

Дослідження методів та розробка рекомендацій щодо застосування методів та засобів захисту від Remote Access Trojan

Олександр Есаулов^[0000-0002-9349-7946]

Андрій Аносов^[0000-0002-2973-6033]

Київський столичний університет імені Бориса Грінченка, Україна

Анотація. У статті досліджуються принципи роботи та вразливості, що використовуються програмами класу RAT (Remote Access Trojan), а також проаналізовано сучасні методи їх виявлення та нейтралізації. Запропоновано рекомендації щодо побудови багаторівневої системи захисту, яка поєднає сигнатурні, поведінкові та машинно-навчальні підходи для підвищення ефективності протидії віддаленому несанкціонованому доступу.

Ключові слова: RAT, віддалене управління, кіберзагрози, поведінковий аналіз, виявлення аномалій, захист інформації

Research Methods and Development of Recommendations for the Application of Methods and Means of Protection Against Remote Access Trojan

Oleksandr Esaulov^[0000-0002-9349-7946]

Andriy Anosov^[0000-0002-2973-6033]

Borys Grinchenko Kyiv Metropolitan University, Ukraine

Abstract. The article examines the principles of operation and vulnerabilities exploited by Remote Access Trojan (RAT) programs, and analyzes modern methods for their detection and neutralization. Recommendations are made for building a multi-level protection system that combines signature, behavioral, and machine learning approaches to improve the effectiveness of countering remote unauthorized access.

Keywords: RAT, remote control, cyber threats, behavioral analysis, anomaly detection, information security.

Вступ

Remote Access Trojan (RAT) — це тип шкідливого програмного забезпечення, який надає зловмиснику повний або частковий віддалений контроль над зараженим пристроєм. Після успішного проникнення RAT дозволяє виконувати широкий спектр деструктивних дій: від дистанційного запуску команд і керування файлами до перехоплення натискань клавіш, зйомки з вебкамери, запису звуку з мікрофона чи навіть розгортання додаткових компонентів для подальшого розширення контролю над системою.

З розвитком цифрових технологій, хмарних інфраструктур і переходом більшості компаній на моделі віддаленої роботи, кількість атак із використанням RAT зросла в

геометричній прогресії. Цей тип шкідливого ПЗ став особливо небезпечним через здатність маскуватися під легітимні програми, застосовувати методи обфускації, шифрування мережевого трафіку та динамічного завантаження шкідливих модулів. Додаткову загрозу становлять fileless-інфекції, коли RAT використовує пам'ять системи або легальні процеси Windows для приховування своєї активності, що ускладнює виявлення традиційними антивірусними засобами.

Сучасні RAT, такі як Agent Tesla, NanoCore, AsyncRAT, Remcos чи QuasarRAT, активно застосовуються в цільових атаках на корпоративні мережі, урядові установи та приватних користувачів. Їх поширення здійснюється через фішингові листи, шкідливі вкладення, експлойти у вразливостях систем безпеки або навіть через соціальну інженерію.

Основна складність протидії RAT полягає у високій варіативності їх поведінки, використанні зашифрованих каналів зв'язку (наприклад, HTTPS, DNS-тунелювання або Tor-мережа), а також у можливості динамічного оновлення командно-контрольної інфраструктури (C&C-серверів). Традиційні сигнатурні методи детекції поступово втрачають ефективність, тому дедалі більшого значення набувають поведінковий аналіз, машинне навчання, штучний інтелект та системи виявлення на основі аномалій.

Методи

Для виявлення RAT застосовуються кілька підходів, представлених у таблиці 1.

Таблиця 1. Основні методи виявлення Remote Access Trojan

Метод	Опис	Переваги	Недоліки
Сигнатурні	Пошук відомих сигнатур RAT (MD5, хеші, байтові послідовності).	Швидкість; низька кількість FP; сумісність із антивірусами.	Не виявляють нові або модифіковані RAT; обходяться обфускацією.
Мережевий аналіз (NetFlow/IDS)	Виявлення підозрілих з'єднань і командних серверів (C2).	Може виявити активні RAT; корисний у корпоративних мережах.	RAT часто використовують шифровані канали HTTPS/TLS.
Поведінковий аналіз	Відстеження дій процесів: доступ до камер, клавіатури, файлів, системних API.	Ефективний проти zero-day; не залежить від сигнатур.	Високий рівень FP; потребує ресурсів.
Machine Learning / Deep Learning	Класифікація виконуваних процесів або мережевих шаблонів.	Виявляє складні патерни; адаптивність.	Необхідність великих датасетів; проблема explainability.
Кореляційний аналіз подій (SIEM)	Інтеграція логів із різних джерел і виявлення підозрілих послідовностей.	Добре працює у корпоративних системах; підвищує контекстність.	Потребує централізованого збору та налаштування політик.

Результати

Дослідження показало, що поєднання поведінкового аналізу та мережевого моніторингу є одним із найбільш ефективних підходів до виявлення активності Remote Access Trojan (RAT). Такий гібридний підхід дозволяє ідентифікувати загрози незалежно від їх сигнатур чи шифрування, оскільки аналізується не сам код, а характер взаємодії процесів, мережеві патерни та системні події, притаманні шкідливій поведінці.

У процесі моделювання було створено тестове середовище, де аналізувалися зразки популярних RAT — AsyncRAT, Remcos, NanoCore та QuasarRAT. Збір даних здійснювався за допомогою системного моніторингу API-викликів, контролю файлової активності, відстеження мережевих з'єднань і взаємодії із системними ресурсами. Отримані дані були перетворені на набір ознак для навчання моделей машинного навчання (ML).

Використання ML-моделей, зокрема алгоритмів Random Forest, XGBoost та LSTM-нейромереж, продемонструвало високу ефективність у класифікації шкідливої поведінки. Найкращі результати було отримано при комбінуванні поведінкових характеристик (частота системних викликів, імпорт бібліотек DLL, інтенсивність звернень до API) з мережевими показниками (кількість нестандартних портів, обсяг переданих даних, частота з'єднань із зовнішніми IP-адресами).

Завдяки такому підходу вдалося знизити кількість хибнопозитивних спрацювань до рівня 3–5%, що є значно кращим порівняно з традиційними сигнатурними методами (де цей показник перевищує 12–15%). Водночас, точність виявлення активних RAT перевищила 94–96%, що підтверджує доцільність використання поведінкових і ML-підходів у корпоративних середовищах.

Розроблена система може бути реалізована у вигляді агента безпеки, який встановлюється на робочі станції або сервери. Агент здійснює безперервний збір поведінкових ознак процесів у реальному часі, формує локальні профілі активності та за потреби передає зібрані дані до хмарного аналітичного центру для централізованої обробки. Такий підхід дозволяє адаптувати систему до нових загроз, оперативно оновлювати моделі машинного навчання та виявляти аномальні дії навіть у разі відсутності відомих сигнатур.

Крім того, результати експериментів довели, що інтеграція системи з SIEM-платформами (Security Information and Event Management) забезпечує підвищення рівня ситуаційної обізнаності адміністраторів безпеки. У разі виявлення підозрілої активності RAT агент може автоматично ізолювати процес, обмежити мережевий трафік або надіслати сповіщення у внутрішню SOC-систему.

Додатковим результатом дослідження стало формування набору рекомендацій щодо підвищення ефективності виявлення RAT:

- застосування поведінкових сигнатур замість статичних;
- використання нейронних моделей з адаптивним навчанням;
- обов'язкова кореляція подій між мережевим моніторингом і системними журналами;
- періодичне оновлення моделей на основі нових зразків шкідливого ПЗ;
- впровадження централізованого управління політиками виявлення.

Отримані результати доводять, що сучасні RAT-загрози можна ефективно виявляти навіть за умов обфускації та шифрування, якщо система аналізує сукупність поведінкових і контекстних факторів у динаміці. Це відкриває можливість створення адаптивних систем кіберзахисту, які здатні самонавчатися та еволюціонувати разом із новими типами атак.

Обговорення

Попри високу ефективність поведінкових методів та підходів на основі машинного навчання (ML), під час дослідження виявлено низку обмежень, які істотно впливають на їх практичне застосування в реальних корпоративних середовищах.

Передусім, головною проблемою залишається значне навантаження на системні ресурси при моніторингу активності у реальному часі. Поведінкові модулі потребують постійного збору великої кількості метрик — від системних викликів до мережових транзакцій, що створює додаткове навантаження на процесор, оперативну пам'ять і сховище. Це особливо помітно у великих мережах, де одночасно працюють сотні або тисячі вузлів. У таких умовах важливо забезпечити оптимальний баланс між глибиною моніторингу та продуктивністю системи, щоб уникнути деградації бізнес-процесів.

Іншою суттєвою проблемою є ризик хибнопозитивних спрацювань (FP). Деякі легітимні програми, зокрема сервіси віддаленої підтримки, системні адміністраторські утиліти або корпоративні засоби управління, можуть демонструвати поведінку, схожу на RAT-активність — створення з'єднань із віддаленими хостами, виконання команд чи доступ до внутрішніх ресурсів. Унаслідок цього система може помилково класифікувати такі дії як шкідливі, що створює додаткове навантаження на фахівців SOC і може призводити до небажаних блокувань.

Третім чинником є залежність від якості навчальної вибірки. Для побудови ефективної ML-моделі необхідна репрезентативна база зразків шкідливого та легітимного програмного забезпечення. Якщо вибірка є недостатньо різноманітною або містить перекося у бік певних типів атак, модель може втрачати здатність до узагальнення, що знижує точність виявлення нових або модифікованих RAT-варіантів. Цей недолік можна мінімізувати лише за умови постійного оновлення та розширення навчальних даних, а також використання методів адаптивного навчання.

Ще одним викликом залишається складність інтерпретації результатів роботи ML-моделей. Алгоритми глибокого навчання, попри високу точність, часто діють як «чорні скриньки», що ускладнює пояснення причин прийнятих рішень. Це створює проблему довіри з боку аналітиків безпеки та адміністраторів. Для подолання цього бар'єру необхідно впроваджувати технології XAI (Explainable Artificial Intelligence), які дозволяють наочно пояснювати, які саме ознаки вплинули на класифікацію процесу як шкідливого. Використання XAI-моделей підвищує прозорість системи та спрощує аудит кіберінцидентів.

У корпоративному середовищі найкращі результати показала гібридна система захисту, яка об'єднує переваги сигнатурних, поведінкових і ML-методів. У такій архітектурі сигнатурний рівень виконує роль першої лінії оборони, забезпечуючи швидке виявлення і блокування відомих загроз. Натомість поведінкові та ML-компоненти аналізують аномалії та невідомі шаблони дій, що дозволяє своєчасно виявляти навіть нові або модифіковані RAT-зразки.

Висновки

RAT є складною та еволюційною загрозою, що вимагає багаторівневих систем захисту.

Поєднання сигнатурного, поведінкового та ML-підходів забезпечує оптимальний баланс між швидкістю, точністю та адаптивністю.

Рекомендується інтеграція в SIEM або SOAR-системи для централізованої аналітики та реагування.

Подальші дослідження мають бути спрямовані на explainable AI-підходи для прозорості виявлення та підвищення довіри користувачів.

Подяки та гранти

Результати наукових досліджень були виконані в рамках науково-дослідної роботи: «Методи та моделі забезпечення кібербезпеки інформаційних систем переробки інформації та функціональної безпеки програмно-технічних комплексів управління критичної інфраструктури» (№ 0122U200483, КСУБГ, м. Київ).

Джерела

1. Yin, K. S., & Levendovszky, J. (2017). Network behavioral features for detecting Remote Access Trojans. *Proceedings of the ACM*. ACM Digital Library. <https://dl.acm.org/doi/10.1145/3171592.3171597>
2. Dong, C., Lu, Z., Cui, Z., Liu, B., & Chen, K. (2020). MBTree: Detecting encryption RAT communication using malicious behavior tree. *arXiv*. <https://arxiv.org/abs/2006.10196>
3. Костюк, Ю., Складанний, П., Рзаєва, С., Мазур, Н., Черевик, В., & Аносов, А. (2025). Особливості реалізації мережевих атак через TCP/IP-протоколи. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29), 571–597. <https://doi.org/10.28925/2663-4023.2025.29.915>
4. Kaya, Y., et al. (2024). ML-based behavioral malware detection is far from a solved problem. *arXiv*. <https://doi.org/10.48550/arXiv.2405.06124>
5. Galli, A. (2024). Explainability in AI-based behavioral malware detection. *Computers & Security*. ScienceDirect. <https://www.sciencedirect.com/science/article/pii/S0167404824001433>
6. Складанний, П., Костюк, Ю., Рзаєва, С., Самойленко, Ю., & Савченко, Т. (2025). Розробка модульних нейронних мереж для виявлення різних класів мережевих атак. *Кібербезпека: освіта, наука, техніка*, 3(27), 534–548. <https://doi.org/10.28925/2663-4023.2025.27.772>
7. Alhazmi, A. H., et al. (2025). A robust and dynamic malware detection and classification approach. *BMC Bioinformatics*. PubMed Central (PMC). <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC12410719/>
8. Rashid, S. J. (2024). Detecting Remote Access Trojan (RAT) attacks based on different LAN analysis methods. *Engineering, Technology & Applied Science Research (ETASR)*. https://www.researchgate.net/publication/384824235_Detecting_Remote_Access_Trojan_RAT_Attacks_based_on_Different_LAN_Analysis_Methods
9. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. *Кібербезпека: освіта, наука, техніка*, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
10. de Oliveira, A. S., et al. (2023). Behavioral malware detection using deep graph convolutional neural networks. *TechRxiv*. <https://www.techrxiv.org/users/660121/articles/675292-behavioral-malware-detection-using-deep-graph-convolutional-neural-networks>
11. Rohini, S., et al. (2024). Malware behaviour analysis and impact quantification. *Computers & Security*. ScienceDirect. <https://www.sciencedirect.com/science/article/abs/pii/S0167404824000361>