

ПРОТИДІЯ КІБЕРАТАКАМ ЧЕРЕЗ АЛГОРИТМІЧНІ ЗАСОБИ

Гуненко В.¹, Світличний В.¹

¹*Харківський національний університет внутрішніх справ
vit.svet@ukr.net*

У статті досліджено сучасні виклики у сфері кібербезпеки та алгоритмічні засоби протидії кібератакам. Розглянуто основні типи атак – DoS, фішинг, malware, SQL-ін'єкції, APT – та методи їх виявлення за допомогою сигнатурного, евристичного й поведінкового аналізу. Особливу увагу приділено використанню машинного навчання для ідентифікації загроз, а також інтеграції алгоритмів у SIEM- та SOAR-системи. Проаналізовано ефективність автоматизованого реагування та проблеми, пов'язані з хибними спрацьовуваннями. Робота акцентує важливість створення комплексних, адаптивних систем кіберзахисту, здатних до самонавчання та оперативного реагування в умовах динамічного цифрового середовища.

Ключові слова: кібербезпека, кібератаки, алгоритмічні методи, машинне навчання, SIEM, SOAR, IDS/IPS, фішинг, DoS, APT, поведінковий аналіз

Вступ

У сучасному цифровому світі, де інформація стала одним із найцінніших активів, проблема захисту від кібератак набула надзвичайної актуальності. Кібератаки – це цілеспрямовані дії, спрямовані на порушення, перехоплення, зміну, знищення або викрадення інформації, що циркулює в інформаційно-комунікаційних системах. Вони можуть мати найрізноманітніший характер, починаючи від примітивних спроб фішингу до складних багаторівневих атак державного рівня, які використовують шкідливе програмне забезпечення, соціальний інжиніринг та вразливості нульового дня. Класифікація кібератак зазвичай здійснюється на основі їх мети, методу реалізації та рівня впливу. Найпоширенішими є атаки на

конфіденційність, цілісність і доступність даних, які можуть реалізовуватись через експлойти, віруси, хробаки, трояни, атакуючі скрипти, ботнети, зловмисні мобільні додатки та інші інструменти. Значна частина атак також базується на використанні людського фактору, коли зловмисник обманом змушує користувача надати доступ до системи або інформації.

Виклад основного матеріалу

У зв'язку з ускладненням технічного середовища, зростанням кількості пристроїв, підключених до мережі, розвитком Інтернету речей, хмарних технологій та віддаленої роботи, традиційні підходи до безпеки втрачають свою ефективність. Сучасні кібератаки здатні обходити статичні засоби захисту, залишаючись невиявленими в системах протягом тривалого часу, а їх наслідки можуть охоплювати не лише технологічну, а й економічну, соціальну та навіть політичну сфери. У таких умовах надзвичайно важливою стає розробка і впровадження ефективних алгоритмів захисту, які здатні виявляти, запобігати, локалізувати та усувати загрози в режимі реального часу. Алгоритмічний підхід до кібербезпеки дозволяє автоматизувати процеси захисту, підвищити швидкість реагування, зменшити навантаження на аналітиків безпеки та забезпечити масштабованість захисних механізмів відповідно до складності системи.

Серед найпоширеніших типів кібератак, що загрожують інформаційним системам у всьому світі, особливе місце займають DoS-атаки, фішинг, зловмисне програмне забезпечення (malware), SQL-ін'єкції та атаки типу APT (Advanced Persistent Threats). DoS-атака (Denial of Service) передбачає перевантаження ресурсів системи – вебсервера, мережі, додатку – запитами до такої міри, що вона припиняє обслуговування легітимних користувачів. У своїй розширеній формі, DDoS (Distributed Denial of Service), ця атака реалізується одночасно з тисяч комп'ютерів, зазвичай заражених вірусами, які об'єднані у ботнет. Такі атаки надзвичайно складно зупинити, оскільки трафік надходить із реальних, хоч і скомпрометованих пристроїв.

Фішинг є ще однією поширеною загрозою, яка спирається на обман користувача. Через електронну пошту, месенджери або фейкові вебсайти зловмисник імітує офіційне повідомлення, закликаючи жертву надати конфіденційну інформацію – логіни, паролі, банківські реквізити. Фішинг небезпечний тим, що може слугувати першим етапом складнішої атаки, наприклад, для впровадження шкідливого програмного забезпечення або отримання доступу до внутрішньої мережі. Malware – загальне поняття для всієї категорії шкідливих програм, які мають на меті вивести з ладу систему, зібрати інформацію або надати віддалений контроль зловмиснику. Сюди входять віруси, трояни, руткіти, шпигунські програми, програми-вимагачі (ransomware), які шифрують дані й вимагають викуп за їх розблокування. Одним із класичних векторів зараження є SQL-ін'єкція – вид атаки на вебдодатки, коли у поле введення користувача зловмисник вводить шкідливий SQL-запит, що дозволяє маніпулювати базою даних: переглядати, змінювати або видаляти інформацію, а іноді й отримувати повний доступ до серверної частини.

Окрему категорію загроз становлять АРТ-атаки – це цілеспрямовані, добре сплановані й довготривалі операції, що зазвичай ініціюються висококваліфікованими групами, зокрема державними. Такі атаки мають на меті проникнення в мережу, закріплення у ній і непомітне переміщення задля збору розвідданих або втручання в критичні процеси. АРТ-загрози складно виявити, оскільки вони маскуються під легітимні процеси, використовують унікальні шкідливі компоненти і поетапну тактику з ретельним уникненням виявлення.

Виявлення кібератак – це ключовий етап у системі захисту, який базується на застосуванні різних алгоритмів і методик. Сигнатурні методи виявлення є найбільш традиційними й полягають у порівнянні поточного трафіку або поведінки об'єктів з базою відомих ознак шкідливої активності – так званих сигнатур. Наприклад, якщо у файлі чи мережевому пакеті виявлено фрагмент коду, який відповідає відомому вірусу, система одразу повідомляє про

загрозу. Цей підхід ефективний для вже відомих атак, але є неефективним щодо нових, ще не задокументованих загроз – атак нульового дня. Саме тому на зміну сигнатурному аналізу приходять евристичні методи, які ґрунтуються на пошуку підозрілої, але не обов'язково вже ідентифікованої поведінки. Наприклад, якщо програма виконує нехарактерні для себе дії – змінює системні файли, створює копії себе у різних директоріях або підключається до зовнішніх серверів без очевидної причини – система може позначити її як потенційно небезпечну.

Ще ефективнішими є поведінкові алгоритми, які аналізують активність об'єктів у системі в динаміці, створюючи профілі «нормальної» поведінки користувачів, програм і пристроїв. Відхилення від цієї норми – наприклад, раптовий сплеск мережевої активності вночі, доступ до критичних файлів зі звичайного користувацького облікового запису або зміни в структурі прав доступу – розцінюються як сигнали можливої атаки. Поведінковий аналіз часто інтегрується зі штучним інтелектом і машинним навчанням, що дозволяє системам самостійно адаптуватися до нових типів загроз, аналізуючи великий обсяг даних у реальному часі. Такий підхід значно підвищує точність виявлення, зменшує кількість хибнопозитивних спрацьовувань і забезпечує проактивне реагування на атаки ще до того, як вони заподіють шкоду. Завдяки комбінації сигнатурних, евристичних і поведінкових алгоритмів сучасні системи кіберзахисту стають здатними до комплексного виявлення загроз у складному та швидкозмінному цифровому середовищі.

Використання машинного навчання для ідентифікації кіберзагроз стало одним із найперспективніших напрямів розвитку сучасної кібербезпеки, оскільки традиційні методи дедалі частіше виявляються недостатньо ефективними перед складними та швидко змінюваними атаками. На відміну від сигнатурних систем, що потребують заздалегідь відомого опису загрози, алгоритми машинного навчання здатні виявляти аномальні патерни в даних, які не підпадають під звичайну поведінку системи або користувача. Це дає змогу виявити нові або

модифіковані атаки, зокрема атаки нульового дня, які раніше не були зафіксовані в базах даних. Суть машинного навчання у сфері безпеки полягає в тому, що модель навчається на великій кількості даних про мережевий трафік, поведінку процесів, вхід/вихід користувачів, журнали системних подій тощо, а потім використовує отримані знання для класифікації майбутніх подій як нормальних або підозрілих.

Методи машинного навчання, що застосовуються в цьому контексті, можуть бути як підконтрольними (supervised learning), коли модель тренується на розмічених даних, так і безпідконтрольними (unsupervised learning), які корисні тоді, коли структура даних невідома, і йдеться про виявлення кластерів або аномалій без попереднього маркування. Наприклад, за допомогою алгоритмів класифікації, таких як дерева рішень, логістична регресія, метод опорних векторів або нейронні мережі, можна точно передбачати, чи є певна сесія з'єднання зловмисною. Безпідконтрольні методи, як-от алгоритм k-середніх або кластеризація DBSCAN, дозволяють ідентифікувати відхилення у великих масивах даних, які можуть бути ранніми ознаками несанкціонованої активності. Крім того, використовуються методи навчання з підкріпленням, які дозволяють системам автоматично вдосконалювати свої рішення шляхом аналізу наслідків своїх дій. Усе це робить машинне навчання надзвичайно корисним інструментом для побудови інтелектуальних систем безпеки, які не просто реагують на інциденти, а здатні передбачати та запобігати загрозам до їх реалізації.

Системи IDS/IPS можуть бути побудовані за різними архітектурами: мережева (Network-based IDS/IPS), яка аналізує мережевий трафік у режимі реального часу, або хостова (Host-based IDS/IPS), яка контролює активність окремих пристроїв або серверів, зокрема операційні системи, журнали подій, доступ до файлів та інші критичні параметри. Основу таких систем становлять алгоритми виявлення загроз: сигнатурні, що орієнтуються на відомі шаблони атак, і аномальні, які виявляють відхилення від нормальної поведінки. З розвитком технологій IDS/IPS

інтегруються з SIEM-системами (Security Information and Event Management), що дозволяє аналізувати загрози у ширшому контексті, враховуючи різні джерела даних, історію подій і поведінкові моделі користувачів.

Алгоритми реагування та автоматизація протидії кібератакам є важливим етапом у забезпеченні безперервного та ефективного захисту інформаційних систем.

У сучасних умовах, коли атаки відбуваються зі стрімкою швидкістю та часто охоплюють численні етапи й вектори впливу, людський фактор виявляється недостатньо швидким для ефективного реагування. Саме тому автоматизовані алгоритми, вбудовані у систему захисту, покликані оперативно виявляти загрозу, визначати її природу, локалізувати зону ураження та здійснювати необхідні дії для нейтралізації або мінімізації наслідків. Такі алгоритми можуть бути запрограмовані на виконання широкого спектру заходів: від блокування підозрілого мережевого трафіку до ізоляції вразливих сегментів системи, від скидання сесій користувача до оновлення політик доступу або виклику резервної інфраструктури. Автоматизоване реагування дозволяє реалізувати концепцію SOAR (Security Orchestration, Automation and Response), в межах якої система безпеки функціонує не лише як детектор, а як активний захисник, що приймає рішення на основі заздалегідь прописаних сценаріїв або динамічних оцінок ситуації. Важливо, що ці алгоритми можуть працювати в режимі 24/7 без зниження швидкості чи втоми, що дає змогу ефективно протистояти навіть масованим або складним атакам з мінімальним людським втручанням.

Однією з головних проблем, яка виникає при використанні автоматизованих алгоритмів виявлення і реагування на загрози, є висока ймовірність помилкових спрацювань – як false positive, так і false negative. Помилкове позитивне спрацювання (false positive) відбувається тоді, коли система ідентифікує легітимну активність як загрозу.

Це може спричинити небажане блокування критичних сервісів, призупинення доступу для користувачів або запуск непотрібних процедур реагування, що, в свою чергу, може призвести до порушень бізнес-процесів і навіть втрати довіри до системи безпеки.

З іншого боку, помилкове негативне спрацювання (false negative) є ще небезпечнішим, адже у такому випадку реальна атака залишається непоміченою, що дозволяє зловмиснику безперешкодно проникати в систему, розширювати присутність, викрадати дані або руйнувати інфраструктуру. Причинами таких помилок можуть бути як обмеженість алгоритму, недосконалість моделей машинного навчання, неякісно підготовлені дані для навчання, так і динамічна зміна поведінки зловмисників, які цілеспрямовано маскують свої дії під легітимні процеси.

Балансування між чутливістю системи до загроз і її толерантністю до помилкових сигналів є складним завданням, що потребує постійного аналізу, донавчання моделей, уточнення правил та участі фахівців із кібербезпеки, які мають коригувати поведінку системи відповідно до конкретного середовища.

Ще одним критично важливим елементом ефективного використання алгоритмів безпеки є їх інтеграція в комплексні SIEM-системи (Security Information and Event Management). SIEM-платформи об'єднують у собі функціональність збору, агрегації, зберігання, аналізу та візуалізації журналів подій з різних джерел – серверів, мережевих пристроїв, прикладного ПЗ, антивірусів, фаєрволів та інших компонентів IT-інфраструктури.

Інтеграція алгоритмів виявлення загроз у SIEM дозволяє отримати цілісну картину стану безпеки організації, розуміти, як взаємопов'язані різні події, виявляти складні атаки, що складаються з багатьох етапів, і оцінювати ризики в контексті бізнесу. Зокрема, SIEM надає змогу виявляти низькорівневі сигнали, які окремо не є загрозою, але у сукупності свідчать про початок складної атаки. Вбудовані алгоритми в SIEM можуть реалізовувати як

прості правила кореляції подій, так і складні сценарії поведінкового аналізу, засновані на штучному інтелекті.

Більше того, сучасні SIEM-системи все частіше інтегруються з SOAR-рішеннями, що дозволяє поєднувати виявлення загроз із автоматизованими сценаріями реагування: надсилання сповіщень, активація ізоляції, створення інцидентів у системах управління, ініціювання процедур розслідування. Завдяки цьому алгоритми стають не лише інструментом технічного аналізу, а й дієвим компонентом стратегічного управління безпекою, який дозволяє керівництву приймати обґрунтовані рішення в умовах постійного потоку подій та атак. Інтеграція алгоритмів у SIEM – це крок до створення єдиної аналітичної платформи, що забезпечує прозорість, гнучкість та ефективність усієї системи кіберзахисту.

Висновок

У сучасних умовах ефективний захист від кібератак неможливий без використання алгоритмічних засобів, здатних виявляти, локалізувати й нейтралізовувати загрози у реальному часі. Застосування машинного навчання, поведінкового аналізу та автоматизованого реагування значно підвищує рівень безпеки, дозволяючи системам адаптуватися до нових типів атак.

Проте важливим залишається завдання мінімізації хибних спрацьовувань і забезпечення гнучкості алгоритмів у різних інфраструктурних середовищах. Інтеграція таких алгоритмів у SIEM/SOAR-платформи забезпечує комплексний підхід до управління кіберзагрозами та стратегічне ухвалення рішень у сфері інформаційної безпеки.

Список використаних джерел

1. Algorithmic Detection Approaches in Cybersecurity: Exploring New Innovations and Techniques: <https://www.researchgate.net/publication/384367638>

2. Methods and Technologies to Detect Cyber Attacks: <https://www.eccu.edu/blog/methods-technologies-detect-cyber-attacks/>

3. Machine Learning Algorithms for Detecting and Preventing Cyber Threats: <https://www.oxjournal.org/machine-learning-algorithms-for-detecting-and-preventing-cyber-threats/>