K прогнозів); для оцінювання

надійності використано метрику очікуваної калібрувальної похибки (Expected Calibration Error, ECE) та пов'язані показники.

У **другому розділі** обґрунтовано й реалізовано метод формування збалансованого навчального датасету смарт-контрактів для задач виявлення вразливостей. Запропоновано комбінований підхід, що поєднує детерміноване інжектування типових вразливостей у безпечний код (pattern-based + статичний аналіз) та контекстно узгоджену генерацію/модифікацію прикладів за допомогою великої мовної моделі. Обидва канали супроводжуються валідацією (компіляція, перевірка статичними аналізаторами, контроль формату та відповідності міток). На виході сформовано збалансований корпус із п'яти категорій (чотири класи вразливостей та безпечні контракти) по 1000 прикладів кожної.

Запропонований комбінований метод формування датасету забезпечує одночасно відтворюваність і контрольованість (детерміновані зміни з фіксованими правилами), реалістичну різноманітність синтаксичних реалізацій (LLM-модифікації, що зберігають логіку та стилістику коду) й збалансованість за класами, що підвищує якість подальшого навчання та порівняльної оцінки моделей.

У **третьому розділі** розроблено комплексний підхід до багатоміткової детекції та слабосупервізованої локалізації вразливостей у смарт-контрактах. Як базові розглянуто моделі BiGRU-ATT та CodeBERT-BiGRU для багатоміткової класифікації. Далі запропоновано модель CodeBERT-GDLA (Graph-aware Dual-Level Attention), що поєднує токенний рівень (CodeBERT + BiGRU) з блоковим рівнем і графовим зсувом. Удосконалено механізм уваги: введено графовий зсув на основі AST (Abstract Syntax Tree), CFG (Control Flow Graph), DFG (Data Flow Graph) графів та навчувані коефіцієнти β , γ для токенного й блочного рівнів, які керують часткою структурної інформації в логітах уваги й підсилюють чутливість до далеких/непрямих залежностей у коді.

Запропоновано CAM-узгоджену увагу: увага моделі узгоджується з Grad-CAM через λ -комбінацію, що забезпечує послідовні, інтерпретовані карти важливості на

токенному та блоковому рівнях. Створено метод слабосупервізованої локалізації, який за допомогою ковзного вікна з top-K агрегацією перетворює логіти на просторові карти релевантності, а причинні регуляризатори дозволяють виділяти мінімально достатні та критично необхідні фрагменти коду без ручної токенної розмітки.

Експериментальні дослідження засвідчили, що запропонована модель CodeBERT-GDLA перевершує базову CodeBERT-BiGRU за macro-F1, macro-AUC і втратою Хеммінга (Hamming loss), а запропонований локалізатор вразливостей досягає високих показників метрик mAP та Hit@K (зокрема для вразливості повторного входу – $mAP \approx 82\%$, $Hit@5 \approx 77\%$, $Hit@7-10 \approx 100\%$) із низькою дисперсією, що підтверджує стабільність локалізації. Проведено абляційний аналіз внеску графових зсувів, дворівневої уваги, причинних регуляризаторів та SAM-узгодженої уваги, а також комплексне дослідження стійкості до збурень та змін, а саме до рефакторингів, перейменувань, форматування й часткової обфускації.

Щоб зменшити ризики надмірної впевненості та підвищити надійність моделей у задачах багатоміткової детекції, у тренувальний процес інтегровано запропонований метод регуляризації для калібрування ймовірнісних оцінок моделей глибокого навчання, який знижує метрику ECE (Expected Calibration Error) і робить ймовірнісні оцінки узгодженішими з реальною частотою подій. Крім того, проведено дослідження стійкості моделі до типових модифікацій коду (рефакторинг, перейменування ідентифікаторів, стилістичні зміни, часткова обфускація), що підтвердило стабільність як класифікації, так і локалізації вразливостей у реалістичних сценаріях використання.

У **четвертому розділі** представлено практичну реалізацію й впровадження програмного засобу аудиту смарт-контрактів на основі моделі CodeBERT-GDLA. Описано вимоги, архітектуру, REST-інтерфейси, користувацький інтерфейс та кейси використання в реальних сценаріях аудиту. Розроблений інструмент інтегрується в CI/CD-конвеєри та може доповнювати існуючі статичні аналізатори.

Наукова новизна одержаних результатів полягає в наступному:

– **вперше запропоновано комбінований метод формування навчального датасету** за рахунок поєднання автоматичного інжектування аномальних конструкцій у коректний код смарт-контрактів із генерацією додаткових прикладів великою мовною моделлю з подальшою верифікацією, що дозволило забезпечити збалансованість вибірки;

– **вперше побудовано модель CodeBERT-GDLA** на основі вдосконалення архітектур BiGRU-ATT і CodeBERT-BiGRU за допомогою запропонованих методів графового зсуву і дворівневої графово-орієнтованої уваги, що дозволило підвищити ефективність виявлення та класифікації вразливостей при інтелектуальному аналізі вихідного коду смарт-контрактів;

– **удосконалено метод регуляризації для калібрування ймовірнісних оцінок багатоміткової класифікації вразливостей смарт-контрактів** за рахунок застосування оптимізації очікуваної калібрувальної похибки, температурного масштабування та фокального приглушення надмірної впевненості безпосередньо в процесі навчання моделі CodeBERT-GDLA, що забезпечує зниження її надмірної впевненості та підвищує надійність результатів інтелектуального аналізу вихідного коду смарт-контрактів;

– **отримав подальший розвиток метод слабосупервізованого виділення вразливих фрагментів коду** за рахунок інтеграції САМ-узгодженої уваги з причинно-орієнтованою регуляризацією, що забезпечує підвищення ефективності локалізації вразливостей при інтелектуальному аналізі вихідного коду смарт-контрактів.

Практичне значення результатів полягає в розробленні програмно-алгоритмічних засобів для автоматизованого аудиту смарт-контрактів, які забезпечують багатоміткове виявлення вразливостей, пояснювану локалізацію фрагментів коду та надійні ймовірнісні оцінки, а також у створенні відтворюваного конвеєра даних і збалансованого датасету для навчання й порівняльної оцінки моделей. Запропоновані моделі й методи впроваджено у діяльності ТОВ НВО

«Діскрет» та НВП «КАРЕ», а також використано в навчальному процесі Національного університету «Одеська політехніка».

Ключові слова: смарт-контракти; машинне навчання; глибоке навчання; автоматизація аудиту; безпека програмного коду; виявлення вразливостей; локалізація вразливостей; багатоміткова класифікація; слабосупервізоване навчання; пояснюваний ШІ; графові подання коду; механізм уваги; причинні регуляризатори; калібрування ймовірностей.

Список публікацій здобувача за темою дисертації

1. Komleva, N. O., & Tereshchenko, O. I. (2023). Requirements for the development of smart contracts and an overview of smart contract vulnerabilities at the Solidity code level on the Ethereum platform. *Herald of Advanced Information Technology*, 6(1), 54–68. DOI: <https://doi.org/10.15276/hait.06.2023.4> (**Scopus**)
2. Tereshchenko, O., & Komleva, N. (2023). Vulnerability detection of smart contracts based on bidirectional GRU and attention mechanism. In *Communications in Computer and Information Science* (Vol. 1980). Springer, Cham. DOI: https://doi.org/10.1007/978-3-031-48325-7_21 (**Scopus**)
3. Tereshchenko, O. I., & Komleva, N. O. (2024). Identification and localization of vulnerabilities in smart contracts using attention vectors analysis in a BERT-based model. *Radio Electronics, Computer Science, Control*, (3), 173–184. DOI: <https://doi.org/10.15588/1607-3274-2024-3-15> (**Web of Science**)
4. Терещенко, О. І. (2025). Комбінований метод локалізації вразливостей у смарт-контрактах на основі Attention та Grad-CAM. *Таврійський науковий вісник. Серія: Технічні науки*, 1(4), 306-316. DOI: <https://doi.org/10.32782/tnv-tech.2025.4.1.30>
5. Терещенко, О. І. (2025). Методи ін'єкції вразливостей у смарт-контракти для формування збалансованих датасетів. *Вісник Херсонського національного технічного університету*. № 3(94), Ч. 2, 456-462. DOI: <https://doi.org/10.35546/kntu2078-4481.2025.3.2.58>
6. Терещенко, О. І. (2025). Виявлення та локалізація вразливостей у смарт-контрактах на основі дворівневої уваги з графовим підкріпленням. *Вісник Кременчуцького національного університету імені Михайла Остроградського*, 4, 272-281. DOI: <https://doi.org/10.32782/1995-0519.2025.4.30>
7. Терещенко, О. І., & Комлева, Н. О. (2023). Огляд наявних інструментів виявлення вразливостей смарт-контрактів на основі машинного навчання. In *Modern*

Information Technology 2023 / Сучасні Інформаційні Технології 2023 (Одеса, 18–19 травня 2023 р.). URL: <http://dspace.opu.ua/jspui/handle/123456789/14202>

8. Tereshchenko, O. (2024). Analysis of the efficacy of using machine learning for detecting vulnerabilities in smart contracts compared to traditional methods. In Proceedings of the XIII International Scientific and Practical Conference “Social Ways of Training Specialists in the Social Sphere and Inclusive Education” (pp. 327–329). Prague, Czech Republic, April 1–3, 2024. URL: <https://eu-conf.com/en/events/social-ways-of-training-specialists-in-the-social-sphere-and-inclusive-education/>

9. Tereshchenko, O. (2024). Improving the security of smart contracts through the integration of machine learning: Challenges and prospects. In Proceedings of the VI International Scientific Conference “Traditional and Innovative Approaches to Scientific Research” (pp. 124–125). Vinnytsia, Ukraine, March 8, 2024. URL: <https://archive.mcnd.org.ua/index.php/conference-proceeding/issue/view/08.03.2024>

10. Tereshchenko, O. I. (2024). Integration of NLP and machine learning methods for smart contract security: A comparison with traditional approaches. In Proceedings of the X International Scientific and Practical Conference “Informatics. Culture. Technology”, September 2024. DOI: <https://doi.org/10.15276/ict.01.2024.31>

11. Терещенко, О. І. (2025). Побудова збалансованого корпусу смарт-контрактів шляхом контрольованої ін’єкції вразливостей. In Proceedings of the VIII International Scientific Conference “Розвиток наукової думки постіндустріального суспільства: сучасний дискурс”, 2025. DOI: <https://doi.org/10.62731/mcnd-17.10.2025>

ABSTRACT

Tereshchenko O.I. Machine Learning Methods and Models for Detecting and Classifying Smart-Contract Vulnerabilities. – Qualification scientific work in the form of manuscript.

Thesis for the PhD degree in specialty 122 Computer science. – Odessa Polytechnic National University, Ministry of Education and Science of Ukraine, Odesa, 2025.

In the introduction, the relevance of the smart contract security problem in blockchain ecosystems (in particular, decentralized finance) is substantiated, as these systems accumulate substantial financial assets while remaining vulnerable to software defects. **The aim of the research** is formulated as improving the effectiveness, according to relevant evaluation metrics, of detecting, classifying, and localizing vulnerabilities in smart contracts through the development of machine learning methods and models that constitute an integrated technological stack for intelligent analysis of smart contract source code.

The object of the research is the processes of detecting, classifying, and localizing vulnerabilities in smart contracts. **The subject of the research** is machine learning methods and models for multi-label detection, classification, and localization of vulnerabilities in smart contracts.

A set of interrelated **research tasks** is formulated, including: the creation of a balanced and representative training dataset of smart contracts with vulnerability annotations; the development of models and methods for multi-label vulnerability classification in smart contracts; improving model reliability and stability through confidence calibration and robustness analysis with respect to code changes; the development of weakly supervised methods for localizing vulnerable code fragments without token-level annotations; and validation of the proposed methods through their implementation in a software tool and evaluation of their effectiveness on real-world smart contracts.

The first chapter presents a systematic review of approaches to smart contract vulnerability detection, covering formal methods, static and dynamic analysis, as well as modern machine learning-based techniques. It is shown that each class of methods has significant limitations that constrain their use in large-scale blockchain ecosystems. Formal methods provide maximal rigor and enable formal proofs of correctness, but are extremely costly, require substantial expert involvement, and scale only to a limited number of critical contracts. Static analysis is the most productive approach in terms of throughput and can rapidly scan large repositories, but suffers from a high rate of false positives and often fails to capture atypical or novel vulnerability patterns. Dynamic analysis brings testing closer to real attack scenarios and is effective at detecting logical flaws, but is slow, does not guarantee full coverage, and heavily depends on predefined interaction scenarios. Machine learning methods demonstrate good scalability and the ability to capture complex patterns; however, in practical settings they often operate as a “black box”, complicating decision interpretation and the localization of vulnerable fragments.

A separate analysis is conducted of transformer-based models for code (CodeBERT/GraphCodeBERT), graph neural networks (GNNs), and hybrid transformer-recurrent architectures. Key practical limitations of their application are identified, including the lack of high-quality fragment-level annotations for localization, low robustness to obfuscation and refactoring without dedicated robustness-enhancing techniques, and a tendency toward overconfident predictions in production environments. An evaluation framework is defined: for detection, micro-F1 and macro-F1, as well as ROC-AUC (Receiver Operating Characteristic – Area Under the Curve), are used; for localization, mAP (mean Average Precision), IoU@ τ (Intersection over Union at threshold τ), and Hit@K (the frequency with which a relevant fragment appears among the top-K predictions) are applied; for reliability assessment, Expected Calibration Error (ECE) and related metrics are employed.

The second chapter substantiates and implements a method for forming a balanced training dataset of smart contracts for vulnerability detection tasks. A combined approach is

proposed that integrates deterministic injection of typical vulnerabilities into safe code (pattern-based methods with static analysis) and context-consistent generation/modification of examples using a large language model. Both channels are accompanied by validation procedures (compilation, checks with static analyzers, and format/label consistency control). As a result, a balanced corpus of five categories (four vulnerability classes and safe contracts) with 1,000 examples per category is produced.

The proposed combined dataset construction method simultaneously ensures reproducibility and controllability (deterministic changes governed by fixed rules), realistic diversity of syntactic implementations (LLM-based modifications that preserve code logic and style), and class balance, thereby improving the quality of subsequent training and comparative model evaluation.

In the third chapter, a comprehensive approach to multi-label detection and weakly supervised localization of vulnerabilities in smart contracts is developed. BiGRU-ATT and CodeBERT-BiGRU models are considered as baselines for multi-label classification. Subsequently, the CodeBERT-GDLA (Graph-aware Dual-Level Attention) model is proposed, which combines a token-level component (CodeBERT + BiGRU) with a block-level component and graph-based biasing. The attention mechanism is enhanced by introducing graph bias derived from AST (Abstract Syntax Tree), CFG (Control Flow Graph), and DFG (Data Flow Graph), along with learnable coefficients β and γ for token- and block-level attention that control the contribution of structural information in attention logits and increase sensitivity to long-range and indirect dependencies in code.

A CAM-aligned attention mechanism is proposed, in which model attention is aligned with Grad-CAM via a λ -combination, yielding consistent and interpretable importance maps at both token and block levels. A weakly supervised localization method is introduced that converts logits into spatial relevance maps using a sliding window with Top-K aggregation, while causal regularizers enable the identification of minimally sufficient and critically necessary code fragments without manual token-level annotation.

Experimental studies demonstrate that the proposed CodeBERT-GDLA model outperforms the baseline CodeBERT-BiGRU in terms of macro-F1, macro-AUC, and Hamming loss, while the proposed vulnerability localizer achieves high mAP and Hit@K scores (in particular for reentrancy vulnerabilities: mAP $\approx 82\%$, Hit@5 $\approx 77\%$, Hit@7–10 $\approx 100\%$) with low variance, confirming localization stability. An ablation study evaluates the contributions of graph biasing, dual-level attention, causal regularizers, and CAM-aligned attention, along with a comprehensive robustness analysis with respect to refactoring, renaming, formatting, and partial obfuscation.

To mitigate the risks of overconfidence and improve model reliability in multi-label detection tasks, a proposed regularization method for calibrating probabilistic outputs of deep learning models is integrated into the training process. This method reduces Expected Calibration Error (ECE) and aligns predicted probabilities with empirical event frequencies. Additionally, robustness to common code modifications (refactoring, identifier renaming, stylistic changes, and partial obfuscation) is evaluated, confirming the stability of both classification and localization under realistic usage scenarios.

The fourth chapter presents the practical implementation and deployment of a smart contract auditing tool based on the CodeBERT-GDLA model. System requirements, architecture, REST interfaces, user interface, and real-world audit use cases are described. The developed tool integrates into CI/CD pipelines and can complement existing static analyzers.

The **scientific novelty** of the obtained results is as follows:

- **for the first time, a combined method for forming a training dataset is proposed**, based on the integration of automatic injection of anomalous constructs into correct smart contract code with the generation of additional samples using a large language model followed by their verification, which ensures dataset balance;
- **for the first time, the CodeBERT-GDLA model is constructed** based on the enhancement of the BiGRU-ATT and CodeBERT-BiGRU architectures through the proposed graph shifting methods and a two-level graph-aware attention mechanism, which

improves the effectiveness of vulnerability detection and classification during intelligent analysis of smart contract source code;

- **the regularization method for calibrating probabilistic estimates in multi-label vulnerability classification of smart contracts is improved** by integrating Expected Calibration Error optimization, temperature scaling, and focal suppression of overconfidence directly into the training process of the CodeBERT-GDLA model, which reduces excessive model confidence and increases the reliability of intelligent smart contract source code analysis results;

- **the method of weakly supervised identification of vulnerable code fragments is further developed** through the integration of CAM-aligned attention with causality-oriented regularization, which enhances the effectiveness of vulnerability localization during intelligent analysis of smart contract source code.

The practical significance of the results lies in the development of software and algorithmic tools for automated smart contract auditing that provide multi-label vulnerability detection, explainable localization of code fragments, and reliable probabilistic estimates, as well as in the creation of a reproducible data pipeline and a balanced dataset for training and comparative evaluation of models. The proposed models and methods have been implemented in the activities of LLC RPA “Diskret” and RPC “KARE”, and have also been used in the educational process of Odesa Polytechnic National University.

Keywords: smart contracts; machine learning; deep learning; audit automation; code security; vulnerability detection; vulnerability localization; multi-label classification; weakly supervised learning; explainable AI; graph representations of code; attention mechanism; causal regularizers; probability calibration.

List of publications of the applicant on the topic of the dissertation

1. Komleva, N. O., & Tereshchenko, O. I. (2023). Requirements for the development of smart contracts and an overview of smart contract vulnerabilities at the Solidity code level on the Ethereum platform. *Herald of Advanced Information Technology*, 6(1), 54–68. <https://doi.org/10.15276/hait.06.2023.4> **(Scopus)**
2. Tereshchenko, O., & Komleva, N. (2023). Vulnerability detection of smart contracts based on bidirectional GRU and attention mechanism. In *Communications in Computer and Information Science* (Vol. 1980). Springer, Cham. https://doi.org/10.1007/978-3-031-48325-7_21 **(Scopus)**
3. Tereshchenko, O. I., & Komleva, N. O. (2024). Identification and localization of vulnerabilities in smart contracts using attention vectors analysis in a BERT-based model. *Radio Electronics, Computer Science, Control*, (3), 173–184. <https://doi.org/10.15588/1607-3274-2024-3-15> **(Web of Science)**
4. Tereshchenko, O. I. (2025). A combined method for vulnerability localization in smart contracts based on attention mechanisms and Grad-CAM. *Tavria Scientific Bulletin. Series: Technical Sciences*, 1(4), 306–316. <https://doi.org/10.32782/tnv-tech.2025.4.1.30>
5. Tereshchenko, O. I. (2025). Methods for vulnerability injection in smart contracts for building balanced datasets. *Bulletin of Kherson National Technical University*, 3(94), Part 2, 456–462. <https://doi.org/10.35546/kntu2078-4481.2025.3.2.58>
6. Tereshchenko, O. I. (2025). Detection and localization of smart contract vulnerabilities using a two-level attention mechanism with graph-based reinforcement. *Bulletin of Kremenchuk Mykhailo Ostrohradskyi National University*, 4, 272–281. <https://doi.org/10.32782/1995-0519.2025.4.30>
7. Tereshchenko, O. I., & Komleva, N. O. (2023). A review of existing machine learning-based tools for smart contract vulnerability detection. In *Modern Information*

Technology 2023 (Odesa, May 18–19, 2023).
<http://dspace.opu.ua/jspui/handle/123456789/14202>

8. Tereshchenko, O. (2024). Analysis of the effectiveness of machine learning approaches for smart contract vulnerability detection compared to traditional methods. In Proceedings of the XIII International Scientific and Practical Conference “Social Ways of Training Specialists in the Social Sphere and Inclusive Education” (pp. 327-329). Prague, Czech Republic, April 1–3, 2024. <https://eu-conf.com/en/events/social-ways-of-training-specialists-in-the-social-sphere-and-inclusive-education/>

9. Tereshchenko, O. (2024). Improving smart contract security through the integration of machine learning: Challenges and prospects. In Proceedings of the VI International Scientific Conference “Traditional and Innovative Approaches to Scientific Research” (pp. 124-125). Vinnytsia, Ukraine, March 8, 2024. <https://archive.mcnd.org.ua/index.php/conference-proceeding/issue/view/08.03.2024>

10. Tereshchenko, O. I. (2024). Integration of NLP and machine learning methods for smart contract security: A comparison with traditional approaches. In Proceedings of the X International Scientific and Practical Conference “Informatics. Culture. Technology”, September 2024. <https://doi.org/10.15276/ict.01.2024.31>

11. Tereshchenko, O. I. (2025). Construction of a balanced smart contract corpus through controlled vulnerability injection. In Proceedings of the VIII International Scientific Conference “Development of Scientific Thought in the Post-Industrial Society: Contemporary Discourse”, 2025. <https://doi.org/10.62731/mcnd-17.10.2025>