



ИСПОЛЬЗОВАНИЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА ДЛЯ ВЫЯВЛЕНИЯ И БЛОКИРОВКИ МОШЕННИЧЕСТВА В СФЕРЕ ТЕЛЕКОММУНИКАЦИЙ

Адилов Рустам Тимурович

Студент компьютерных наук, Университет ИНХА в Ташкенте

E-mail: rstmadylov@gmail.com

ORCID: 0009-0004-7250-4627

<https://doi.org/10.5281/zenodo.21872993>

ARTICLE INFO

Qabul qilindi: 06-avgust 2026 yil

Ma'qullandi: 08-avgust 2026 yil

Nashr qilindi: 10-avgust 2026 yil

KEYWORDS

искусственный интеллект;
машинное обучение;
телекоммуникационное
мошенничество; IRSF; SIM-
box; обнаружение аномалий;
графовые нейронные сети;
мониторинг в реальном
времени; Random Forest.

ABSTRACT

Статья посвящена анализу основных типов мошенничества, характерных для телекоммуникационной отрасли (IRSF, SIM-box, wangiri-звонки, подмена SIM-карты, CLI-спуфинг), и применению искусственного интеллекта и машинного обучения для их выявления и блокировки. Данные международных отраслевых организаций (CFCA, GLF) свидетельствуют, что в 2023 году глобальные потери операторов от мошенничества достигли 38,95 млрд долларов США, а по последним оценкам этот показатель вырос до 41,82 млрд долларов США. В работе выполнено сравнение показателей качества различных моделей: от правило-ориентированных систем, где полнота обнаружения не превышает 40%, до гибридных решений на базе Random Forest с точностью 92,3% и графовых нейронных сетей, обнаруживающих до 85% мошеннических сетей SIM-карт. Установлено, что переход от классических правил к моделям ИИ повышает полноту обнаружения на 45-50 процентных пунктов и одновременно сокращает долю ложных срабатываний в 3-6 раз. Заключительная часть статьи посвящена ключевым проблемам отрасли - дисбалансу данных, интерпретируемости моделей, адаптивности мошеннических схем - и перспективным направлениям дальнейших исследований.

Введение

Организованное мошенничество продолжает наносить операторам связи один из самых чувствительных финансовых ударов среди всех отраслей экономики. Согласно отчёту Communications Fraud Control Association (CFCA), в 2023 году совокупные потери операторов связи от мошеннических действий достигли отметки примерно в 38,95 млрд долларов США - рост на 12% относительно 2021 года, что соответствует около 2,5%

суммарной выручки отрасли [1]. Более свежие расчёты, опирающиеся на те же данные CFCA, дают уже свыше 41,82 млрд долларов США совокупных потерь; при этом специалисты отрасли подчёркивают, что подобная статистика фиксирует лишь прямой ущерб самих операторов и не учитывает убытки, которые несут конечные абоненты от телефонного мошенничества, а эта величина, по некоторым оценкам, оказывается ещё значительнее [4].

При этом убытки распределены между видами мошенничества крайне неравномерно. Наибольшую долю занимает международное мошенничество с разделением доходов (International Revenue Share Fraud, IRSF), которое в 2023 году обошлось операторам примерно в 6,23 млрд долларов [1]. Далее следует мошенничество через подключение к частным телефонным станциям (PBX fraud) с ежегодным ущербом порядка 4-7 млрд долларов, а также схемы «одного гудка» (wangiri), приносящие отрасли около 2,23 млрд долларов потерь [3]. Особую озабоченность вызывает динамика мошенничества с подменой SIM-карт (SIM-swap): с 2020 года количество зафиксированных случаев увеличилось примерно на 400%, а годовой ущерб превысил 2 млрд долларов [2]. Не менее тревожна и ситуация с роботизированными звонками и синтезированными голосовыми атаками на основе ИИ (AI-driven robocalls, voice cloning) по прогнозам, к 2025 году они могут обойтись отрасли примерно в 76 млрд долларов; параллельно доля операторов, сообщающих о высоком уровне CLI-спуфинга, за год выросла с 49% до 55%, а доля операторов с высоким уровнем спам-трафика - с 46% до 61% [3].

Долгое время основным инструментом защиты операторов оставались статичные правила и пороговые значения - контроль числа международных вызовов за интервал времени, отслеживание аномальной длительности соединений, проверка подозрительных направлений трафика. Подобные механизмы легко внедряются, но плохо справляются с постоянно меняющейся тактикой мошенников и генерируют большое количество ложных тревог, из-за чего службы фрод-мониторинга перегружены, а доверие к автоматическим оповещениям падает. По имеющимся данным, правило-ориентированные решения обнаруживают не более 40% случаев мошенничества с SIM-картами, а доля ложных срабатываний при этом достигает около 30% [11]. Именно эти ограничения объясняют растущий интерес к искусственному интеллекту (ИИ) и машинному обучению (МО): такие модели способны находить скрытые закономерности в больших объёмах данных о вызовах (call detail records, CDR), сигнальном трафике (SS7/SIP) и поведении абонентов, опираясь на признаки активности SIM-карт, смены IMEI, географической мобильности и структуры международных соединений [5].

В настоящей статье ставится задача обобщить сложившиеся практики применения искусственного интеллекта для обнаружения и блокировки телекоммуникационного мошенничества, дать количественное сравнение их результативности по показателям accuracy, precision и recall, приводимых в опубликованных исследованиях, а также предложить обобщённую архитектуру интеллектуальной системы противодействия мошенничеству, применимую операторами связи на практике.

Анализ литературы и методология

Материал статьи носит характер аналитического обзора: изучен корпус научных публикаций и отраслевых отчётов 2021-2026 гг., посвящённых применению ИИ для обнаружения телекоммуникационного мошенничества; рассмотренные методы далее распределены по типу решаемой задачи, набору используемых признаков и заявленным метрикам качества (accuracy, precision, recall, F1-score). Источники отбирались по запросам «telecom fraud detection», «machine learning», «IRSF», «SIM-box», «graph neural network fraud», «SIM swap detection». В результате отбора сформировалось четыре смысловых блока: классические алгоритмы обучения с учителем, гибридные схемы «правила + ИИ», модели глубокого обучения на потоковых данных и графовые архитектуры.

Чаще всего в рассмотренных работах используется обучение с учителем на размеченных записях CDR. Так, для распознавания обхода через SIM-box (SIM-box bypass) применяются классификаторы, обучаемые на признаках вызовов фиксированного и мобильного операторов: пример подобной системы, реализованной на данных камерунского оператора связи, продемонстрировал, что анализ структуры международного трафика достаточно надёжно отличает легитимные соединения от обходных схем через SIM-фермы [5]. В свою очередь, Terzi, Sağiroğlu и Kılınç демонстрируют, что сочетание технологий больших данных с алгоритмами машинного обучения даёт возможность обрабатывать заметно большие объёмы телекоммуникационного трафика, чем традиционные системы обнаружения мошенничества, не жертвуя при этом скоростью анализа [6].

Самостоятельной сложностью выступает дисбаланс классов: доля мошеннических операций в общем трафике обычно не превышает 1-2%, что снижает способность стандартных классификаторов находить редкие случаи. Именно этому вопросу посвящена работа Krasić и Čelar, где для повышения полноты обнаружения на несбалансированных наборах данных предлагаются методы взвешивания классов и передискретизации [7]. Yehya и Salhab, в свою очередь, сравнивают несколько алгоритмов машинного обучения применительно к задаче обнаружения телекоммуникационного мошенничества и приходят к выводу, что ансамблевые методы - случайный лес и градиентный бустинг - устойчиво опережают по точности одиночные линейные модели [8].

Наглядные количественные оценки приводит исследование гибридной системы обнаружения SIM-мошенничества, где правило-ориентированная предфильтрация сочетается с классификатором Random Forest. На синтетическом наборе данных, воспроизводящем поведенческие характеристики абонентов (частота смены IMEI, географическая мобильность, интенсивность звонков и сообщений), система показала точность (accuracy) 92,3%, precision 90,7%, recall 89,5% и F1-меру 90,1%, тогда как дерево решений и метод опорных векторов (SVM), взятые отдельно, дали accuracy лишь 85,6% и 88,1% соответственно; сама по себе правило-ориентированная предфильтрация сократила долю пропущенных случаев мошенничества примерно на 18% [2]. При этом итоговая точность модели на базе Random Forest превышала 90% при времени обработки запроса менее 5 секунд, что говорит о пригодности гибридного подхода для систем, работающих в режиме реального времени [2].

Более сложный класс решений образуют архитектуры глубокого обучения на базе вариационных автокодировщиков и генеративно-состязательных сетей (VAE-GAN), которые обучаются на потоковых данных CDR, дополненных цифровым отпечатком устройства (device fingerprinting). По предварительным результатам, подобная архитектура поднимает precision примерно до 93%, а recall примерно до 90% относительно классических автокодировщиков, сохраняя при этом задержку обработки, приемлемую для биллинговых систем реального времени [10]. Существенным преимуществом такого подхода является устойчивость к «дрейфу концепта» (concept drift): модель непрерывно дообучается на новых потоках вызовов, что особенно важно при постоянно меняющейся тактике обходных схем SIM-box и IRSF.

Отдельную, четвёртую группу составляют модели, учитывающие сетевую (графовую) природу данных. Мошеннические SIM-карты и номера редко действуют поодиночке, обычно они образуют связанные группы за счёт общих устройств IMEI, мест совместной активации или повторяющихся шаблонов звонков, поэтому графовые модели машинного обучения (Graph ML / GNN), анализируя связи между узлами графа абонентов, способны обнаруживать целиком мошеннические сети, а не только отдельные подозрительные события. По оценкам отраслевых источников, такие графовые модели выявляют до 85% мошеннических сетей SIM-карт при доле ложных срабатываний около 5%, тогда как правило-ориентированные системы обнаруживают лишь порядка 40% случаев при FPR около 30% [11]. Для среднего по размеру оператора связи потери от SIM-мошенничества могут достигать 35 млн долларов в год, что делает переход от точечных правил к сетевому анализу связей экономически обоснованным решением [11].

В большинстве рассмотренных гибридных решений в роли базового классификатора используется случайный лес (Random Forest) — ансамбль деревьев решений, где итоговое предсказание формируется путём голосования отдельных деревьев; такая конструкция снижает риск переобучения и позволяет оценивать вклад каждого признака в результат [12]. Методологически работа сводится к сопоставлению показателей accuracy, precision, recall и F1-score, заявленных непосредственно в первоисточниках; собственных экспериментов не проводилось, так как работа имеет обзорный характер, а доступа к реальным наборам CDR-данных узбекистанских операторов связи на момент подготовки статьи не было.

Результаты

На основании проанализированных источников можно описать обобщённую архитектуру интеллектуальной системы противодействия телекоммуникационному мошенничеству; она включает четыре последовательных уровня.

1. Сбор данных. На этом уровне объединяются записи о вызовах (CDR), данные сигнализации (SS7/SIP), сведения о SIM-картах и устройствах (IMEI), а также поведенческие показатели абонента: как часто меняются устройства, в каких географических точках проявляется активность, каковы типичные шаблоны SMS-трафика.

2. Подготовка данных и построение признаков. Здесь рассчитываются агрегированные показатели: интенсивность международных вызовов за выбранный

интервал, распределение вызовов по направлениям, частота смены IMEI, отклонения в географической мобильности абонента.

3. Моделирование. На этом уровне применяются как самостоятельные алгоритмы (случайный лес, градиентный бустинг, автоэнкодеры для поиска аномалий), так и гибридные схемы, сочетающие правило-ориентированную фильтрацию с моделями машинного обучения, а также графовые нейронные сети, предназначенные для выявления скоординированных мошеннических групп.

4. Принятие решения и реагирование. Каждому событию присваивается показатель уверенности (confidence score); при высоком значении показателя событие либо блокируется автоматически, либо в режиме, близком к реальному времени, направляется в очередь на проверку специалисту службы фрод-мониторинга.

Ниже количественно охарактеризован масштаб проблемы: приведены данные о финансовых потерях по основным видам телекоммуникационного мошенничества, взятые из проанализированных отраслевых источников.

Совокупные глобальные потери отрасли в 2023 году составили 38,95 млрд \$, что соответствует 2,5% выручки отрасли [1]. По последней оценке совокупные глобальные потери достигли 41,82 млрд \$ [4]. Потери от IRSF (мошенничество с разделением доходов) составляют 6,23 млрд \$ в год [1]. Потери от PBX-мошенничества оцениваются в 4-7 млрд \$ в год [3]. Потери от wangiri (схема «одного гудка») составляют около 2,23 млрд \$ в год [3]. Потери от SIM-swap мошенничества превышают 2 млрд \$ в год, при этом их число выросло примерно на 400% с 2020 года [2]. Прогнозируемые потери от AI-роботизированных звонков к 2025 году оцениваются в 76 млрд \$ (глобально) [3]. Доля операторов с высоким уровнем CLI-спуфинга выросла до 55% в 2024 году (с 49% в 2023 году) [3]. Доля операторов с высоким уровнем спам-звонков выросла до 61% в 2024 году (с 46% в 2023 году) [3]. Среднегодовые потери среднего оператора связи от SIM-мошенничества достигают 35 млн \$ [11].

Далее сопоставлены показатели качества (accuracy, precision, recall), заявленные авторами рассмотренных исследований для различных алгоритмических подходов. Поскольку данные получены на разных выборках, строгая статистическая сопоставимость показателей не гарантируется, однако прослеживается устойчивая тенденция: по мере перехода от правило-ориентированных систем к моделям машинного и глубокого обучения, а затем к графовым архитектурам, полнота обнаружения мошенничества последовательно растёт, а доля ложных срабатываний снижается.

Правило-ориентированная система демонстрирует полноту обнаружения (recall) около 40% при доле ложных срабатываний (FPR) около 30% [11]. Дерево решений (Decision Tree) показывает accuracy 85,6% [2]. Метод опорных векторов (SVM) показывает accuracy 88,1% [2]. Гибридная модель Random Forest (правила + ML) демонстрирует accuracy 92,3%, precision 90,7% и recall 89,5% ($F1 = 90,1\%$) [2]. Архитектура VAE-GAN (dual-encoder, потоковые данные) достигает precision около 93% и recall около 90% [10]. Графовые модели Graph ML / GNN (сети SIM-карт) показывают recall около 85% при FPR около 5% [11].

Сопоставление показателей демонстрирует, что замена правило-ориентированного подхода (recall около 40%, FPR около 30%) гибридной моделью

Random Forest даёт прирост полноты обнаружения более чем на 49 процентных пунктов, а доля ложных срабатываний при этом сокращается в 3-6 раз в зависимости от конкретной реализации. Лучшее сочетание точности и полноты среди рассмотренных подходов показывают графовые модели (recall около 85% при FPR около 5%) и гибридные архитектуры VAE-GAN (precision около 93%, recall около 90%) это говорит в пользу совместного применения сетевого анализа связей между абонентами и моделей глубокого обучения на потоковых данных.

Отдельным результатом анализа стала оценка экономической целесообразности внедрения ИИ-решений. Поскольку ежегодный ущерб отрасли от IRSF составляет 6,23 млрд долларов [1], а переход от правило-ориентированного подхода (recall 40%) к графовым моделям (recall 85%) позволяет сократить долю необнаруженных мошеннических сетей SIM-карт с 60% до 15% [11], можно условно оценить потенциальный эффект от масштабного внедрения подобных моделей в миллиарды долларов ежегодно предотвращённых потерь даже с учётом расходов на разработку, обучение моделей и построение инфраструктуры обработки данных в реальном времени.

Обсуждение

Несмотря на существенный прирост эффективности, который даёт применение искусственного интеллекта в выявлении телекоммуникационного мошенничества, этот подход сопряжён с рядом методологических и практических ограничений. Первое из них устойчивый дисбаланс классов: подтверждённые случаи мошенничества обычно составляют не более 1-2% обучающей выборки, что вынуждает прибегать к специальным техникам взвешиванию ошибок, синтетической передискретизации, обучению с учётом стоимости ошибки [7].

Второе ограничение связано с тем, что мошеннические схемы постоянно подстраиваются под внедряемые средства защиты: заметный годовой рост доли операторов, фиксирующих высокий уровень CLI-спуфинга (с 49% до 55%) и спама-трафика (с 46% до 61%), подтверждает, что статичные модели быстро теряют актуальность и требуют регулярного дообучения вместе с мониторингом дрейфа данных [3]. Архитектуры типа VAE-GAN частично снимают эту проблему за счёт постоянной адаптации к потоковым данным, но для этого им требуются существенные вычислительные мощности и развитая инфраструктура потоковой обработки [10].

Третья проблема касается интерпретируемости сложных моделей - глубоких сетей и графовых архитектур: специалистам по управлению мошенничеством нужно понятное обоснование каждого решения о блокировке, а это стимулирует развитие объяснимого искусственного интеллекта (explainable AI) применительно к данной области. Показательно, что гибридные решения, объединяющие прозрачные правила с моделью Random Forest, изначально закладываются с учётом требования объяснимости и визуализируют значимость признаков для каждого конкретного решения, облегчая работу аналитиков и укрепляя доверие к системе [2].

Наконец, нельзя игнорировать регуляторные и этические аспекты внедрения подобных систем: обработка данных о местоположении и поведении абонентов обязана соответствовать законодательству о защите персональных данных, а ошибочная блокировка добросовестного абонента способна одновременно нанести репутационный

и финансовый ущерб и клиенту и оператору. По этой причине на практике события с высоким риском чаще передаются на приоритетную проверку специалисту, а не блокируются автоматически, полностью автоматическое блокирование сохраняется только для случаев с максимальным показателем уверенности модели.

Заключение

Представленный анализ показывает, что методы искусственного интеллекта и машинного обучения существенно эффективнее классических правило-ориентированных систем при выявлении и блокировке телекоммуникационного мошенничества. При общеотраслевых потерях порядка 39-42 млрд долларов США в год [1, 4] переход от правило-ориентированных систем (recall около 40%, FPR около 30%) к гибридным моделям на базе Random Forest (accuracy 92,3%, precision 90,7%, recall 89,5%) и к графовым нейронным сетям (recall до 85% при FPR около 5%) даёт кратный рост эффективности обнаружения и одновременно снижает нагрузку на службы фрод-мониторинга за счёт уменьшения числа ложных срабатываний.

В числе перспективных направлений будущих исследований можно назвать: методы работы с сильно несбалансированными данными без потери качества классификации; объяснимые модели, обеспечивающие прозрачность принимаемых решений для служб фрод-мониторинга и регуляторов; федеративное обучение, дающее операторам связи возможность обмениваться выявленными шаблонами мошенничества без передачи персональных данных абонентов; дальнейшее развитие графовых и гибридных VAE-GAN архитектур для одновременного выявления как координированных мошеннических сетей, так и единичных аномальных событий; а также построение потоковых архитектур, способных распознавать мошенничество в реальном времени с минимальной задержкой обработки запроса.

Список литературы:

1. Communications Fraud Control Association. Global Fraud Loss Survey 2023. CFCA, 2024. URL: <https://cfca.org/telecommunications-fraud-increased-12-in-2023-equating-to-an-estimated-38-95-billion-lost-to-fraud/>
2. Koli D. H., Limbhare S. B., Nikum M. V. AI-Driven SIM Card Fraud Detection System // International Journal of Advanced Research in Computer and Communication Engineering (IJARCCE). 2025. Vol. 14, No. 10. P. 216-221. DOI: 10.17148/IJARCCE.2025.141033
3. National Anti-Fraud Platform. Technology for Secure Collaboration & Compliance: GLF Fraud Report 2024 statistics. AB Handshake, 2025. URL: <https://abhandshake.com/community/telecom-fraud-prevention-solution/>
4. TheStreet. The \$41 Billion Telecom Fraud Secret. TheStreet, 2026. URL: <https://www.thestreet.com/technology/the-41-billion-telecom-fraud-secret>
5. Djomadji E. M. D., Basile K. I., Christian T. T., Djoko F. V. K., Sone M. E. Machine Learning-Based Approach for Identification of SIM Box Bypass Fraud in a Telecom Network Based on CDR Analysis: Case of a Fixed and Mobile Operator in Cameroon // Journal of Computer and Communications. 2023. Vol. 11, No. 2. DOI: 10.4236/jcc.2023.112010
6. Terzi D. S., Sağiroğlu Ş., Kılınç H. Telecom Fraud Detection with Big Data Analytics // International Journal of Data Science. 2021. Vol. 6, No. 3. P. 191. DOI: 10.1504/ijds.2021.121090

7. Krsić I., Čelar S. Telecom Fraud Detection with Machine Learning on Imbalanced Dataset // 2022 International Conference on Software, Telecommunications and Computer Networks (SoftCOM). IEEE, 2022. P. 1-6.
8. Yehya B. A., Salhab N. Telecommunications Fraud Machine Learning-based Detection // 2023 4th International Conference on Data Analytics for Business and Industry (ICDABI). IEEE, 2023. P. 656-661. DOI: 10.1109/ICDABI60145.2023.10629612
9. Subex. Telecom Fraud in 2026: Types, Emerging Risks & How AI-First Prevention Stops Revenue Leakage. Subex, 2026. URL: <https://www.subex.com/article/telecom-fraud-in-2026-types-emerging-risks-how-ai-first-prevention-stops-revenue-leakage/>
10. Banu V. Real-Time Fraud Detection in Telecom Charging Systems Using AI // International Journal of Emerging Trends in Computer Science and Information Technology (IJETCSIT). 2025. P. 571-582. DOI: 10.56472/ICCSAIML25-163
11. Kumo AI. SIM Fraud Detection: Telecom AI Use Case. Kumo AI, 2025. URL: <https://kumo.ai/solutions/industries/telecom/fraud-detection/>
12. Breiman L. Random Forests // Machine Learning. 2001. Vol. 45, No. 1. P. 5-32. DOI: 10.1023/A:1010933404324

