

МЕТОДЫ КЛАССИФИКАЦИИ СЕТЕВОГО ТРАФИКА С ИСПОЛЬЗОВАНИЕМ МАШИННОГО ОБУЧЕНИЯ

Т. Т. Сафиуллин

tuleubay.safiullin@mail.ru;

Научный руководитель – М. С. Абрамович, кандидат физико-математических наук

В статье анализируются результаты применения методов машинного обучения для классификации сетевого трафика. Приводятся конкретные сценарии и проблемы использования. Анализируются статьи, посвященные машинному обучению в сфере анализа трафика. Приводятся результаты собственных экспериментов. Приведены выводы об особенностях применения методов машинного обучения и перспективах их дальнейшего использования.

Ключевые слова: классификация сетевого трафика; машинное обучение; компьютерные сети.

ВВЕДЕНИЕ

Существует много методик, используемых для анализа сетевого трафика, таких как самоподобие и TES, которые основаны на анализе системы связи и обнаружении атак [1]. В настоящее время многие компании используют машинное обучение для выявления различных типов вредоносного поведения пользователей, включая массовую регистрацию поддельных аккаунтов, мошеннические транзакции и кражу личных данных, где безопасность может быть достигнута с помощью предложенных приложений и методов машинного обучения [2].

Основная концепция машинного обучения состоит в том, что машины способны постоянно учиться и работать с огромными массивами данных с помощью классификаторов и алгоритмов. Классификаторы считаются основой машинного обучения, и задача их состоит в том, чтобы классифицировать наблюдения. Алгоритмы же, в свою очередь, строят модели поведения и используют эти модели в качестве основы для будущих прогнозов, используя новые входные данные. Сила инструментов машинного обучения заключается в обнаружении и анализе сетевых атак без необходимости их точного описания, как было сказано ранее. Машинное обучение может помочь в решении наиболее распространенных задач, включая регрессию, прогнозирование и классификацию в эпоху чрезвычайно большого количества данных и нехватки кадров в области кибербезопасности.

Методы машинного обучения применяются во многих аспектах эксплуатации и управления сетями, где производительность системы может быть оптимизирована, а ресурсы могут быть лучше использованы.

Более того, кластеризация и классификация извлекают шаблоны из пакетов данных, которые могут быть использованы во многих приложениях, таких как анализ безопасности и профилирование пользователей [3]. Кроме того, существует множество приложений для анализа трафика на основе алгоритмов машинного обучения, например выявление аномалий с помощью признаков, описывающих поведение пользователя [4]. Сочетание алгоритмов машинного обучения и анализа трафика является актуальной темой в связи с эффективностью инструментов машинного обучения, которая заключается в обнаружении и анализе сетевых атак без необходимости их точного описания [3].

В данной статье рассматриваются различные методы машинного обучения для классификации сетевого трафика и проводится сравнительный анализ их точности.

РЕЗУЛЬТАТЫ И ИХ ОБСУЖДЕНИЕ

В [5] поставленной целью являлось обнаружение как известных, так и ранее неизвестных угроз безопасности. Была разработана система классификаторов, с помощью которой было обнаружено 2090 новых образцов вредоносного программного обеспечения. Это достигается частично за счет набора признаков (полезная нагрузка пакета, параметры запроса URL, время запроса, адреса веб-страниц), построенной для каждой группы HTTP-потоков пользователя, посещающего определенный хост, и частично за счет матрицы самоподобия признаков, вычисленной для каждой группы.

В [6] основной задачей являлась классификация трафика. Для этого был использован усовершенствованный классификатор деревьев решений – C5.0, который показал способность различать 7 различных приложений в тестовом наборе из 1 622 710 неизвестных случаев со средней точностью 92.5%. В качестве признаков были выбраны количество входящих/исходящих/общих байтов полезной нагрузки, соотношение входящих и исходящих пакетов данных, среднее значение, минимум, максимум, первый квартиль, медиана, третий квартиль, стандартное отклонение входящего/исходящего/общего размера полезной нагрузки и ряд других.

В работе был построен классификатор, использующий бустинг над решающими деревьями. Набор данных был взят из базы данных Канадского института кибербезопасности [8], состоящий из 3 119 345 записей, 83 признаков, которые содержат 15 меток классов (1 легитимная метка + 14 меток атак). Для этого набора данных было имитировано поведение 25 пользователей на основе протоколов HTTP, HTTPS, FTP, SSH и протокола электронной почты. Также набор данных содержит

информацию об атаках (XSS, FTP/SSH-Patator, PortScan, Brute Force, DoS, DDoS и др.) в виде данных о трафике за пять дней.

В таблице 1 представлены результаты классификаторов [4-8].

Таблица 1

Оценки точности классификаторов

Классификаторы:	Точность классификации:
Деревья решений C5.0	92.5%
Бустинг деревьев решений C5.0	99.3-99.9%
Случайный лес	96%
Метод опорных векторов	97.4%
к-ближайших соседей	97.1%
Наивный байесовский классификатор	60.3%

Для отбора информативных признаков использовалась двухэтапная процедура:

- на первом этапе с помощью встроенного механизма метода `sklearn.ensemble.RandomForestClassifier` (атрибут `feature_importances`), реализующего энтропийный подход к оценке важности признаков, отобрано 20 признаков;
- на втором этапе исключено 10 признаков, которые очень сильно коррелируют с другими.

На этапе выбора модели классификации была произведена оценка качества наиболее распространенных моделей машинного обучения.

Для каждого из рассмотренных классификаторов определены гиперпараметры, обеспечивающие наибольшую точность классификации.

Качество ответов классификаторов (моделей) сравнивалось с использованием следующих метрик:

- доля правильных ответов (accuracy);
- точность (precision, насколько можно доверять классификатору);
- полнота (recall, как много объектов класса «есть атака» определяет классификатор);
- F1-мера (гармоническое среднее между точностью и полнотой).

Проведено тестирование случайного леса на данных, когда вредоносные атаки также моделировались. В этом случае оценка полноты (recall) составила 0.946 и F1- меры – 0.97.

В таблице 2 приведены полученные значения метрик качества, усредненные по результатам 5 итераций кросс-валидаций лучших моделей машинного обучения.

Таблица 2

Оценки точности классификации лучших моделей машинного обучения

Модель	Accuracy	Precision	Recall	F1-мера
к-ближайших соседей	0.971	0.942	0.961	0.969
Деревья решений	0.975	0.973	0.946	0.969
Случайный лес	0.971	0.978	0.943	0.970
Адаптивный бустинг	0.978	0.962	0.965	0.973
Логистическая регрессия	0.955	0.939	0.914	0.963

Как следует из таблицы 2, все лучшие модели машинного обучения показали высокие оценки точности классификации.

При проведении экспериментов были выделены следующие проблемы использования:

- Классификация сетевого трафика требует периодического обновления и обучения модели, кроме того, необходимо определить задачу модели. Этот процесс огромен, если смотреть с точки зрения вредоносного трафика, изменения типов и особенностей современных кибератак. Это требует определенного времени для оценки данных и определенного объема данных.
- Существует множество критериев, которые необходимо учитывать при определении эффективности методов обнаружения типа атаки, и крайне сложно объективно оценить работу каждого метода [9].

ЗАКЛЮЧЕНИЕ

В статье рассматривается использование методов машинного обучения в анализе трафика. Представлен краткий обзор и сравнение некоторых существующих подходов машинного обучения, используемых в анализе трафика. Построен собственный классификатор и проведены эксперименты. Несмотря на большую роль машинного обучения, в статье показано, что оно все еще имеет некоторые проблемы в виде временной сложности.

БИБЛИОГРАФИЧЕСКИЕ ССЫЛКИ

1. Rueda, A. A survey of traffic characterization techniques in telecommunication networks. In Proceedings of 1996 Canadian Conference on Electrical and Computer Engineering. 1996. IEEE.
2. Chakraborty, A., J.S. Banerjee, and A. Chattopadhyay, Non-uniform quantized data fusion rule for data rate saving and reducing control channel overhead for cooperative spectrum sensing in cognitive radio networks. Wireless Personal Communications, 2019. 104(2): p. 837-851.

3. *Cheng, Y., et al., Bridging machine learning and computer network research: a survey. CCF Transactions on Networking, 2019. 1(1- 4): p. 1-15.*
4. *Mukkamala, S., G. Janoski, and A. Sung. Intrusion detection: support vector machines and neural networks. in proceedings of the IEEE International Joint Conference on Neural Networks (ANNIE), St. Louis, MO. 2002.*
5. *Bartos, K., M. Sofka, and V. Franc. Optimized invariant representation of network traffic for detecting unseen malware variants. In 25th {USENIX} Security Symposium ({USENIX} Security 16). 2016.*
6. *Bujlow, T., T. Riaz, and J.M. Pedersen. A method for classification of network traffic based on C5.0 Machine Learning Algorithm. in 2012 international conference on computing, networking and communications (ICNC). 2012.*
7. *Laskov, P., et al. Learning intrusion detection: supervised or unsupervised? In International Conference on Image Analysis and Processing. 2005.*
8. *Panigrahi R., Borah S. A detailed analysis of CICIDS2017 dataset for designing Intrusion Detection Systems. International Journal of Engineering & Technology, vol 7, no 3.24, 2018, p. 479-482.*
9. *Buczak, A.L. and E. Guven, A survey of data mining and machine learning methods for cybersecurity intrusion detection. IEEE Communications Surveys & Tutorials, 2015. 18(2): p. 1153-1176.*