

мандатным меткам, что упрощает анализ инцидентов. Настройка аудита централизованная, через графическую утилиту, с привязкой к пользователям и группам [2].

Анализ показал принципиально разные подходы к безопасности. Классические дистрибутивы (Debian, RHEL) используют адаптивную модель на базе дискреционного доступа: она гибка, но требует ручной настройки администратором, иначе система уязвима даже для root. Astra Linux SE реализует проактивную модель, где защита изначально встроена в архитектуру на основе формальной модели МРОСЛ.

Список использованных источников

1. Оценка стойкости механизма, реализующего мандатную сущностно-ролевую модель разграничения прав доступа в операционных системах семейства GNU Linux / А. И. Катасонов, С. И. Штеренберг, А. Ю. Цветков // Вестник Санкт-Петербургского государственного университета технологии и дизайна. Серия 1: Естественные и технические науки. – 2020. – № 2. – С. 50-56. – DOI 10.46418/2079-8199_2020_2_8. – EDN EUMWWI

2. Преимущества использования СЗИ в ОС Astra Linux Special Edition [Электронный ресурс] // habr.com. – URL: <https://habr.com/ru/companies/astralinux/articles/670060/>

Гилева Дарья Евгеньевна, бакалавриат СПбГУТ, daryagileva.2004@yandex.ru.

Гельфанд Артем Максимович, стар. препод. каф. ЗСС, gelfand.am@sut.ru

Ковцур Максим Михайлович, к.т.н., доцент, доцент каф. ИБКС, kovcur.mm@sut.ru

УДК 004.7.056; 004.7:004.8

СИСТЕМА ОБНАРУЖЕНИЯ АНОМАЛИЙ НА ОСНОВЕ АНАЛИЗА ЖУРНАЛА СОБЫТИЙ

Е.А. Атарская, А.М. Вульфин, А.Д. Кириллова
Уфимский университет науки и технологий, г. Уфа

Ключевые слова: интеллектуальный анализ данных, машинное обучение, анализ логов, межсетевые экраны.

Растущее число кибератак и появление новых уязвимостей требуют комплексных и адаптивных подходов к защите. Для обеспечения безопасности применяются различные системы и подходы, в том числе межсетевые экраны уровня веб-приложений (WAF).

Внедрение алгоритмов машинного обучения (МО) в WAF значительно расширяет возможности защиты веб-приложений, позволяя автоматизировать обнаружение аномалий и новых векторов атак, позволяет сократить количество ложных срабатываний и адаптироваться к меняющимся условиям работы, что повышает точность обнаружения угроз и снижает нагрузку на специалистов по информационной безопасности.

Обнаружение вредоносного трафика может быть реализовано с помощью анализа логов (журналов работы) WAF при анализе запросов прикладного уровня [1]. Использование данного подхода позволит обнаруживать аномальные HTTP- и HTTPS-запросы с высокой точностью при низком количестве ложных срабатываний.

Анализ журналов будет осуществляться с помощью моделей МО, способных выявлять аномалии, а их обучение будет проводиться на наборе данных на основе журналов WAF ModSecurity.

Целью является повышение эффективности выявления аномалий сетевого трафика на основе анализа журналов работы межсетевого экрана уровня веб-приложений. Под эффективностью понимается F1-мера и оперативность анализа.

Разработанная система обнаружения аномалий в логах WAF включает взаимосвязанные модули, которые обрабатывают набор данных с HTTP-запросами: приложение HTTP-requests sender (использует алгоритмы преобразования HTTP-запросов в curl-запросы и их отправки в WAF (ModSecurity)), модуль Log converter (конвертирует логи WAF в формат xlsx-таблицы), модуль Dataset creator (использует алгоритм выделения признаков из логов в формате xlsx-таблицы с целью формирования набора данных для МО), а также компонент Machine Learner, в рамках которого происходит разработка моделей, выявляющих аномалии в логах WAF.

Для формирования обучающей и тестовой выборок использован набор данных CSIC [2], содержащий как легитимные, так и аномальные HTTP-запросы типов GET, POST, PUT, что позволяет смоделировать реальные сценарии работы веб-приложений. Особое внимание было уделено этапу извлечения информативных признаков из логов, включая числовое кодирование метода запроса, анализ структуры пути, параметров и сообщений Mod_Security Apache. Такой подход обеспечивает комплексную оценку каждого запроса и позволяет учитывать как синтаксические, так и семантические особенности трафика.

Вычислительный эксперимент по обнаружению аномалий в журналах работы WAF с использованием подготовленного набора данных проводился на основе алгоритмов MO Isolation Forest, One Class Support Vector Machine (SVM), DBSCAN и Local Outlier Factor (LOF).

Перед обучением моделей выполнена нормализация набора данных: 20% вредоносных сессий из тестовой выборки перенесены в обучающую, а количество аномалий в тестовой выборке сбалансировано с количеством нормальных сессий для устранения дисбаланса классов.

Таблица 1 – Результаты машинного обучения моделей, выявляющих аномалии

	Isolation Forest	One Class SVM	DBSCAN	LOF
TP	1 321 (41%)	1 295 (40%)	1 378 (43%)	1 495 (46%)
TN	1 479 (46%)	1 491 (46%)	1 258 (39%)	1 320 (41%)
FP	132 (4%)	120 (4%)	353 (11%)	291 (9%)
FN	290 (9%)	316 (10%)	233 (7%)	116 (4%)
F1-score	0,86	0,85	0,83	0,86

Средняя оценка качества работы моделей МО после обучения на наборе, состоящем из 3222 наблюдений, в котором 50% наблюдений являются аномалиями, достигает в F1-score 0,85 (табл. 1).

Сравнение результатов обнаружения вредоносного трафика моделью Isolation Forest и WAF показывает, что доли ложноположительных предсказаний (FPR) составили 8% и 0%, а истинно положительных (TPR) – 82% и 29% соответственно.

В результате вычислительного эксперимента наилучший результат в классификации аномального трафика среди моделей машинного обучения показала модель на основе алгоритма Isolation Forest (0,86 F1), так как модель LOF (0,86 F1) характеризуется высокой долей ложноположительных срабатываний (FP). Кроме того, по сравнению с Mod_Security, Isolation Forest выявила аномалии на 53% эффективнее.

Результаты экспериментов подтверждают, что интеграция методов машинного обучения в процесс анализа журналов WAF способствует более эффективному обнаружению аномалий. Разработанную систему предлагается интегрировать с межсетевым экраном веб-приложений. Дальнейшее развитие этого направления связано с использованием более сложных моделей анализа и интеграцией с другими источниками событий, связанных с безопасностью. Это повысит устойчивость и адаптивность систем безопасности к современным киберугрозам.

Список использованных источников

1. Shaheed A., Kurdy M.H.D.B. Web application firewall using machine learning and features engineering / Security and Communication Networks, 2022, T. 2022, № 1, C. 5280158.
2. Applebaum S., Gaber T., Ahmed A. Signature-based and machine-learning-based web application firewalls: A short survey / Procedia Computer Science, 2021, T. 189, C. 359-367.

Атарская Елена Андреевна, аспирант каф. вычислительной техники и защиты информации, atarskaya.ea@ugatu.su.

Вульфин Алексей Михайлович, д.т.н., профессор каф. вычислительной техники и защиты информации, vulfin.alexey@gmail.com.

Кириллова Анастасия Дмитриевна, к.т.н., доцент каф. вычислительной техники и защиты информации, kirillova.ad@ugatu.su.

УДК 004.056, 004.75:004.8

СИСТЕМА ОБНАРУЖЕНИЯ СЕТЕВЫХ АТАК В ПРОМЫШЛЕННОЙ СЕТИ

А.М. Вульфин, А.Д. Кириллова, Э.Е. Садыкова
Уфимский университет науки и технологий, г. Уфа

Ключевые слова: сетевой трафик, алгоритмы объяснимого искусственного интеллекта, большие языковые модели

Разработка систем безопасности промышленных сетей и мониторинга