

Обушев Антон Константинович
магистрант МТУСИ, Москва, Россия

Ларин Александр Иванович
доцент кафедры ИСУиА и КиИБ МТУСИ, к.т.н., доцент, Москва, Россия

ИНТЕЛЛЕКТУАЛЬНЫЕ СИСТЕМЫ В УПРАВЛЕНИИ БЕЗОПАСНОСТЬЮ: АНАЛИЗ МЕТОДОВ ВЫЯВЛЕНИЯ ПОВЕДЕНЧЕСКИХ АНОМАЛИЙ ПОЛЬЗОВАТЕЛЕЙ ДЛЯ АВТОМАТИЗИРОВАННОГО УПРАВЛЕНИЯ ИНФОРМАЦИОННОЙ БЕЗОПАСНОСТЬЮ

Аннотация. Современные системы защиты информации всё чаще сталкиваются не с сигнатурными атаками, а с трудноформализуемыми поведенческими отклонениями, возникающими в рамках легитимной активности пользователей. Это смещает фокус с детерминированных правил на вероятностные и обучаемые модели, способные выявлять аномалии без явной разметки данных. В центре внимания оказываются методы класса anomaly detection, применяемые в архитектурах UEBA и интегрируемые в SIEM-платформы.

В статье выполнен сравнительный анализ ключевых подходов к выявлению поведенческих аномалий — от статистических моделей до методов глубокого обучения. Особое внимание уделяется формализации задач, ограничениям обучающих выборок и особенностям применения моделей в условиях высокой размерности и шумных данных. Рассматриваются алгоритмы One-Class SVM, Isolation Forest, автоэнкодеры и скрытые марковские модели, с акцентом на их математические основания и практическую интерпретацию в контексте информационной безопасности.

Ключевые слова: информационная безопасность, поведенческий анализ пользователей, UEBA, обнаружение аномалий, машинное обучение, One-Class SVM, Isolation Forest, автоэнкодеры, скрытые марковские модели, информационная безопасность, SIEM.

Abstract. Modern information security systems increasingly face not signature-based attacks, but rather difficult-to-formalize behavioral anomalies arising from legitimate user activity. This shifts the focus from deterministic rules to probabilistic and trainable models capable of detecting anomalies without explicit data labeling. This paper focuses on anomaly detection methods used in UEBA architectures and integrated into SIEM platforms.

This article provides a comparative analysis of key approaches to behavioral anomaly detection, from statistical models to deep learning methods. Particular attention is paid to task formalization, limitations of training sets, and the specifics of model application in high-dimensional and noisy data environments. One-Class SVM, Isolation Forest, autoencoders, and hidden Markov models are discussed, with an emphasis on their mathematical foundations and practical interpretation in the context of information security.

Keywords: information security, user behavior analysis, UEBA, anomaly detection, machine learning, One-Class SVM, Isolation Forest, autoencoders, hidden Markov models, information security, SIEM.

Введение

Практика эксплуатации корпоративных информационных систем постепенно смещает акценты: если ранее основную нагрузку несли сигнатурные механизмы обнаружения атак, то сегодня значительная доля инцидентов связана с действиями, формально не нарушающими правил доступа. Речь идёт о сценариях, в которых злоумышленник действует под легитимной учётной записью — будь то компрометация, инсайдерская активность или автоматизированные атаки, имитирующие поведение пользователя [1]. В подобных условиях классические средства контроля, основанные на заранее заданных правилах, оказываются недостаточными.

Проблема усугубляется тем, что «аномальность» поведения редко имеет однозначное определение. Одно и то же действие — например, доступ к критическим данным в нерабочее время — может быть как признаком атаки, так и следствием нормальной рабочей нагрузки. В результате задача

обнаружения смещается в область анализа контекста и динамики поведения, где статические политики уступают место моделям, способным учитывать индивидуальные паттерны активности.

Дополнительное усложнение вносит дефицит размеченных данных. Инциденты безопасности относительно редки и разнообразны, что делает невозможным построение полноценных обучающих выборок для классических задач классификации. На практике это приводит к доминированию методов, работающих в парадигме обучения без учителя или с частичным контролем, где модель формирует представление о «норме», а отклонения интерпретируются как потенциальные аномалии.

Именно в этом контексте развиваются интеллектуальные системы управления безопасностью, в том числе решения класса *UEBA (User and Entity Behavior Analytics)*, интегрируемые в *SIEM-платформы (Security Information and Event Management)* [10]. Их задача — не просто фиксировать события, а выявлять скрытые закономерности в поведении пользователей и автоматически адаптировать уровень реагирования [9]. Однако эффективность таких систем определяется не только точностью моделей, но и их устойчивостью к изменению поведения, интерпретируемостью и способностью работать в условиях высокой неопределённости.

Переход к подобным подходам требует систематизации используемых методов и критического анализа их применимости. Это особенно важно, поскольку на практике выбор алгоритма часто определяется не теоретическими основаниями, а удобством внедрения или доступностью реализации, что не всегда коррелирует с качеством выявления аномалий.

Постановка задачи выявления поведенческих аномалий

Формализация задачи обнаружения поведенческих аномалий оказывается сложнее, чем это может показаться на уровне интуиции. В классическом понимании аномалией считается наблюдение [2], существенно отклоняющееся от некоторого «нормального» распределения данных. Однако в контексте информационной безопасности сама норма является

динамической и зачастую индивидуальной для каждого пользователя или группы.

Ключевая трудность заключается в том, что распределение нормальных данных заранее неизвестно и может изменяться во времени. Более того, сами признаки часто коррелированы и подвержены шуму, что делает явное моделирование плотности затруднительным. В результате задача приобретает форму либо оценки плотности (*density estimation*), либо построения границы области нормальности (*boundary estimation*).

Отсутствие размеченных данных усиливает эту неопределённость. В отличие от классических задач классификации, где модель опирается на размеченные данные, в данном случае обучение чаще всего происходит только на выборке, отражающей нормальное поведение. Это сразу накладывает ограничения: модель не «видит» примеры атак и вынуждена ориентироваться лишь на структуру обычной активности.

Ситуация дополнительно осложняется тем, что сама норма оказывается нестабильной. Поведение пользователя со временем меняется — иногда постепенно, иногда довольно резко. В результате то, что ещё недавно выглядело как отклонение, через некоторое время воспринимается как допустимое действие [16]. На практике это приводит к необходимости использовать адаптивные подходы, которые могут подстраиваться под новые паттерны без полного переобучения модели.

Не менее важен и способ представления данных. Поведенческие характеристики можно описывать по-разному: через агрегированные показатели за период, последовательности действий или, например, структуры взаимодействий между объектами [13]. От выбранного формата зависит и сама постановка задачи. Если данные представлены в виде последовательностей, логично использовать вероятностные модели, учитывающие порядок событий. Когда же речь идёт об агрегированных признаках, чаще применяются методы, работающие с геометрией пространства признаков и границами между «нормой» и отклонениями.

Задача выявления поведенческих аномалий находится на пересечении нескольких классов задач машинного обучения [12]:

- **классификация** — при наличии частичной разметки;
- **оценка плотности** — при моделировании распределения нормы;
- **поиск выбросов** — при выявлении точек, удалённых от основной массы данных.

На практике это приводит к тому, что границы между подходами размываются [1], а конкретная реализация определяется компромиссом между точностью, интерпретируемостью и вычислительной сложностью.

Классификация методов выявления поведенческих аномалий

Несмотря на разнообразие алгоритмов, применяемых для анализа поведенческих данных, их удобно рассматривать через призму того, каким образом моделируется «норма» и как определяется отклонение от неё. В этом смысле различия между методами носят не только технический, но и концептуальный характер [1, 22]: одни стремятся явно аппроксимировать распределение данных, другие — выделить геометрию пространства нормальных наблюдений, третьи — реконструировать поведение и анализировать ошибки восстановления [14].

Статистические методы исторически являются отправной точкой. В их основе лежит предположение о существовании некоторого распределения, описывающего нормальное поведение.

На практике используются как параметрические модели (например, гауссовские смеси), так и непараметрические оценки плотности [2]. Однако в задачах информационной безопасности такие подходы быстро сталкиваются с ограничениями: реальные данные редко подчиняются простым распределениям, а высокая размерность приводит к «размыванию» плотности (эффект проклятия размерности).

Методы классического машинного обучения частично обходят эту проблему, переходя от явного моделирования плотности к построению

разделяющих границ. Здесь можно выделить два подподхода. Первый — алгоритмы *one-class learning*, которые формируют компактную область, содержащую нормальные данные [4]. Второй — методы, основанные на изоляции или локальной плотности (*Isolation Forest, LOF*), где аномалии интерпретируются как точки, легко отделяемые от основной массы наблюдений [3, 5]. В отличие от статистических моделей, такие методы менее чувствительны к форме распределения, но их поведение становится сложнее интерпретировать в терминах вероятности.

Методы глубокого обучения предлагают иной взгляд на задачу, смещая акцент с явного моделирования на обучение представлений. Наиболее характерный пример — автоэнкодеры, где модель обучается реконструировать входные данные, а ошибка восстановления используется как мера аномальности [6].

Более сложные архитектуры — в частности, вариационные автоэнкодеры и рекуррентные сети — дают возможность учитывать временную структуру поведения, что особенно важно при анализе пользовательской активности [7]. За счёт этого удаётся уловить зависимости, которые не видны в статичных признаках. Впрочем, за такую выразительность приходится платить: возрастают требования к объёму данных, увеличивается вычислительная нагрузка, а сами решения становятся менее прозрачными с точки зрения интерпретации.

Модели вероятностного типа, ориентированные на последовательности, например скрытые марковские модели, занимают в этом ряду скорее промежуточную позицию [8]. Они описывают поведение как стохастический процесс, в котором каждое действие связано с предыдущим, пусть и не напрямую.

В таком подходе аномалии проявляются не как отдельные «выбросы», а как маловероятные последовательности событий. Это особенно заметно в сценариях, где критичен именно порядок действий: сами по себе операции

могут выглядеть допустимыми, но их комбинация или последовательность уже выходит за рамки привычного поведения.

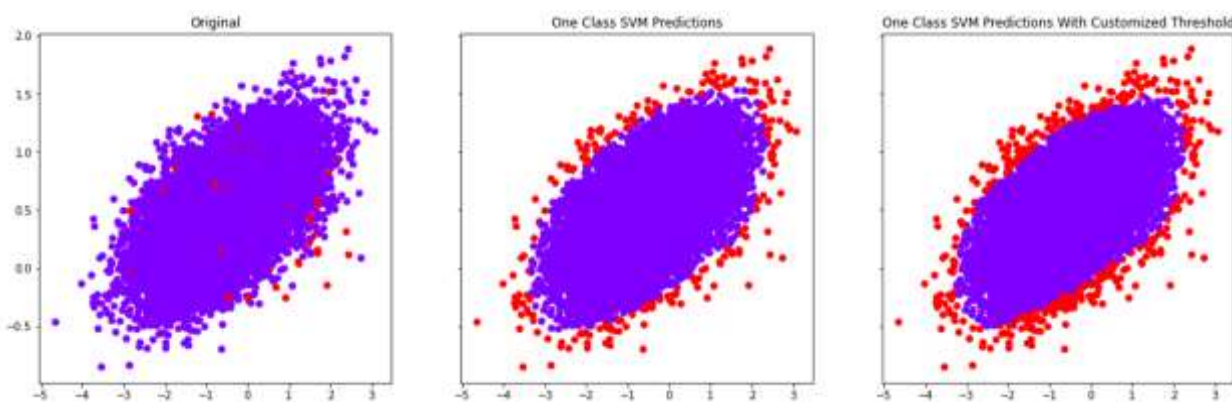
Если рассматривать эти классы методов в совокупности, становится заметно, что различия между ними постепенно стираются. Например, автоэнкодеры можно интерпретировать как неявные модели плотности, а методы изоляции — как приближение к оценке локальной структуры данных. В практических системах ИБ это приводит к гибридизации подходов: комбинируются геометрические, вероятностные и нейросетевые методы, что позволяет компенсировать слабые стороны каждого из них.

При этом выбор конкретного класса редко определяется только теоретическими соображениями. Более значимыми оказываются факторы устойчивости к изменению поведения, масштабируемости и способности объяснить результат — особенно в контексте автоматизированного принятия решений.

Анализ методов выявления поведенческих аномалий

Различие между подходами особенно проявляется на уровне их формализации: каждая модель по-своему отвечает на вопрос, что считать «отклонением». Ниже рассмотрены четыре метода, которые часто используются в задачах поведенческой аналитики и при этом опираются на разные математические принципы [21].

One-Class SVM



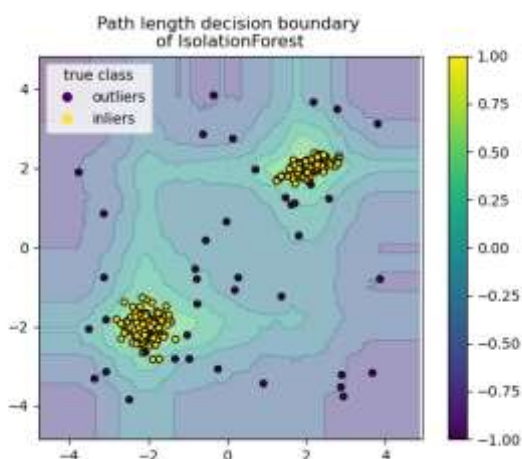
Идея *One-Class SVM* сводится к построению границы, отделяющей нормальные данные от начала координат в преобразованном пространстве

признаков [4]. Фактически модель ищет область минимального объёма, содержащую большинство наблюдений.

В контексте ИБ это означает, что поведение пользователя считается нормальным, если оно лежит внутри сформированной области. На практике поведение метода во многом определяется настройками. В частности, выбор ядра и корректное масштабирование признаков оказывают заметное влияние на результат. При увеличении размерности данные начинают «рассыпаться» в пространстве, и построенная граница становится менее наглядной — интерпретировать её уже не так просто.

Isolation Forest

В отличие от подходов, где сначала пытаются описать нормальное поведение, здесь используется другая логика: аномальные объекты проще «отделить» от остальных с помощью случайных разбиений пространства [5].



Алгоритм формирует ансамбль случайных деревьев. На каждом шаге пространство делится по случайному признаку, и для каждой точки фиксируется длина пути до листа, в котором она оказывается изолированной.

Сильная сторона метода в том, что он не требует явного задания распределения данных. За счёт этого он сравнительно устойчив к сложным и неоднородным выборкам. В задачах информационной безопасности это особенно заметно при работе с агрегированными характеристиками — например, статистиками активности. В то же время временные зависимости он

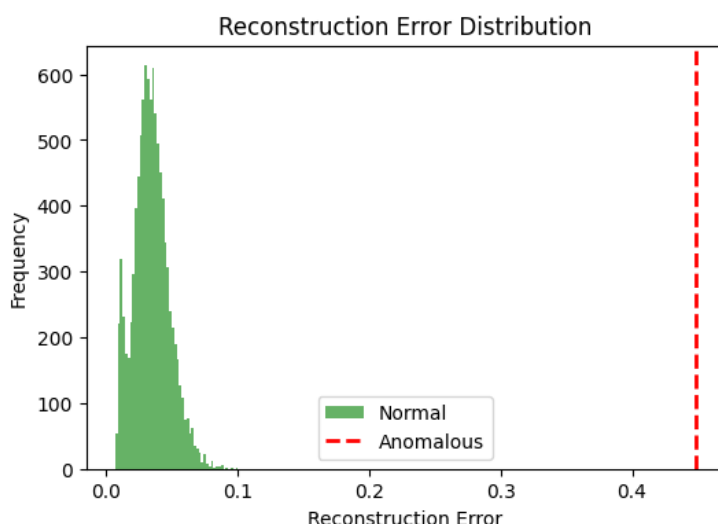
практически не учитывает, если они не были заранее отражены в признаках [15].

Автоэнкодеры

Автоэнкодер — это нейронная сеть, которая обучается восстанавливать входные данные, пропуская их через сжатое представление [6]. Логика здесь довольно простая: если модель «привыкла» к нормальному поведению, она будет воспроизводить его без существенных искажений. А вот на нетипичных примерах качество восстановления, как правило, заметно падает.

В задачах информационной безопасности такие модели оказываются полезными за счёт способности улавливать сложные, нелинейные зависимости между признаками. Это особенно важно, когда поведение пользователя нельзя описать простыми правилами. Однако на практике возникает неприятный эффект: если в обучающей выборке уже присутствуют аномалии, сеть может частично «выучить» и их, снижая чувствительность модели.

Есть и другая проблема — интерпретация результатов. Высокая ошибка реконструкции сигнализирует о возможном отклонении, но сама по себе не объясняет, в чём именно оно заключается. Определить, какие признаки стали причиной аномалии, зачастую оказывается нетривиальной задачей.



Скрытые марковские модели (HMM)

Если рассматривать поведение пользователя как цепочку действий, логично переходить к моделям, которые учитывают временную структуру [8].

В рамках НММ предполагается, что наблюдаемые события зависят от скрытых состояний, которые напрямую не наблюдаются, но определяют динамику процесса.

Для задач информационной безопасности это даёт важное преимущество: появляется возможность анализировать не отдельные действия, а их порядок. Например, последовательность обращений к ресурсам может выглядеть подозрительно именно из-за своей структуры, даже если каждое действие по отдельности допустимо.

В то же время такие модели требуют аккуратной настройки. Число скрытых состояний приходится подбирать, и при усложнении поведения пользователей вычислительная нагрузка быстро растёт. Это ограничивает масштабируемость, особенно в больших системах.

На практике это приводит к тому, что один и тот же тип аномалии фиксируется разными моделями по-разному. Поведение, которое выглядит редким, но при этом сохраняет внутреннюю структуру, может остаться незамеченным автоэнкодером, но при этом оказаться маловероятным с точки зрения последовательностной модели и быть обнаруженным именно там.

Сравнительный анализ методов выявления аномалий

Сопоставление различных подходов позволяет выявить не только их сильные и слабые стороны, но и ограничения, которые становятся критичными в реальных системах информационной безопасности. Ниже приведено обобщённое сравнение рассмотренных методов.

Таблица 1

Результаты сравнительного анализа методов выявления поведенческих аномалий

Метод	Преимущества	Недостатки	Где применяется
<i>One-Class SVM</i>	хорошо отделяет «нормальное» поведение от отклонений	- требует аккуратной настройки - плохо работает на очень больших данных	анализ стабильного поведения пользователей, где паттерны мало меняются

	• работает даже без примеров атак • учитывает сложные зависимости	• сложно объяснить, почему поведение признано аномальным	
<i>Isolation Forest</i>	• быстро работает на больших объёмах данных • не требует сложной настройки • хорошо находит редкие события	• не учитывает последовательность действий • результат может немного «плавать» из-за случайности	первичный поиск подозрительных действий в логах и событиях
Автоэнкодеры	• находят сложные и скрытые зависимости • подходят для большого числа признаков • хорошо выявляют нестандартные комбинации действий	• требуют много данных • могут «привыкнуть» к аномалиям	UEBA-системы, анализ сложного поведения пользователей
<i>HMM</i>	• учитывают порядок действий • позволяют анализировать сценарии поведения • дают понятную вероятностную оценку	• сложно настраивать • плохо работают при очень сложном поведении • ресурсоёмкие	анализ последовательностей действий (например, цепочки доступа к ресурсам)

Результаты проведенного анализа показывают, что ни один из методов не является универсальным. *Isolation Forest* справляется с отсеком аномалий в больших потоках данных на первых стадиях, но может генерировать значительное число ложных срабатываний, не учитывая контекст. *НММ* или автоэнкодеры способны фиксировать сложные зависимости, требуя большего объёма данных и вычислительных ресурсов.

Граница между преимуществами и недостатками определяется контекстом использования. Низкая интерпретируемость нейросетевых моделей критична в системах, где требуется объяснение решений (особенно при расследовании инцидентов), но менее значима в задачах предварительной фильтрации событий.

В системах ИБ используют комбинирование методов: простые алгоритмы используются для быстрого отбора подозрительных событий, более сложные модели уточняют оценку [21]. Каскадный подход позволяет нивелировать ограничения отдельных алгоритмов, что усложняет архитектуру системы и процессы её настройки.

Методы выявления аномалий в автоматизированных системах ИБ в практическом применении

Практическая ценность рассмотренных методов проявляется в составе комплексных систем управления безопасностью — в платформах класса *UEBA* и *SIEM* [10]. Именно здесь происходит переход от абстрактных моделей к операционным решениям, влияющим на процессы обнаружения и реагирования.

В типичной архитектуре *SIEM* система агрегирует события из различных источников: сетевые журналы, действия пользователей, обращения к приложениям, системные вызовы. Однако без интеллектуальной обработки этот поток остаётся набором слабо связанных сигналов. Встраивание моделей обнаружения аномалий позволяет перейти к анализу поведения как целостного процесса, где отдельные события интерпретируются в контексте исторической активности [12].

UEBA-решения, в свою очередь, делают акцент на профилировании пользователей и сущностей [19]. Для каждого субъекта формируется модель нормального поведения, которая затем используется для оценки текущих действий. На этом уровне методы *anomaly detection* интегрируются следующим образом:

- на этапе агрегации признаков — формируются поведенческие векторы (частоты, временные паттерны, корреляции) [13];
- на этапе моделирования — применяются алгоритмы (например, *Isolation Forest* или автоэнкодеры) для построения профиля нормы;
- на этапе оценки риска — вычисляется скор аномальности, который затем преобразуется в индикатор риска;
- на этапе реагирования — система инициирует уведомление или автоматическое действие.

На практике часто используется каскадная схема: быстрые и простые методы (например, деревья изоляции) выполняют предварительную фильтрацию, после чего более сложные модели уточняют оценку. Такой подход в целом позволяет держать разумный баланс между скоростью обработки данных и качеством выявления отклонений. Однако при переходе от теории к реальным системам довольно быстро становятся заметны ограничения, которые в моделях обычно не проявляются.

Прежде всего, проблема ложных срабатываний никуда не исчезает. Даже незначительные изменения в поведении пользователя — например, смена рабочего графика или появление новой задачи — могут быть восприняты системой как отклонение [10]. Когда пользователей много, это приводит к накоплению сигналов, с которыми аналитики просто не успевают работать. В результате уровень доверия к системе постепенно снижается [17].

Не менее сложной оказывается ситуация с интерпретацией результатов. В моделях вроде автоэнкодеров или ансамблей деревьев не всегда понятно, что

именно стало причиной высокого аномального сора. Для задач информационной безопасности это принципиально: без внятного объяснения трудно принимать решения и тем более разбирать инциденты.

Важным учитываемым аспектом является учет изменения поведения со временем. Пользовательские паттерны не остаются неизменными, и модель, обученная на исторических данных, со временем начинает хуже отражать текущую реальность [22]. Чтобы компенсировать этот эффект, приходится либо регулярно переобучать модель, либо внедрять механизмы адаптации, что усложняет сопровождение системы [16].

Наконец, многое упирается в подготовку данных. Большинство моделей работает уже с агрегированными признаками, и именно на этом этапе закладывается значительная часть качества. Если признаки выбраны неудачно или чрезмерно упрощены, даже сложные алгоритмы не смогут компенсировать эти ограничения.

В реальных системах всё чаще уходят от использования одного алгоритма в пользу комбинированных решений. Формируются гибридные архитектуры, в которых сочетаются разные подходы [23]:

- одновременно применяются несколько моделей, дополняющих друг друга;
- учитывается контекст — например, роль пользователя или тип ресурса;
- используется обратная связь от аналитиков для корректировки работы системы.

Полностью устранить ограничения это не позволяет, но поведение системы становится заметно устойчивее в реальных условиях эксплуатации.

Заключение

Анализ развития методов выявления поведенческих аномалий показывает, что основным направлением не является замена одних алгоритмов на другие. Скорее наблюдается постепенное усложнение и комбинирование подходов. Описание «нормы» одной универсальной моделью на практике

ограничено тем, что само понятие «нормы» нестабильно, зависит от контекста и со временем может меняться.

Различные методы по-разному реагируют на данную неопределённость. Геометрические подходы вроде One-Class SVM работают, пока структура данных остаётся относительно стабильной, но теряют устойчивость при увеличении размерности и изменении распределений. Методы изоляции лучше справляются со сложными выборками, но их результаты не всегда легко интерпретировать. Нейросетевые модели дают большую гибкость и способны выявлять сложные зависимости, но требуют значительных ресурсов и редко обеспечивают прозрачность решений. Вероятностные модели позволяют анализировать поведение как процесс, однако их применение ограничивается сложностью настройки и масштабируемости.

Опыт внедрения в системах защиты информации показывает, что одних математических свойств алгоритма недостаточно. Значение имеют качество признаков, устойчивость к изменению поведения пользователей и насколько модель вписывается в процессы мониторинга и реагирования. Эффективность системы складывается из совокупности факторов и не определяется одним выбранным методом. Интеллектуальные системы выступают как элементы широкой инфраструктуры, где присутствует согласованность различных компонентов.

Заслуживает внимания проблема объяснимости. Чем сложнее модель, тем больше разрыв между точностью и интерпретируемостью. Усложнение моделей ограничивает применение некоторых методов в критически значимых системах. На практике это приводит к компромиссным решениям, где используются комбинация моделей с различным уровнем прозрачности.

Перспективное направление развития связано с построением гибридных архитектур, способных адаптироваться к изменяющемуся поведению пользователей и обеспечивать удовлетворяемый уровень интерпретации результатов. Вероятно, именно сочетание различных подходов, дополненное

механизмами обратной связи и контекстного анализа, станет основой следующего поколения систем управления информационной безопасностью.

Список литературы

1. Chandola V., Banerjee A., Kumar V. Anomaly Detection: A Survey // ACM Computing Surveys. — 2009. — Vol. 41, No. 3. — С. 1–58.
2. Aggarwal C. C. Outlier Analysis. — New York: Springer, 2017. — 446 p.
3. Breunig M. M. et al. LOF: Identifying Density-Based Local Outliers // ACM SIGMOD Record. — 2000. — Vol. 29, No. 2. — С. 93–104.
4. Schölkopf B. et al. Estimating the Support of a High-Dimensional Distribution // Neural Computation. — 2001. — Vol. 13, No. 7. — С. 1443–1471.
5. Liu F. T., Ting K. M., Zhou Z.-H. Isolation Forest // IEEE ICDM. — 2008. — P. 413–422.
6. Goodfellow I., Bengio Y., Courville A. Deep Learning. — Cambridge: MIT Press, 2016. — 775 С.
7. Kingma D. P., Welling M. Auto-Encoding Variational Bayes // ICLR. — 2014.
8. Rabiner L. R. A Tutorial on Hidden Markov Models // Proceedings of the IEEE. — 1989. — Vol. 77, No. 2. — С. 257–286.
9. Sommer R., Paxson V. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection // IEEE Symposium on Security and Privacy. — 2010. — С. 305–316.
10. Shashanka M., Shenoy P., Aggarwal C. Advances in User and Entity Behavior Analytics // IEEE Security & Privacy. — 2020. — Vol. 18, No. 5. — С. 46–54.
11. Khaliq A. et al. Intrusion Detection Using Machine Learning: A Survey // IEEE Access. — 2020.
13. Feng W. et al. Multi-Granularity User Anomalous Behavior Detection // Applied Sciences. — 2024. [Электронный ресурс]. Режим доступа: <https://www.mdpi.com/2076-3417/15/1/128> (Дата обращения 01.04.2026).
14. Artioli P. et al. A Comprehensive Investigation of Clustering Algorithms for Anomaly Detection // Frontiers in Big Data. — 2024. [Электронный ресурс].

Режим доступа: <https://www.frontiersin.org/journals/big-data/articles/10.3389/fdata.2024.1375818/full> (Дата обращения 23.03.2026).

15. Fuentes J., Ortega-Fernandez I. et al. Cybersecurity Threat Detection Based on a UEBA Framework Using Deep Autoencoders // arXiv. — 2025. [Электронный ресурс]. Режим доступа: <https://arxiv.org/abs/2505.11542> (Дата обращения 03.04.2026).

16. Moriano P. et al. Adaptive Anomaly Detection for Identifying Attacks in Cyber-Physical Systems: A Systematic Literature Review // arXiv. — 2024. [Электронный ресурс]. Режим доступа: <https://arxiv.org/abs/2411.14278> (Дата обращения 27.03.2026).

17. Huang Z. et al. IDU-Detector: A Synergistic Framework for Robust Masquerader Attack Detection // arXiv. — 2024. [Электронный ресурс]. Режим доступа: <https://arxiv.org/abs/2411.06172> (Дата обращения 27.03.2026).

18. Ahmad W. et al. Leveraging Machine Learning for Internal User Behaviour Analysis // International Journal of IT & Information Systems. — 2025. [Электронный ресурс]. Режим доступа: <https://journals.tultech.eu/index.php/ijitis/article/download/254/244> (Дата обращения 06.04.2026).

19. Ibraheem I. O. et al. Detecting Malicious Insider Threats through Anomaly-Based User Behaviour Analytics // Southern African Journal of Security. — 2025.

20. Tynymbayev B. et al. Applying Hidden Markov Models for User and Entity Behavior Analytics // ScienceDirect. — 2026. [Электронный ресурс]. Режим доступа: www.sciencedirect.com/science/article/pii/S2772941926000311 (Дата обращения 06.04.2026).

21. Баранов А. В., Лысенко А. А. Методы выявления аномалий в задачах информационной безопасности // Информационная безопасность. — 2021. — №3. — С. 45–52.

22. Кузнецов Н. А., Смирнов Д. В. Анализ поведенческих моделей пользователей в системах защиты информации // Вестник МГТУ. — 2022. — №5. — С. 78–86.

23. Марков А. С., Фадеева Н. В. Интеллектуальные методы анализа данных в задачах кибербезопасности // Вестник НГУ. Серия: Информационные технологии. — 2023.