

Тогрул Сиявуш оглы Асадов

студент-магистрант Университета ADA и Университета им. Дж. Вашингтона (США)

E-mail: tasadov7572@ada.edu.az

ПРИМЕНЕНИЕ МЕТОДОВ МАШИННОГО ОБУЧЕНИЯ В ОБЕСПЕЧЕНИИ КИБЕРБЕЗОПАСНОСТИ: ВОЗМОЖНОСТИ И ВЫЗОВЫ

Резюме

Статья посвящена исследованию применения методов машинного обучения (ML) в области кибербезопасности, рассматривая как возможности, так и вызовы, с которыми сталкиваются организации при их использовании. В условиях интенсивной цифровизации и роста объёмов данных традиционные методы защиты информационных систем оказываются недостаточными для борьбы с современными киберугрозами, такими как целенаправленные атаки (APT), фишинг и вредоносное ПО. В статье выделяется ключевая роль ML в автоматическом выявлении, анализе и прогнозировании угроз, а также в обработке больших данных в реальном времени. Особое внимание уделяется вызовам, связанным с интерпретацией моделей, атакующими примерами, а также правовыми и этическими проблемами использования таких технологий.

Ключевые слова: цифровизация, машинное обучение, кибербезопасность, искусственный интеллект.

UOT: 336

JEL: E58, F31, F38, G18, E42

DOI: <https://doi.org/10.54414/ZPKH5165>

Введение

В условиях стремительной цифровизации всех сфер жизни вопросы обеспечения кибербезопасности приобретают критическую важность. Рост объёмов данных, распространение облачных технологий, развитие Интернета вещей (IoT) и удалённого доступа увеличивают как количество, так и сложность киберугроз. Традиционные методы защиты информационных систем становятся недостаточно эффективными перед лицом постоянно эволюционирующих атак, включая APT, фишинг, атаки на основе вредоносного ПО и уязвимости нулевого дня.

В этих условиях особую значимость приобретает внедрение методов машинного обучения (ML) как инструмента для автоматического выявления, анализа и прогнозирования киберугроз. Следует отметить, что машинное обучение представляет собой область искусственного интеллекта (ИИ), в которой компьютеры способны обучаться и совершенствовать свои действия, не опираясь на

заранее прописанные инструкции. В современном применении ИИ в бизнесе и технологиях машинное обучение часто выступает ключевым компонентом, что приводит к смешению и взаимозаменяемому использованию этих двух понятий [5]. Как показывают источники, изученные в данном направлении, эффективность модели машинного обучения определяется двумя основными факторами: качеством исходных данных (принцип «что вложишь — то и получишь») и корректным выбором алгоритма в зависимости от конкретной задачи. Подбор алгоритма зависит от природы данных и цели применения. В контексте кибербезопасности применяются различные методы машинного обучения:

- Деревья решений — для идентификации и классификации атак;
- Методы снижения размерности — для устранения шума и нерелевантной информации;

- Кластеризация К-средних — для выявления вредоносного ПО;
- k-ближайших соседей (kNN) — для аутентификации по распознаванию лиц;
- Линейная регрессия — для прогнозов в области сетевой безопасности;
- Метод опорных векторов (SVM) — для анализа и прогнозирования подозрительных IP- и порт-адресов [8].

Выбор алгоритма для моделей машинного обучения зависит от типа доступных данных и конкретной задачи. Актуальность выбранного исследования обусловлена необходимостью научного осмысления потенциала машинного обучения в сфере информационной безопасности, выявления его преимуществ и ограничений, а также анализа условий, необходимых для эффективного и безопасного внедрения этих технологий в реальные киберсреды. Проведение такого анализа способствует формированию научно обоснованных подходов к построению устойчивых и интеллектуальных систем киберзащиты

Обзор литературы

В ряде современных исследований отмечается возрастающая роль искусственного интеллекта и методов машинного обучения в обеспечении кибербезопасности. Наиболее распространённые направления применения включают:

- веб- и DNS-фильтрацию, где алгоритмы ИИ анализируют трафик и URL-адреса с целью выявления вредоносных сайтов и фишинговых атак;
- управление уязвимостями посредством автоматической приоритизации и оценки рисков на основе данных CVE и истории исправлений и др. [2]. В работе С.Бермана рассматриваются методы обучения на поведенческих признаках для классификации вредоносных программ с использованием моделей случайного леса и нейросетевых подходов [4]. Заслуживает особый интерес исследование [7], в котором был предложен ML-подход к выявлению фишинговых писем на основе анализа содержимого и URL-адресов, с использованием методов классификации SVM и логистической регрессии. В области

обнаружения сетевых вторжений активно применяются алгоритмы кластеризации и обнаружения аномалий. В работе Шона и других предложена модель глубокого обучения, совмещающая автоэнкодеры и кластеризацию для детектирования аномального поведения в сетевом трафике [9]. Методы ML находят применение и в системах управления уязвимостями. Например, в исследовании [3] рассматривается подход к прогнозированию степени критичности уязвимостей на основе анализа метаданных CVE и применения градиентного бустинга. Следует отметить, что важное место занимает обнаружение мошенничества. К примеру, в работе [6] представляют обзор алгоритмов выявления аномалий для задач финансового мошенничества, включая использование ансамблей деревьев решений и нейросетевых структур. Существуют исследования, активно анализирующие гибридные и объяснимые модели машинного обучения, в которых предлагается алгоритм LIME для интерпретации результатов классификации, что особенно актуально для систем, работающих в критически важных инфраструктурах [10].

Таким образом, проведенный обзор литературы подтверждает значительный потенциал применения методов машинного обучения в различных аспектах кибербезопасности, что делает необходимым дальнейшее исследование как технических возможностей, так и этико-правовых ограничений таких систем.

Проведенные исследования дают возможность классифицировать распространенные угрозы и соответствующие методы машинного обучения, которые могут быть использованы для их предотвращения или обнаружения в контексте обеспечения кибербезопасности.



Тип угрозы	Описание угрозы	Подходящий алгоритм/метод машинного обучения
Вредоносное ПО (Malware)	Программы, направленные на разрушение или контроль системы, утечку данных.	Алгоритмы классификации (например, Random Forest, SVM), нейронные сети для анализа поведения ПО, методы анализа бинарных файлов.
Атаки на отказ в обслуживании (DoS/DDoS)	Перегрузка системы запросами для ее недоступности.	Алгоритмы обнаружения аномалий (например, Isolation Forest), алгоритмы кластеризации для выявления необычного трафика.
SQL-инъекции	Атаки, которые используют уязвимости в SQL-запросах для несанкционированного доступа к базе данных.	Методика анализа входных данных с помощью классификаторов, обучение на аномальных запросах.
Фишинг	Мошенничество через электронную почту или веб-сайты для кражи личных данных.	Алгоритмы обработки текста (NLP), классификация URL, методы машинного обучения для анализа электронной почты.
Атаки типа Man-in-the-Middle (MITM)	Прерывание или перехват связи между двумя сторонами с целью изменения данных.	Алгоритмы для обнаружения аномалий в сетевом трафике, методы криптографической аутентификации.
Инсайдерские угрозы	Угрозы, исходящие от внутренних пользователей (например, сотрудники).	Алгоритмы анализа поведения (например, кластеризация, анализ аномалий), обнаружение аномалий в действиях пользователей.
Анализ вредоносного кода	Разбор и классификация файлов и кода на наличие вредоносных элементов.	Использование нейронных сетей для анализа файлов, машинное обучение для анализа бинарных данных и их поведения.
Фальсификация данных	Применение различных методов для изменения или подмены данных.	Алгоритмы для анализа целостности данных, методы машинного обучения для обнаружения изменений в данных (например, с использованием метода опорных векторов).
Угрозы в облачных сервисах	Атаки, направленные на облачные вычисления и утечку данных.	Алгоритмы мониторинга и анализа облачных данных, кластеризация для анализа аномального поведения в облаке.
Анализ уязвимостей программного обеспечения	Поиск и эксплуатация уязвимостей в программном обеспечении.	Алгоритмы статического анализа кода, обучение на примерах уязвимостей, алгоритмы для выявления уязвимостей в исходном коде.

Источник: составлено автором на основе изучения соответствующей научной литературы

Потенциальные возможности для Азербайджана.

В этом контексте особый интерес представляет исследование потенциала применения методов машинного обучения в области кибербезопасности в Азербайджане. С учётом развития цифровой экономики, расширения электронной коммерции и цифровизации государственного сектора, страна сталкивается с новыми вызовами в сфере защиты информационных систем. Это создаёт как потребность в локализованных решениях, так и

благоприятные условия для научных исследований и технологических разработок. Следует отметить, что утвержденная 19 марта 2025 года Президентом Азербайджанской Республики «Стратегия искусственного интеллекта Азербайджана на 2025–2028 годы» — является всесторонним документом, направленным на ускорение цифровой трансформации страны. Особое внимание уделяется вопросам кибербезопасности, что особенно актуально в условиях роста угроз в

цифровом пространстве. Стратегия предусматривает развитие национальной инфраструктуры, защиту персональных данных, регулирование использования ИИ и повышение устойчивости государственных систем. Комплексный подход, охватывающий правовое регулирование, инфраструктуру, подготовку кадров и поддержку бизнеса, создает прочную основу для формирования безопасной и эффективной ИИ-экосистемы. Успешная реализация стратегии укрепит позиции Азербайджана как технологически развитого и киберустойчивого государства [1].

С учётом изложенного, можно выделить ряд приоритетных направлений развития и исследований, способствующих эффективному внедрению технологий машинного обучения в обеспечение кибербезопасности в Азербайджане:

1. Создание национальных дата-сетов для обучения моделей

Большинство современных моделей машинного обучения требуют обширных обучающих выборок. Однако данные, полученные из других стран, не всегда учитывают локальную специфику:

- Исследования могут быть направлены на формирование открытых и анонимизированных дата-сетов, включающих данные о киберугрозах, инцидентах и вредоносных активностях, характерных для Азербайджана.

- Создание специализированных наборов данных для секторов с высокой киберугрозой: финансовый сектор, инфраструктура, образование.

2. Разработка мультиязычных моделей (в т.ч. на азербайджанском языке)

Многие фишинговые и вредоносные сообщения ориентированы на локальные языки. Важно разрабатывать и обучать модели:

- для фильтрации фишинга и спама на азербайджанском языке;
- для анализа подозрительного контента в социальных сетях и мессенджерах, которые активно используются в стране.

3. Интеграция ML в национальные CERT и SOC

- Разработка ML-моделей для автоматического мониторинга, корреляции и прогнозирования угроз в рамках национального

центра реагирования на киберинциденты (CERT).

- Создание адаптивных SIEM-систем с ML-модулями для предприятий и государственных структур.

4. Обучение специалистов

- Исследование методов использования ML в автоматизированном обучении пользователей: модели, которые отслеживают действия и предлагают меры предосторожности.

- Анализ поведения пользователей для профилактики инцидентов из-за "человеческого фактора" — одной из главных причин взломов.

5. Сотрудничество университетов и индустрии

- Развитие исследовательских лабораторий по ML в кибербезопасности в вузах.

- Разработка совместных проектов с национальными телеком-компаниями, банками и госсекторами, нацеленными на локальные решения.

6. Разработка нормативно-правовой базы для использования ИИ в кибербезопасности

- Исследования в области регулирования и этики использования ML в кибербезопасности, включая вопросы ответственности при автоматическом принятии решений.

Заключение

Как показывают исследования, современные угрозы, такие как целенаправленные атаки, фишинг и вредоносное ПО, требуют использования передовых технологий, способных эффективно справляться с объемом данных и сложностью угроз. В этом контексте машинное обучение представляется мощным инструментом для автоматического выявления аномалий, анализа поведения и прогноза новых видов атак, что значительно улучшает защиту информационных систем. Однако, несмотря на обещающие перспективы, использование ML в кибербезопасности сталкивается с рядом вызовов, таких как сложности в интерпретации моделей, риски атакующих примеров, а также этические и правовые проблемы. По этой причине, на наш взгляд, необходимо продолжить исследования, направленные на усовершенствование методов машинного обучения и их адаптацию к реальным условиям. Важным шагом

в этом направлении является создание локализованных дата-сетов, разработка мультиязычных моделей, а также интеграция ML в национальные системы реагирования на киберинциденты. Реализация предложенных инициатив в рамках национальных стратегий цифровизации и искусственного интеллекта, таких как «Стратегия искусственного интеллекта Азербайджана на 2025–2028 годы», создаст прочную основу для повышения уровня киберзащиты и устойчивости страны в цифровом пространстве.

Таким образом, применение методов машинного обучения в кибербезопасности имеет большой потенциал для развития и защиты цифровых инфраструктур, но требует комплексного подхода, включающего технические, этические и правовые аспекты, а также сотрудничество между государственными структурами, академией и промышленностью для эффективного внедрения этих технологий.

Список использованной литературы:

1. “Azərbaycan Respublikasının 2025–2028-ci illər üçün süni intellekt Strategiyası”nın təsdiq edilməsi haqqında Azərbaycan Respublikası Prezidentinin Sərəncamı / Bakı şəhəri, 19 mart 2025-ci il № 530. URL: <https://e-qanun.az/framework/59218>
2. Козлова Н.Ш., Довгаль В.А. Анализ применения искусственного интеллекта и машинного обучения в кибербезопасности URL: <file:///C:/Users/HP/Downloads/analiz-primeneniya-iskusstvennogo-intellekta-i-mashinnogo-obucheniya-v-kiberbezopasnosti.pdf>
3. Amer E., Zelin I. A dynamic Windows malware detection and prediction method based

on contextual understanding of API call sequence / Computers & Security

Volume 92, May 2020. URL: <https://doi.org/10.1016/j.cose.2020.101760>

4. Berman, S. et al. A survey of deep learning methods for cyber security. Information, 2019, 10 (4), 122 p. <https://doi.org/10.3390/info10040122>

5. Brown S. Machine learning, explained. 2021. URK: <https://mitsloan.mit.edu/ideas-made-to-matter/machine-learning-explained>

6. Carcillo, F., et al. Combining unsupervised and supervised learning in credit card fraud detection / Information Sciences, 2019. 557, pp.317–331. URL: <https://www.sciencedirect.com/science/article/abs/pii/S0020025519304451>

7. Jain, A. K., Gupta, B. B. Phishing detection: Analysis of visual similarity based approaches. Security and Privacy. 2027. URL: https://www.researchgate.net/publication/312205924_Phishing_Detection_Analysis_of_Visual_Similarity_Based_Approaches

8. Machine learning (ML) in cybersecurity. 2025. URL: <https://www.sailpoint.com/identity-library/how-ai-and-machine-learning-are-improving-cybersecurity>

9. Shone, N. et al. A deep learning approach to network intrusion detection. IEEE Transactions on Emerging Topics in Computational Intelligence, 2018, 2(1), 41–50. <https://doi.org/10.1109/TETCI.2017.2772792>

10. Ribeiro, M. T. et al. "Why should I trust you?": Explaining the predictions of any classifier. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 2016, pp. 1135–1144. <https://doi.org/10.1145/2939672.2939778>

Toğrul Siyavuş oğlu ƏSƏDOV

ADA Universiteti və Corc Vaşinqton Universitetinin (ABŞ) magistrantı

E-mail: tasadov7572@ada.edu.az**KİBERTƏHLÜKƏSİZLİYİN TƏMİNİNDƏ MAŞIN ÖYRƏNMƏSİ METODLARININ
TƏTBİQİ: İMKANLAR VƏ ÇAĞIRIŞLAR****Xülasə**

Məqalə kibertəhlükəsizlik sahəsində maşın öyrənməsi (MÖ) metodlarının tətbiqini araşdırır və onların istifadəsi zamanı təşkilatların qarşılaşdıqları imkanları və çağırışları təhlil edir. Rəqəmsallaşmanın sürətlənməsi və məlumat həcmələrinin artması şəraitində informasiya sistemlərinin ənənəvi mühafizə üsulları müasir kibertəhdidlərə — məqsədli hücumlar (APT-lər), fişinq və zərərli proqram təminatı kimi hallara qarşı yetərli olmur. Məqalədə maşın öyrənməsinin kibertəhlükələrin avtomatik aşkarlanması, təhlili və proqnozlaşdırılması, eləcə də böyük həcmli məlumatların real vaxt rejimində işlənməsi baxımından mühüm rol vurğulanır. Eyni zamanda, modellərin şərh olunması, təcavüzkar nümunələr və bu texnologiyaların tətbiqi ilə bağlı hüquqi və etik məsələlər kimi mühüm çağırışlara da xüsusi diqqət yetirilir.

Açar sözlər: rəqəmsallaşma, maşın öyrənməsi, kibertəhlükəsizlik, süni intellekt.

Toghrul Siyavush ASADOV

Master's student at ADA University and The George Washington University (USA)

E-mail: tasadov7572@ada.edu.az**APPLICATION OF MACHINE LEARNING METHODS IN ENSURING
CYBERSECURITY: OPPORTUNITIES AND CHALLENGES****Abstract**

This article explores the application of machine learning (ML) methods in the field of cybersecurity, addressing both the opportunities and challenges faced by organizations in their implementation. In the context of rapid digitalization and the exponential growth of data, traditional approaches to information system protection are increasingly insufficient to counter modern cyber threats such as advanced persistent threats (APT), phishing, and malware. The article highlights the critical role of ML in the automated detection, analysis, and prediction of threats, as well as its capacity to process large-scale data in real time. Particular attention is given to the challenges associated with model interpretability, adversarial examples, and the legal and ethical issues surrounding the use of these technologies.

Keywords: digitalization, machine learning, cybersecurity, artificial intelligence.