

УДК 004.056

DOI: 10.26467/2079-0619-2026-29-3-59-70

Защита корпоративных информационных систем авиапредприятий от атак нулевого дня: подход к обнаружению вторжений на основе глубокого неконтролируемого обучения

Н.О. Машошин¹, А.А. Кулешов²

¹ООО «СофтТелематика», г. Москва, Россия

²АО «Научно-производственное предприятие «Аэросила», г. Ступино, Россия

Аннотация: В работе рассматривается актуальная задача обнаружения атак нулевого дня в корпоративных информационных системах авиапредприятий, относящихся к объектам критической информационной инфраструктуры. Показано, что традиционные сигнатурные средства защиты, включая классические системы обнаружения вторжений (IDS), обладают фундаментальной ограниченной эффективностью в условиях ранее неизвестных и целенаправленных кибервоздействий, что подтверждается реальными инцидентами в авиационном секторе. Поиск эффективных средств детектирования атак нулевого дня, а также комплексных целенаправленных воздействий (АРТ) остается актуальным, поскольку их скрытность и уникальность сводят к минимуму эффективность сигнатурных методов обнаружения. В данной работе предложена нейросетевая модель обнаружения вторжений, ориентированная на применение в центрах обеспечения безопасности (SOC) авиапредприятий и основанная на методах глубокого неконтролируемого обучения. Модель реализована в виде глубокого автоэнкодера, обучаемого исключительно на легитимном сетевом трафике, что позволяет формировать устойчивое представление нормального поведения системы и выявлять статистические отклонения без предварительного знания сигнатур атак. Экспериментальная валидация проведена на наборе данных CICIDS2018 с использованием метрик F1-score, ROC-AUC, precision и recall. Предложенный подход продемонстрировал F1-меру 0,81 и ROC-AUC 0,844, превзойдя показатели классических алгоритмов неконтролируемого обнаружения аномалий (Isolation Forest и One-Class SVM). Полученные результаты подтверждают применимость разработанной модели в качестве проактивного аналитического компонента гибридных систем обнаружения вторжений и ее потенциал для повышения киберустойчивости информационных систем авиационного сектора. Модель может быть интегрирована в существующие SOC-платформы авиапредприятий для дополнения сигнатурного анализа поведенческим контекстом.

Ключевые слова: кибербезопасность, авиапредприятие, атака нулевого дня, обучение без учителя, глубокое обучение, автоэнкодер, обнаружение аномалий, SOC, ICAO.

Для цитирования: Машошин Н.О., Кулешов А.А. Защита корпоративных информационных систем авиапредприятий от атак нулевого дня: подход к обнаружению вторжений на основе глубокого неконтролируемого обучения // Научный вестник МГТУ ГА. 2026. Т. 29, № 3. С. 59–70. DOI: 10.26467/2079-0619-2026-29-3-59-70

Protection of corporate information systems of aviation enterprises from zero-day attacks: an intrusion detection approach based on deep unsupervised learning

N.O. Mashoshin¹, A.A. Kuleshov²

¹SoftTelematika LLC, Moscow, Russia

²Scientific-Production Enterprise "Aerosila", Stupino, Russia

Abstract: The paper considers the urgent task of detecting zero-day attacks in corporate information systems of aviation enterprises classified as critical information infrastructure (CII). It is shown that traditional signature-based protection tools, including classical intrusion detection systems (IDS), have fundamentally limited effectiveness in conditions of previously unknown and targeted cyberattacks, which is confirmed by real incidents in the aviation sector. The search for effective methods of detecting zero-day

attacks, as well as advanced persistent threats (APT) remains an urgent task, since their secrecy and uniqueness minimize the effectiveness of signature-based detection methods. In this paper a neural-network intrusion detection model is proposed, focused on application in aviation security operation centers (SOC) of aviation enterprises and based on deep unsupervised learning methods. The model is implemented in the form of a deep autoencoder trained exclusively on legitimate network traffic, which makes it possible to form a stable representation of normal behavior of the system and identify statistical deviations without prior knowledge of attack signatures. Experimental validation was performed on the CICIDS2018 dataset using the metrics F1-score, ROC-AUC, precision, and recall. The proposed approach demonstrated the F1-measure of 0.81 and a ROC-AUC of 0.844, exceeding the performance of classical unsupervised algorithms for uncontrolled anomaly detection (Isolation Forest and One-Class SVM). The results obtained confirm the applicability of the developed model as a proactive analytical component of hybrid intrusion detection systems and its potential to increase the cyber resilience of aviation-sector information systems. The model can be integrated into existing SOC platforms of aviation enterprises to complement signature-based analysis with a behavioral context.

Keywords: cybersecurity, aviation enterprise, zero-day attack, unsupervised learning, deep learning, autoencoder, anomaly detection, SOC, ICAO.

For citation: Mashoshin, N.O., Kuleshov, A.A. (2026). Protection of corporate information systems of aviation enterprises from zero-day attacks: an intrusion detection approach based on deep unsupervised learning. Civil Aviation High Technologies, vol. 29, no. 3, pp. 59–70. DOI: 10.26467/2079-0619-2026-29-3-59-70

Введение

Цифровая трансформация авиационной отрасли привела к формированию сложных, высокоинтегрированных корпоративных информационных систем авиапредприятий, обеспечивающих как управление полетной деятельностью, бронирование и продажу билетов, техническое обслуживание воздушных судов, так и взаимодействие с пассажирами и внешними контрагентами [1]. Рост степени автоматизации и взаимосвязанности данных систем сопровождается существенным расширением поверхности атаки, что делает авиапредприятия привлекательной целью для злоумышленников.

Реальный инцидент, произошедший в июле 2025 года с ПАО «Аэрофлот», наглядно продемонстрировал уязвимость корпоративных ИТ-инфраструктур авиационных компаний: в результате целенаправленной хакерской атаки была нарушена работоспособность значительного числа серверов и рабочих станций, что привело к дестабилизации операционной деятельности и отмене рейсов¹. Подобные инциденты подтверждают, что последствия кибератак в авиационной отрасли выходят за рамки информационной

безопасности и напрямую затрагивают вопросы устойчивости бизнеса и безопасности полетов.

Традиционные системы обнаружения вторжений (IDS), основанные на сигнатурном анализе, демонстрируют ограниченную эффективность в условиях атак нулевого дня, а также при сложных целенаправленных атаках (APT) [2–4]. Указанные типы угроз характеризуются высокой степенью адаптивности, скрытности и использованием ранее неизвестных векторов воздействия, что делает невозможным их своевременное выявление средствами традиционного сигнатурного анализа. Дополнительным ограничением является специфика авиационной отрасли, связанная со строгими регламентами, сертификационными требованиями и необходимостью обеспечения непрерывности бизнес-процессов, что затрудняет частое обновление программных компонентов и сигнатурных баз данных средств защиты.

В ответ на такие вызовы Международная организация гражданской авиации (ICAO) в стратегических документах по кибербезопасности подчеркивает необходимость перехода от реактивных методов защиты к проактивным аналитическим подходам, ориентированным на выявление аномалий и неизвестных угроз на ранних стадиях².

¹ Сбой на взлете. «Аэрофлот» подвергается беспрецедентной хакерской атаке [Электронный ресурс] // Коммерсантъ. 2025. URL: <https://www.kommersant.ru/doc/7923869> (дата обращения: 25.11.2025).

² Cybersecurity Action Plan. 2nd ed. // ICAO, 2022. 44 p.

Расширение аналитических возможностей систем защиты связано с применением методов машинного обучения. Однако значительная часть существующих решений в области IDS опирается на обучение с учителем и предполагает наличие размеченных выборок атак. Подобные модели ориентированы преимущественно на классификацию заранее известных сценариев вторжений и, несмотря на высокие показатели точности в лабораторных условиях, демонстрируют ограниченную способность к выявлению принципиально новых типов аномалий. В связи с этим особую актуальность приобретают методы машинного обучения без учителя, позволяющие формировать модель нормального функционирования информационной системы и обнаруживать статистически значимые отклонения без использования размеченных данных [5, 6]. Развитие данного направления отражено в работах по одноклассовой классификации на основе нейронных сетей [7], а также в исследованиях, направленных на повышение устойчивости моделей без учителя к шуму и выбросам в данных [8]. В обзоре современных подходов к глубокому обнаружению аномалий [9] подчеркивается преимущество нелинейных представлений при анализе сложных многомерных распределений.

В области сетевой кибербезопасности предложены различные архитектуры IDS на основе глубокого обучения, включая сверточные и рекуррентные нейронные сети [10, 11]. Тем не менее проблема обнаружения ранее неизвестных сценариев вторжений остается открытой, что подтверждается результатами сравнительного анализа методов неконтролируемого выявления аномалий в многомерных данных [12].

Среди методов неконтролируемого обучения перспективным направлением является применение глубоких автоэнкодеров, способных выявлять аномальное поведение на основе анализа ошибок реконструкции входных данных [8, 13]. Для корпоративных сетей авиапредприятий такой подход представляет практический интерес, поскольку в условиях появления атак нулевого дня и постоянно

модифицируемых сценариев вторжений формирование полной базы известных атак является концептуально недостижимой задачей.

Основная научная задача исследования заключается в разработке и валидации методологии обнаружения атак нулевого дня в корпоративных информационных системах авиапредприятий на основе глубокого неконтролируемого обучения.

Целью работы является экспериментальная оценка способности модели глубокого автоэнкодера к обнаружению аномалий, то есть к выявлению признаков несанкционированного воздействия в сетевом трафике корпоративных сетей авиапредприятий. В качестве критериев эффективности используются количественные метрики качества классификации, включая F1-меру, точность и полноту.

Научная новизна и вклад работы заключаются в следующем.

1. Разработана методология обнаружения атак нулевого дня в корпоративных сетях авиапредприятий на основе глубокого автоэнкодера, реализующая парадигму обучения без учителя, ориентированная на обучение исключительно на данных легитимного трафика.

2. Предложен подход к интеграции системы обнаружения аномалий в существующую инфраструктуру центров обеспечения кибербезопасности (SOC) авиапредприятий, сочетающий преимущества сигнатурных и неконтролируемых методов анализа.

3. Сформирован метод оценки эффективности системы обнаружения, согласованный с принципами проактивной кибербезопасности, рекомендованными ICAO, и требованиями операционной устойчивости авиационных организаций.

Предлагаемый подход имеет практическое значение для повышения уровня защищенности корпоративных информационных систем авиапредприятий и может быть использован в качестве компонента интеллектуальных систем мониторинга безопасности авиационного сектора.

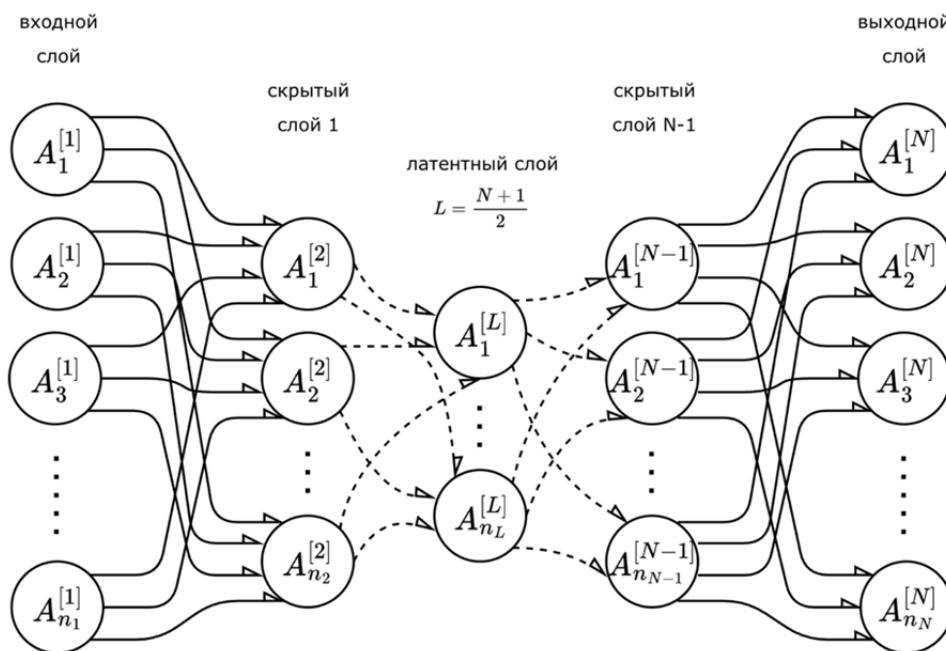


Рис. 1. Архитектура глубокого автоэнкодера: нижний индекс – порядковый номер нейрона, верхний индекс – порядковый номер слоя

Fig. 1. Deep autoencoder architecture: the lower index is the sequence number of the neuron, the upper index is the sequence number of the layer

Методология

Архитектура математической модели автоэнкодера

В основе предлагаемого подхода лежит глубокий автоэнкодер – нейронная сеть, целью обучения которой является реконструкция входных данных на выходе. Архитектура автоэнкодера (АЕ), представленная на рис. 1, включает две симметричные функциональные части: энкодер и декодер.

Энкодер осуществляет отображение входного вектора признаков $x \in X \in R^{n_1}$ через серию скрытых слоев нейронной сети в низкоразмерное латентное представление:

$$z = f(x) \in R^{n_L}, n_L \ll n_1, \quad (1)$$

где $f(\cdot)$ – нелинейное отображение, реализуемое последовательностью скрытых слоев нейронной сети, n_L – размерность латентного

представления, $L = \frac{N+1}{2}$, N – общее число слоев нейронной сети.

Декодер выполняет функцию обратного преобразования и формирует восстановленный вектор $\hat{x} = g(z)$, где $g(\cdot)$ – отображение, аппроксимирующее обратную функцию энкодера.

Обучение модели заключается в минимизации функции потерь реконструкции, в качестве которой используется среднеквадратичная ошибка (MSE). Для отдельного вектора x ошибка определяется как

$$L(x, \hat{x}) = \frac{1}{n} \sum_{i=1}^n (x_i - \hat{x}_i)^2, \quad (2)$$

где n – размерность вектора.

Для мини-батча из m образцов $X = \{x^{(1)}, \dots, x^{(m)}\}$ используется усреднение

$$L_b(X_b, \hat{X}_b) = \frac{1}{m} \sum_{i=1}^m L(x_i, \hat{x}_i). \quad (3)$$

Ключевым свойством, используемым для обнаружения аномалий, является то, что мо-

дель, обученная исключительно на легитимном трафике, минимизирует ошибку (2) для данных, выбранных из того же распределения. Таким образом, происходит обучение эффективного реконструирования только тех данных, которые соответствуют нормальному поведению сети. Аномальные образцы (например, векторы атак) статистически отличаются от легитимного распределения, что приводит к значительному росту ошибки реконструкции (2). Этот рост служит детектирующим признаком: если $L(x, \hat{x}) > \tau$, где τ – оптимально выбранный порог, то образец классифицируется как аномалия. Таким образом, архитектура автоэнкодера позволяет реализовать детектор, не требующий данных об атаках на этапе обучения, что критически важно для предметной области с отсутствием размеченных векторов атак.

Набор данных и преобработка

В ходе исследования для тестирования экспериментальной оценки предложенной модели использовался публичный набор данных CICIDS2018, сформированный Канадским институтом кибербезопасности и широко применяемый для моделирования сетевого трафика корпоративных информационных систем [14]. Этот набор данных содержит как легитимный сетевой трафик, так и различные типы сетевых атак, что позволяет применять его для задач обнаружения аномалий или вторжений.

С целью повышения обобщающей способности модели и устойчивости процесса обучения была реализована поэтапная процедура преобработки данных. На первом этапе осуществлен отбор признаков, в ходе которого были исключены низкоинформативные и избыточные признаки, включая служебные метаданные (например, временные метки), а также производные признаки с высокой степенью взаимной корреляции. Это позволило снизить размерность входного пространства и уменьшить риск переобучения модели. Далее проведена семантическая группировка

портов назначения: вместо абсолютных числовых значений порты были классифицированы по трем взаимодополняющим категориям:

- 1) функциональное назначение сервиса (web, mail, база данных);
- 2) категория в соответствии с диапазонами IANA (well-known, registered, ephemeral);
- 3) контекст безопасности (например, категория suspicious для портов, часто ассоциируемых с вредоносной активностью).

Такая группировка позволяет модели оперировать обобщенными типами сетевых сервисов, что повышает ее интерпретируемость и устойчивость к изменениям в сетевой инфраструктуре. На заключительном этапе для устранения асимметрии распределений и статистических выбросов, характерных для сетевого трафика, применен трехступенчатый конвейер нормализации к данным. Он включал использование устойчивого масштабирования на основе межквартильного размаха (RobustScaler), последующее преобразование распределений с использованием метода Яо – Джонсона (PowerTransformer), а также ограничение значений признаков в диапазоне $[-4, 4]$ для обеспечения численной устойчивости обучения нейронной сети. Принципиально важно, что все параметры преобразователей оценивались исключительно на подвыборке легитимного трафика обучающего набора данных, что полностью исключает информационную утечку и корректно моделирует условия реальной эксплуатации системы обнаружения аномалий.

Для визуальной оценки результатов преобработки и анализа структуры признакового пространства была выполнена проекция данных на первые три главных компонента методом PCA. На рис. 2 представлена трехмерная визуализация, из которой видно, что легитимный трафик формирует сложную структуру, состоящую из нескольких пересекающихся кластеров, в то время как аномальные наблюдения распределены более диффузно и частично перекрываются с областями нормальной активности. Такое распределение подтверждает отсутствие простой линейной делимости между классами и обосно-

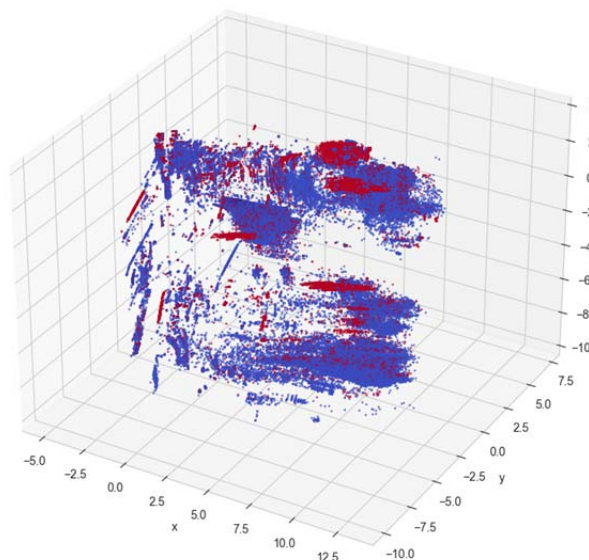


Рис. 2. 3D-проекция признакового пространства (PCA).
Цветом выделены классы: нормальный трафик (синий) и атаки (красный)
Fig. 2. 3D-feature space projection (PCA principal component analysis).
Classes are highlighted in color. Normal traffic (blue) and attacks (red)

вызывает применение нелинейных методов обнаружения аномалий, в частности автоэнкодеров.

Контрольные методы для сравнительного анализа

Для объективной оценки эффективности предложенного глубокого автоэнкодера были выбраны два классических алгоритма обнаружения аномалий, функционирующих в парадигме обучения без учителя.

1. One-Class SVM (OC-SVM) [15] – метод, основанный на построении разделяющей гиперплоскости в пространстве признаков, охватывающей основную массу наблюдений, соответствующих нормальному режиму функционирования системы. Алгоритм обучается исключительно на данных легитимной активности и широко применяется для задач обнаружения аномалий в условиях отсутствия размеченных примеров атак.

2. Isolation Forest (IF) [16] – ансамблевый алгоритм, использующий принцип изоляции аномальных наблюдений. Метод основан на построении набора случайных деревьев ре-

шений, в которых аномальные объекты, как правило, изолируются на меньшей глубине по сравнению с нормальными данными. Isolation Forest является одним из наиболее распространенных методов неконтролируемого обнаружения аномалий.

Оба контрольных алгоритма были обучены и протестированы на одном и том же наборе данных с применением идентичной процедуры предобработки, описанной в «набор данных и предобработка». Настройка гиперпараметров моделей (включая коэффициент контаминации для Isolation Forest, а также тип ядра и параметр ν для OC-SVM) осуществлялась на валидационной выборке, содержащей исключительно легитимный сетевой трафик. Такой подход обеспечивает корректность и сопоставимость результатов сравнительного анализа, представленного в разделе экспериментальных исследований.

Процедура обучения и оценки

Целями экспериментальной процедуры являлись оценка способности предложенной модели на основе автоэнкодера выявлять аномальные состояния сетевого трафика,

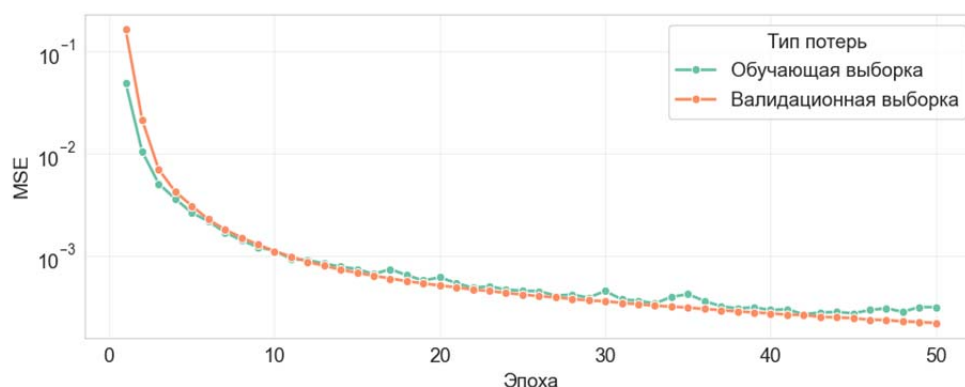


Рис. 3. Процесс обучения автоэнкодера – график функции потерь на обучающей и валидационных выборках
Fig. 3. The autoencoder training process is a graph of the loss functions on training and validation samples

а также анализ устойчивости и воспроизводимости полученных результатов. Все вычислительные эксперименты были реализованы с использованием языка программирования Python версии 3.13. Для построения и обучения нейронной сети применялась библиотека PyTorch, для предобработки данных и расчета метрик качества – библиотека scikit-learn, для работы с массивами данных – NumPy и Pandas.

Архитектура автоэнкодера подбиралась эмпирически на основе серии предварительных экспериментов. В итоговой конфигурации использовалась симметричная структура, включающая шесть скрытых слоев в энкодере с размерностями [256, 192, 144, 108, 81, 64] и латентным представлением размерностью 32. Декодер имел зеркальную архитектуру. Обучение модели осуществлялось исключительно на подвыборке легитимного сетевого трафика с использованием оптимизатора Adam. В качестве функции потерь использовалась среднеквадратичная ошибка реконструкции (MSE). Анализ кривых обучения (рис. 3) свидетельствует о стабильной сходимости модели и отсутствии признаков переобучения.

Ключевым этапом настройки детектора аномалий являлось определение порогового значения функции потерь реконструкции τ , используемого для принятия решения о наличии аномалии. Значение порога определялось на отдельной валидационной выборке, содержащей исключительно данные нор-

мальной сетевой активности, по критерию максимизации индекса Юдена (Youden's J statistic). Такой подход соответствует реальным условиям эксплуатации систем обнаружения вторжений, при которых информация о возможных атаках на этапе конфигурации системы отсутствует.

На рис. 4 представлено распределение ошибок реконструкции автоэнкодера для легитимного трафика и атакующих воздействий в логарифмическом масштабе. Видно, что ошибки реконструкции нормального трафика преимущественно сконцентрированы в области малых значений ошибки реконструкции, тогда как для атак характерны существенно большие значения ошибки и более сложная, многомодальная структура распределения в силу разнообразности классов атак.

Несмотря на частичное перекрытие распределений в области малых значений MSE, между ними наблюдается выраженный статистический сдвиг, что подтверждает информативность ошибки реконструкции в качестве признака аномального поведения. Вертикальной пунктирной линией на рис. 4 обозначено выбранное пороговое значение τ .

Окончательная оценка эффективности разработанной модели проводилась на тестовой выборке, включающей 14 различных типов атак. Для всестороннего анализа использовался набор стандартных метрик качества классификации: макроусредненная F1-мера, площадь под ROC-кривой (ROC-AUC), точность (precision) и полнота (recall). Применен-

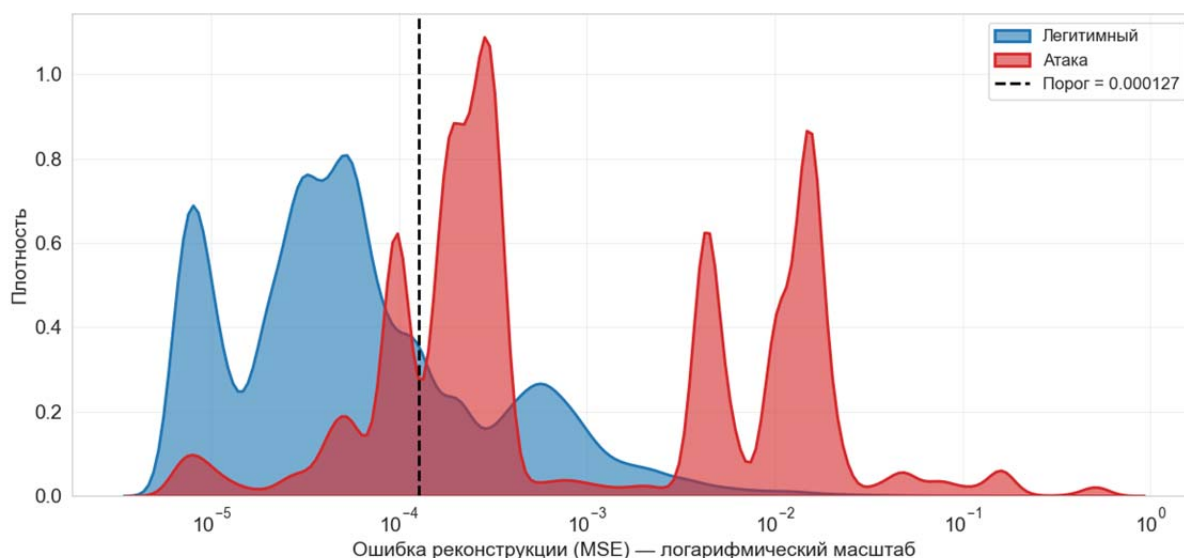


Рис. 4. Распределение ошибок реконструкции (логарифмический масштаб)
Fig. 4. Distribution of reconstruction errors (log scale)

ние нескольких метрик позволило объективно оценить как общую эффективность детектирования, так и баланс между ошибками первого и второго рода.

Результаты и обсуждение

Экспериментальная оценка предложенного автоэнкодера на тестовой выборке, содержащей 14 типов атак, показала его устойчивую и сбалансированную эффективность при решении задачи обнаружения аномалий. В частности, модель продемонстрировала макроусредненную F1-меру 0,81 и значение площади под ROC-кривой (ROC-AUC), равное 0,844.

Анализ показателей по классам свидетельствует о сопоставимом качестве детектирования легитимного трафика и атакующих воздействий. Для легитимного трафика значения точности и полноты составили 0,84 и 0,78 соответственно, тогда как для атак — 0,78 и 0,81. Такое соотношение метрик указывает на сбалансированный характер принимаемых решений и отсутствие выраженного смещения в сторону ошибок первого либо второго рода. Дополнительно стабильная сходимость функции потерь в процессе обучения, а также отчетливое разделение рас-

пределений ошибок реконструкции (рис. 4) подтверждают корректность обучения модели и обоснованность выбранного порога детектирования.

Для оценки относительной эффективности предложенного подхода было проведено сравнение с классическими методами неконтролируемого обнаружения аномалий — Isolation Forest и One-Class SVM [15]. Результаты сравнительного анализа представлены в табл. 1. Как видно из полученных данных, автоэнкодер демонстрирует более высокие значения макроусредненной F1-меры и ROC-AUC по сравнению с контрольными алгоритмами. Превышение составляет порядка 5–8 процентных пунктов по F1-score.

Данное преимущество может быть объяснено способностью глубокой нейронной архитектуры эффективно моделировать сложные нелинейные зависимости в многомерном пространстве признаков и формировать более точное представление нормального поведения сетевого трафика. В отличие от методов, основанных на геометрических предположениях или локальной изоляции наблюдений, автоэнкодер адаптируется к статистической структуре данных, что особенно важно в условиях высокой вариативности и шумности сетевого трафика.

Таблица 1
Table 1

Сравнительные метрики методов неконтролируемого обнаружения аномалий
Comparative metrics of uncontrolled anomaly detection methods

Метод	F1-score (макро)	ROC-AUC
Автоэнкодер (AE)	0,81	0,844
Isolation Forest (IF)	0,75	0,801
One-Class SVM (OC-SVM)	0,72	0,775

Заключение

В рамках проведенного исследования показано, что автоэнкодер, обучаемый в парадигме неконтролируемого обучения, может эффективно использоваться для обнаружения аномальных состояний сетевого трафика и выявления ранее неизвестных атак в корпоративных информационных системах авиапредприятий. Полученные экспериментальные результаты подтверждают возможность применения данного подхода в качестве проактивного аналитического модуля в составе гибридных систем обнаружения вторжений, функционирующих в центрах мониторинга безопасности (SOC).

Ключевым преимуществом предложенного решения является отсутствие необходимости в использовании размеченных данных атак и регулярного обновления сигнатурных баз данных. Это обстоятельство имеет принципиальное значение для авиационной отрасли, характеризующейся жесткими регламентами эксплуатации, высокими требованиями к отказоустойчивости и ограниченными возможностями оперативного обновления программных компонентов. Интеграция оценки аномальности, формируемой автоэнкодером, в существующие SIEM-системы позволяет дополнять традиционные сигнатурные оповещения поведенческим контекстом, что способствует более точной приоритизации инцидентов.

Вместе с тем следует отметить ограничения проведенного исследования: экспериментальная оценка модели выполнялась с использованием синтетического набора данных

CICIDS2018, который, несмотря на широкое распространение в научных исследованиях, не в полной мере отражает особенности трафика и эксплуатационные условия информационных систем авиационной отрасли. В связи с этим практическое внедрение предложенного подхода требует дополнительной валидации на реальных операционных данных.

Перспективными направлениями дальнейших исследований являются повышение интерпретируемости моделей автоэнкодеров за счет разработки методов анализа вклада отдельных признаков в формирование ошибки реконструкции, а также адаптация предложенной методологии к условиям промышленной эксплуатации в составе действующих SOC. Реализация указанных направлений позволит сократить время анализа инцидентов и повысить доверие специалистов информационной безопасности к результатам, формируемым методами машинного обучения.

Список литературы

1. Аврамчиков В.М., Тимохович А.С., Рожнов И.П. Цифровая трансформация в авиационной отрасли: возможности и перспективы [Электронный ресурс] // Вестник Евразийской науки. 2024. Т. 16, № 3. 13 с. URL: <https://esj.today/PDF/04ECVN324.pdf> (дата обращения: 23.11.2025).
2. Scarfone K., Mell P. Guide to intrusion detection and prevention systems (IDPS): NIST Special Publication 800-94. Gaithersburg: National Institute of Standards and Technology, 2007. 127 p.

3. **Ahmed M., Mahmood A.N., Hu J.** A survey of network anomaly detection techniques // *Journal of Network and Computer Applications*. 2016. Vol. 60. Pp. 19–31. DOI: 10.1016/j.jnca.2015.11.016

4. **Buczak A.L., Guven E.** A survey of data mining and machine learning methods for cyber security intrusion detection [Электронный ресурс] // *Communications Surveys & Tutorials*. 2016. Vol. 18, no. 2. Pp. 1153–1176. DOI: 10.1109/COMST.2015.2494502 (дата обращения: 23.11.2025).

5. **Chalapathy R., Chawla S.** Deep learning for anomaly detection: A survey [Электронный ресурс] // *ACM Computing Surveys*. 2019. 50 p. DOI: 10.48550/arXiv.1901.03407 (дата обращения: 23.11.2025).

6. **Pang G.** Deep learning for anomaly detection: A review / G. Pang, C. Shen, L. Cao, A. van den Hengel [Электронный ресурс] // *ACM Computing Surveys*. 2020. 36 p. DOI: 10.1145/3439950 (дата обращения: 23.11.2025).

7. **Ruff L.** Deep one-class classification / L. Ruff, R.A. Vandermeulen, N. Görnitz, L. Deecke // *Proceedings of the 35th International Conference on Machine Learning (ICML)*, 2018. Vol. 80. Pp. 4393–4402.

8. **Zhou C., Paffenroth R.** Robust deep autoencoders for unsupervised anomaly detection // *Proceedings of the 23rd ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD)*, 2017. Pp. 665–674. DOI: 10.1145/3097983.3098052

9. **Mirsky Y.** Kitsune: an ensemble of autoencoders for online network intrusion detection / Y. Mirsky, T. Doitshman, Y. Elovici, A. Shabtai [Электронный ресурс] // *Network and Distributed Systems Security (NDSS) Symposium 2018*, 15 p. DOI: 10.14722/ndss.2018.23204 (дата обращения: 23.11.2025).

10. **Javaid A.** A deep learning approach for network intrusion detection system / A. Javaid, Q. Niyaz, W. Sun, M. Alam [Электронный ресурс] // *9th EAI International Conference on Bio-inspired Information and Communications Technologies*. 2016. DOI: 10.4108/eai.3-12-2015.2262516 (дата обращения: 23.11.2025).

11. **Shone N.** A deep learning approach to network intrusion detection / N. Shone,

T.N. Ngoc, V.D. Phai, Q. Shi [Электронный ресурс] // *IEEE Transactions on Emerging Topics in Computational Intelligence*. 2018. Vol. 2, no. 1. Pp. 41–50. DOI: 10.1109/TETCI.2017.2772792 (дата обращения: 23.11.2025).

12. **Goldstein M., Uchida S.** A comparative evaluation of unsupervised anomaly detection algorithms for multivariate data [Электронный ресурс] // *PLoS ONE*. 2016. Vol. 11, no. 4. ID: e0152173. DOI: 10.1371/journal.pone.0152173 (дата обращения: 23.11.2025).

13. **Sakurada M., Yairi T.** Anomaly detection using autoencoders with nonlinear dimensionality reduction [Электронный ресурс] // *Proceedings of the MLSDA 2014 2nd Workshop on Machine Learning for Sensory Data Analysis*, 2014. Pp. 4–11. DOI: 10.1145/2689746.2689747 (дата обращения: 23.11.2025).

14. **Kumar A.** A survey of CICIDS 2017 and CSE-CIC-IDS2018 datasets [Электронный ресурс] // *AIP Conference Proceedings*. 2024. Vol. 3085. ID: 050001. DOI: 10.1063/5.0222131 (дата обращения: 23.11.2025).

15. **Schölkopf B.** Estimating the support of a high-dimensional distribution / B. Schölkopf, J.C. Platt, J. Shawe-Taylor, A.J. Smola, R.C. Williamson // *Neural Computation*. 2001. Vol. 13, no. 7. Pp. 1443–1471. DOI: 10.1162/089976601750264965

16. **Liu F.T., Ting K.M., Zhou Z.-H.** Isolation Forest [Электронный ресурс] // *2008 Eighth IEEE International Conference on Data Mining (ICDM)*, 2008. Pp. 413–422. DOI: 10.1109/ICDM.2008.17 (дата обращения: 23.11.2025).

References

1. **Avramchikov, V.M., Timokhovich, A.S., Rozhnov, I.P.** (2024). Digital transformation in the aviation industry: opportunities and prospects. *The Eurasian Scientific Journal*, vol. 16, no. 3, 13 p. Available at: <https://esj.today/PDF/04ECVN324.pdf> (accessed: 23.11.2025). (in Russian)

2. **Scarfone, K., Mell, P.** (2007). Guide to intrusion detection and prevention systems (IDPS): NIST Special Publication 800-94.

Gaithersburg: National Institute of Standards and Technology, 127 p.

3. **Ahmed, M., Mahmood, A.N., Hu, J.** (2016). A survey of network anomaly detection techniques. *Journal of Network and Computer Applications*, vol. 60, pp. 19–31. DOI: 10.1016/j.jnca.2015.11.016

4. **Buczak, A.L., Guven, E.** (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *Communications Surveys & Tutorials*, vol. 18, no. 2, pp. 1153–1176. DOI: 10.1109/COMST.2015.2494502 (accessed: 23.11.2025).

5. **Chalapathy, R., Chawla, S.** (2021). Deep learning for anomaly detection: A survey. *ACM Computing Surveys*, vol. 54, no 2, 50 p. DOI: 10.48550/arXiv.1901.03407 (accessed: 23.11.2025).

6. **Pang, G., Shen, C., Cao, L., van den Hengel, A.** (2020). Deep learning for anomaly detection: A review. *ACM Computing Surveys*, 36 p. DOI: 10.1145/3439950 (accessed: 23.11.2025).

7. **Ruff, L., Vandermeulen, R., Görnitz, N., Deecke, L.** (2018). Deep one-class classification. In: *Proceedings of the 35th International Conference on Machine Learning (ICML)*, vol. 80, pp. 4393–4402.

8. **Zhou, C., Paffenroth, R.** (2017). Robust deep autoencoders for unsupervised anomaly detection. In: *Proceedings of the 23rd ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD)*, pp. 665–674. DOI: 10.1145/3097983.3098052

9. **Mirsky, Y., Doitshman, T., Elovici, Y., Shabtai, A.** (2018). Kitsune: an ensemble of autoencoders for online network intrusion detection. In: *Network and Distributed Systems Security (NDSS) Symposium*, 15 p. DOI: 10.14722/ndss.2018.23204 (accessed: 23.11.2025).

10. **Javaid, A., Niyaz, Q., Sun, W., Alam, M.** (2016). A deep learning approach for network intrusion detection system. In: *9th EAI International Conference on Bio-inspired Information and Communications Technologies*. DOI: 10.4108/eai.3-12-2015.2262516 (accessed: 23.11.2025).

11. **Shone, N., Ngoc, T.N., Phai, V.D., Shi, Q.** (2018). A deep learning approach to network intrusion detection. In: *IEEE transactions on emerging topics in computational intelligence*, vol. 2, no. 1, pp. 41–50. DOI: 10.1109/TETCI.2017.2772792 (accessed: 23.11.2025).

12. **Goldstein, M., Uchida, S.** (2016). A comparative evaluation of unsupervised anomaly detection algorithms for multivariate data. *PLoS ONE*, vol. 11, no. 4, ID: e0152173. DOI: 10.1371/journal.pone.0152173 (accessed: 23.11.2025).

13. **Sakurada, M., Yairi, T.** (2014). Anomaly detection using autoencoders with nonlinear dimensionality reduction. In: *Proceedings of the MLSDA 2014 2nd Workshop on Machine Learning for Sensory Data Analysis*, pp. 4–11. DOI: 10.1145/2689746.2689747 (accessed: 23.11.2025).

14. **Kumar, A.** (2024). A survey of CICIDS 2017 and CSE-CIC-IDS2018 datasets. In: *AIP Conference Proceedings*, vol. 3085. ID: 050001. DOI: 10.1063/5.0222131 (accessed: 23.11.2025).

15. **Schölkopf, B., Platt, J.C., Shawe-Taylor, J., Smola, A.J., Williamson, R.C.** (2001). Estimating the support of a high-dimensional distribution. *Neural Computation*, vol. 13, no. 7, pp. 1443–1471. DOI: 10.1162/089976601750264965

16. **Liu, F.T., Ting, K.M., Zhou, Z.-H.** (2008). Isolation Forest. In: *2008 Eighth IEEE International Conference on Data Mining (ICDM)*, pp. 413–422. DOI: 10.1109/ICDM.2008.17 (accessed: 23.11.2025).

Сведения об авторах

Машошин Никита Олегович, инженер по машинному обучению, ООО «СофтТелематика», аспирант кафедры основ радиотехники и защиты информации МГТУ ГА, nikita.mashoshin@yandex.ru.

Кулешов Александр Анатольевич, доктор технических наук, заместитель генерального директора по гражданской авиатехнике АО «НПП "Аэросила"», kaa.aerosila@mail.ru.

Information about the authors

Nikita O. Mashoshin, Machine Learning Engineer, SoftTelematika LLC, Postgraduate Student of the Radio Engineering Fundamentals and Information Security Chair, Moscow State Technical University of Civil Aviation, nikita.mashoshin@yandex.ru.

Alexander A. Kuleshov, Doctor of Technical Sciences, Deputy General Director for Civil Aviation Products of JSC Scientific-Production Enterprise "Aerosila"; kaa.aerosila@mail.ru.

Поступила в редакцию	02.02.2026	Received	02.02.2026
Одобрена после рецензирования	10.03.2026	Approved after reviewing	10.03.2026
Принята в печать	21.05.2026	Accepted for publication	21.05.2026