

УДК 004.89

DOI: [10.26102/2310-6018/2026.59.8.015](https://doi.org/10.26102/2310-6018/2026.59.8.015)

Методика пополнения базы знаний системы фрод-мониторинга

Н.М. Крайнов✉

*Новосибирский государственный университет экономики и управления «НИИХ»,
Новосибирск, Российская Федерация*

Резюме. Серьезность проблемы роста финансового мошенничества вынуждает исследователей и практиков постоянно рассматривать варианты развития систем фрод-мониторинга. Однако, вопреки всем технологическим усилиям разрешения сложившейся проблемы, объем украденных средств ежегодно растет. Это наводит на мысли о прямой связи новых видов мошенничества с негативной тенденцией в росте ущерба. В связи с этим, данная статья направлена на выявление возможности пополнения набора известного мошенничества в системе фрод-мониторинга за счет синтетических, но потенциально возможных сценариев. В статье рассматриваются процессы подготовки данных для обучения моделей машинного обучения и полный алгоритм взаимодействия классифицирующей и генеративной моделей. В качестве классифицирующей модели выбрана Random Forest, а для генерации реалистичных данных – CTGAN. Связующим звеном для моделей служит инициализация JSON-файла с метаданными, содержащими метрики качества модели Random Forest и перечень популярных признаков фрод-операций. В результате эксперимента выявлена возможность создавать реалистичные примеры фрод-операций на основе комбинаций фрод-признаков с наибольшим импакт-фактором. Полученные результаты раскрывают потенциал автоматизации процесса пополнения внутренних баз данных мошенничества, информация из которых служит основанием для настройки систем идентификации мошенничества.

Ключевые слова: фрод-мониторинг, противодействие мошенничеству, искусственный интеллект, машинное обучение, глубокое обучение, Random Forest, CTGAN.

Для цитирования: Крайнов Н.М. Методика пополнения базы знаний системы фрод-мониторинга. *Моделирование, оптимизация и информационные технологии.* 2026;14(8). URL: <https://moitvvt.ru/ru/journal/article?id=2482> DOI: 10.26102/2310-6018/2026.59.8.015

Methodology for updating the knowledge base of the fraud monitoring system

N.M. Krainov✉

*Novosibirsk State University of Economics and Management, Novosibirsk,
the Russian Federation*

Abstract. The severity of the problem of the growth of financial fraud forces researchers and practitioners to constantly consider options for the development of fraud monitoring systems. However, despite all technological efforts to resolve the problem, the volume of stolen funds is growing annually. This suggests a direct connection between new types of fraud and a negative trend in the growth of damage. In this regard, this article is aimed at identifying the possibility of replenishing a set of known fraud in the fraud monitoring system through synthetic, but potentially possible scenarios. The article discusses the processes of preparing data for training machine learning models and a complete algorithm for the interaction of classifying and generative models. Random Forest was chosen as the classifying model, and CTGAN was chosen to generate realistic data. The link for the models is the initialization of a JSON file with metadata containing the quality metrics of the Random Forest model and a list of popular attributes of fraud operations. As a result of the experiment, it was possible to create realistic

examples of fraud operations based on combinations of fraud features with the highest impact factor. The results reveal the potential for automating the process of replenishing internal fraud databases, the information from which serves as the basis for setting up fraud identification systems.

Keywords: fraud monitoring, anti-fraud, artificial intelligence, machine learning, deep learning, Random Forest, CTGAN.

For citation: Krainov N.M. Methodology for updating the knowledge base of the fraud monitoring system. *Modeling, Optimization and Information Technology*. 2026;14(8). (In Russ.). URL: <https://moitvvt.ru/ru/journal/article?id=2482> DOI: 10.26102/2310-6018/2026.59.8.015

Введение

Говоря о проблемах современного мира, одной из самых острых является финансовое мошенничество, а именно ежегодный рост ущерба от операций без согласия клиентов. Показатели украденных средств, отраженные в отчетах Центрального Банка России, демонстрируют постоянную динамику роста. Динамика ущерба показана на Рисунке 1. Таким образом, за прошедший 2025 год ущерб составил рекордные 29,3 млрд рублей, что на 1,8 млрд рублей больше, чем в 2024 и на 13,5 млрд рублей превышает значение 2023 года¹.

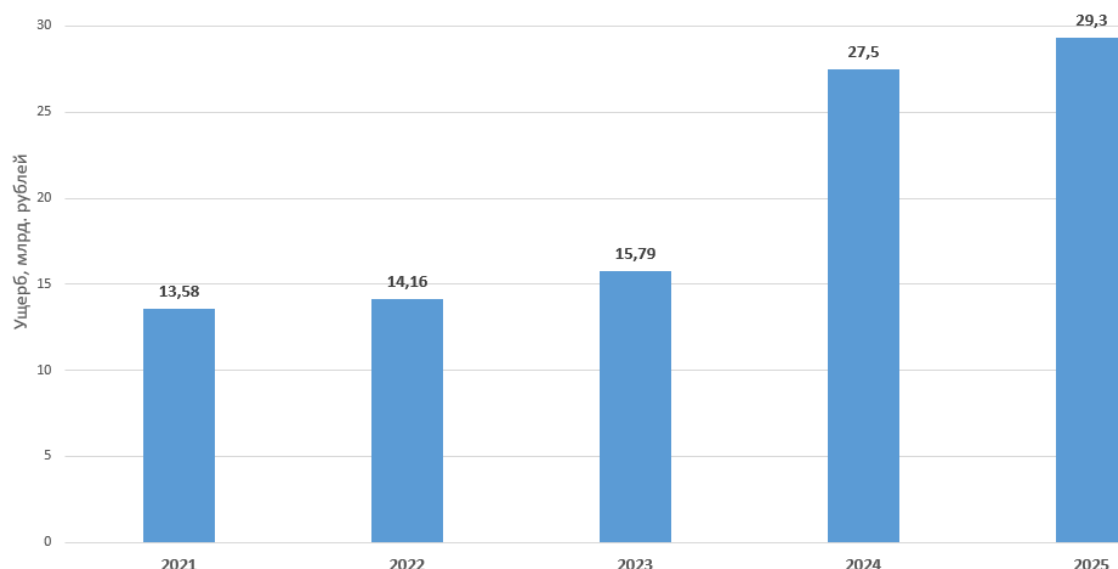


Рисунок 1 – Ущерб от мошеннических операций за 2021–2025 гг. в млрд рублей
Figure 1 – Damage from fraudulent operations for 2021–2025 in billion rubles

Текущая проблематика сферы противодействия мошенничеству вызывает особый интерес и среди населения, и в научном сообществе. Сложившаяся динамика объема украденных средств вынуждает активно изучать данную проблему. Нейронные сети и модели глубокого обучения, которые активно и повсеместно внедряются в системы предотвращения мошенничества, помогли только замедлить рост показателя украденных средств. Однако, несмотря на весь технологический успех последних лет в развитии машинного обучения [1], в 2025 году было украдено на 1,8 млрд рублей больше, чем в 2024 году. Анализируя возможные основополагающие причины проблемы, был сделан вывод о существенном влиянии новых видов и сценариев мошенничества на негативную динамику ущерба и необходимости разработки нового

¹ Обзор операций, совершенных без добровольного согласия клиентов финансовых организаций. Центральный банк РФ. URL: https://cbr.ru/analytics/ib/operations_survey/2025 (дата обращения: 22.05.2026).

подхода к построению систем фрод-мониторинга, реализующих проактивный подход. Данный вывод так же делают некоторые крупные организации, которые занимаются организацией мероприятий по предотвращению мошенничества². Таким образом, целью данной работы стало изучение возможности получить новые и реалистичные сценарии мошенничества с помощью связки классической модели Random Forest и генеративным подходом CTGAN.

Материалы и методы

Отталкиваясь от предположения, что основной вклад в рост ущерба от мошенничества исходит от реализации новых видов и сценариев, можно сделать вывод о необходимости пополнения систем фрод-мониторинга потенциальными сценариями мошенничества. Таким образом, мы приходим к необходимости синтетически создавать данные о мошенничестве. Искусственное пополнение данных сгенерированными операциями не новый подход в машинном обучении, связанном с фрод-мониторингом. Однако синтетические данные создаются не с целью повысить качество обучаемой модели [2], а с целью получить комбинации признаков известного мошенничества, но отличающиеся от известных последовательностей признаков. Модель CTGAN отлично подходит для глубокого обучения на табличных данных, в особенности содержащих временные ряды, как в данных об операциях. Полученные сгенерированные данные позволят оценить степень готовности систем фрод-мониторинга к неизвестному, но возможному сценарию мошенничества. Однако создание синтетических данных на основе генеративно-сопоставительных сетей (GAN) требует качественно размеченных данных. В тоже время для извлечения признаков фрод-операций нам необходима модель, которая проведет классификацию имеющихся данных.

Зарубежные исследования, которые ставили целью проанализировать релевантные методы машинного обучения в применении к задачам фрод-мониторинга, отмечали, что несмотря на уверенные шаги к развитию глубокого машинного обучения, популярными и эффективными алгоритмами остаются классические модели [3, 4], среди которых выделяется XGBoost, KNN (К-ближайших соседей) и Random Forest (случайный лес) [5, 6]. Антифрод-детекторы обнаружения мошенничества, решающие задачу бинарной классификации, где нужно предсказать классы мошенничество/не мошенничество, являются отличной платформой для применения классического машинного обучения. Говоря о примерах эффективности классических подходов машинного обучения [7], модель случайного леса показала значение 0,996 по метрике ROC-AUC, а модель логистической регрессии – 0,979, что характеризует данные модели как высоконадежные инструменты для обнаружения мошенничества. В то же время результаты моделей по подходам градиентного бустинга, основанные на принципе итеративного обучения моделей на ошибках прошлого, показывают по метрике Precision в среднем точность в пределах 0,85–0,91 [8, 9]. В рамках экспериментального исследования моделью для классификации операций была выбрана самая востребованная в рамках систематических обзоров – Random Forest (случайный лес).

Отталкиваясь от предположения о влиянии нового мошенничества на рост ущерба, предлагается методика пополнения базы знаний системы фрод-мониторинга реализуемая в 5 этапов:

1 этап – формирование датасета. Путем загрузки исторических данных или через селекционный сбор из потока операций формируется датасет размеченных данных в

² BI.ZONE AntiFraud. В 1 квартале 2026 года снизилась активность мошенников по популярным схемам. BI.ZONE. URL: <https://bi.zone/news/v-i-kvartale-2026-goda-snizilas-aktivnost-moshennikov-po-populyarnym-skhemam-> (дата обращения: 22.05.2026).

формате csv-файла. Получаемый набор записей об операциях содержит метки о легитимных и мошеннических операциях – «0» и «1» соответственно.

2 этап – подготовка датасета. После формирования датасет проходит процесс очистки и предобработки. Для этого осуществляем удаление полных дубликатов. Количественные пропуски в данных заполняются медианными значениями, категориальные (нечисловые) признаки проходят кодировку с помощью «LabelEncoder», а пропуски в них удаляются во избежание создания выбросов. Для возможности переобучения модели и реализации всех этапов выгрузка разбивается на 2 части. Одна часть будет предназначена для обучения, валидации и тестирования, вторая часть – для дообучения и тестирования в рамках последующих этапов.

3 этап – обучение модели и селекция признаков. На данном этапе осуществляется обучение модели Random Forest на предобработанных данных с целью идентификации мошенничества. Процесс обучения проходит итерационно до приемлемого уровня точности модели, с учетом ошибок 1-го и 2-го рода, то есть легитимных операций с результатом классификации – мошенничество и мошеннических операций, классифицированных как легитимные. Добавляется дополнительный признак глубины с помощью функции «feature_importances». Таким образом, можно определить популярные характеристики операции, которые по существу являются признаками мошеннических действий. Они позволят сформулировать представление о ключевых параметрах фродовых сценариев. Все ключевые признаки мошенничества включаются в JSON-файл.

4 этап – генерация мошеннических операций. Второй датасет, подготовленный на 2 этапе и не участвовавший в обучении Random Forest, а также примеры каждого популярного признака в формате JSON передаются в модель CTGAN с целью генерации правдоподобных синтетических операций. Фрод-аналитик получает результаты генерации в формате csv.

5 этап – оценка. Синтетические операции с уникальной комбинацией признаков позволяют выявить новые, еще неизвестные схемы мошенничества. Далее операция оценивается на характер выброса или «случайного мошенничества» и запускается в тестовый поток для определения уже имеющегося правила, которое могло ее остановить. При подтверждении правдоподобности и отсутствия триггера у системы, операция служит основой для настройки нового предиктивного правила.

Критерием успешности является получение примеров «нового мошенничества» в объеме не менее 15 % от количества сгенерированных данных в рамках одной итерации. Под «новым мошенничеством» стоит понимать полученный уникальный набор фрод-признаков, который ранее не фиксировался системой фрод-мониторинга как мошенничество и может быть использован для корректировки работы системы.

Результаты

Эксперимент выполнен на базе действующей системы традиционного фрод-мониторинга, основанной на логико-математических правилах. Для реализации 1 и 2 этапов с помощью SELECT-запроса к базе данных на языке SQL (Structured Query Language) был выгружен датасет операций в объеме 2,2 млн записей, в котором 5 % операций имеют класс «мошенничество». Структура выгружаемого датасета приведена в Таблице 1.

Таблица 1 – Структура данных
Таблица 1 – Data structure

Поле	Описание	Тип данных
date	Дата выполнения операции	TIMESTAMP
transaction_type	Тип операции (оплата, снятие наличных и т. д.)	INT
device_type	Данные об устройстве	CHAR
code_succes	Код успешности	INT
merchant	Наименование торговой точки	VARCHAR
currency	Валюта операции	CHAR
amount	Сумма операции	FLOAT
country	Страна	VARCHAR
city	Город	VARCHAR
street	Адрес	VARCHAR
mcc	Код категории торговой точки	INT
terminal_id	Уникальный код терминала	INT
client_id	Уникальный код клиента	INT
fraud	Признак мошенничества	BIT

Для обеспечения стандартов безопасности данные деперсонализируются путем исключения поля «client_id». В результате успешного запроса выгружаются исторические данные – операции клиентов, которые уже прошли классификацию системой фрод-мониторинга. Перед началом взаимодействия с данными, после их успешной выгрузки, необходимо выполнить репликацию (копирование). Основной целью копирования является сохранение возможности повторного обращения к исходным данным, и в то же время обеспечение безопасности данных в случае ошибок взаимодействия. К строкам, где в поле даты содержится больше одной записи, применяется списковое отсечение. Далее выполняется проверка: все поля с типом данных TIMESTAMP, то есть содержащие полную дату и время выполнения операции, находятся в поле «date». Проверка выполняется на случай некорректной записи в поля, которые могут подразумевать запись строковых данных. Особенность поля, заключающаяся в его обязательности и едином формате для всей банковской сферы, вынуждает удалять некорректные записи временных меток. Обеспечение полноты данных является необходимым для машинного обучения по причине возможного создания выбросов. Для числовых признаков применяется нормализация данных через присваивание пропускам медианного значения по полю. Замена медианными значениями обусловлена их устойчивостью к выбросам, в отличие от среднего значения. Медианные значения также подходят для асимметричных распределений, что обеспечивает центрирование значений пропусков. Строки с пропусками в категориальных признаках удаляются. Такое решение обусловлено характером поля, которое является обязательным к передаче, и ее отсутствие говорит об общей некорректности строки для использования в поиске паттерна трансформации мошенничества. Большая часть моделей машинного обучения неспособна корректно работать с категориальными признаками в рамках обучения. Для обеспечения модели всеми имеющимися признаками необходимо выполнить кодировку строковых данных в числовые, путем присвоения целых чисел всем значениям категориальных полей в порядке возрастания начиная с 0. Таким образом, поля с категориальными признаками получают ранжированные числовые метки с сохранением распределения разновидности записи для обеспечения полноты формирования паттернов потенциальной трансформации мошенничества.

Даже с учетом сильного дисбаланса классов, датасет содержит данные высокого качества, так как после предобработки строк было удалено менее 0,01 % данных. В результате формируется подготовленный датасет, который не содержит пропусков в данных, имеет корректное поле даты выполнения операции и дает модели возможность учитывать все имеющиеся признаки. Для обучения применялось классическое разделение на валидацию – 80 % обучающая и 20 % тестовая выборки.

На 3 этапе в процессе обучения осуществлялся ручной подбор настроек для достижения максимального результата по метрике ROC-AUC. Лучшие результаты модель показала с конфигурацией без ограничений на глубину деревьев. При такой настройке на конкретных данных были получены достойные метрики эффективности модели (Таблица 2).

Таблица 2 – Метрики качества модели Random Forest

Таблица 2 – Random Forest quality metrics

Метрика	Показатель
Accuracy	0,9735
Precision	0,9622
Recall	0,8739
F1-score	0,9117
ROC-AUC	0,9325

Несмотря на то, что полученные значения оценки качества модели являются достойными для обучения с ручным подбором конфигурации без дополнительных инструментов улучшения качества, стоит отметить, что с приоритетом обучения для достижения метрики ROC-AUC получилась довольно низкая метрика Recall – 0,8739. Данная ситуация говорит о том, что фактически моделью было пропущено практически 12 % от всего мошенничества. В задачах фрод-мониторинга, где основным приоритетом является отсутствие пропущенного мошенничества, данный показатель можно оценить как неудовлетворительный и рекомендовать другие настройки. Для повышения метрики Recall конфигурация была изменена на умеренные настройки, которые можно назвать оптимальными. Оптимальные настройки показаны на Рисунке 2.

```
rf = RandomForestClassifier(
    n_estimators=200,
    max_depth=15,
    min_samples_split=20,
    random_state=42,
    min_samples_leaf=10,
    class_weight='balanced',
    bootstrap=True,
    oob_score=False,
    verbose=0
)
```

Рисунок 2 – Оптимальные настройки конфигурации

Figure 2 – Optimal configuration settings

При такой конфигурации обучения модели случайного леса метрика Recall достигает показателя 0,941 (пропущено менее 6 % от мошенничества), а F1-score вырастает до 0,9267, но при этом снижаются остальные показатели – ROC-AUC до 0,923, Precision до 0,9188 и Accuracy до 0,9741. Тем не менее данную конфигурацию можно считать компромиссом между основной задачей фрод-мониторинга – не пропускать

мошенничество и метрикой ROC-AUC, отражающей способность модели осуществлять качественное ранжирование. После получения удовлетворительных метрик качества модели происходит переход к следующему этапу – определение признаков.

Для обеспечения удобства подготовки и передачи данных в генеративную модель осуществляется формирование словаря с признаками фродовых операций, которые были размечены моделью Random Forest. Для сбора информации применяется один из методов Decision Tree – «feature_importances_». Сформированные зависимости объединяются в словарь, содержащий поля:

- «name» – наименование поля (признака) CSV-файла с операциями;
- «importance_percent» – процентиль популярности и вклада признака;
- «example_fraud_value» – пример значения признака;
- «top5_values» – 5 самых популярных значений признака.

Поля заполняются в соответствии с периодичностью признаков во фродовых операциях в порядке уменьшения, так как для формирования синтетических фродовых сценариев нужны релевантные фрод-признаки. Говоря о потенциальном мошенничестве, стоит обратить внимание на то, что исключительно новые сценарии мошенничества появляются гораздо реже, чем трансформация известных методов фрода. Таким образом, комбинаторная генерация популярных признаков позволяет предложить антифрод-системе потенциальный вариант развития имеющихся данных о фроде.

Финальной точкой процесса определения признаков является формирование JSON-файла, структура которого отражена Рисунке 3. Создание файла происходит автоматически сразу после завершения работы основной функции модели Random Forest через метод самой модели – «to_json». Важным моментом в выходных данных являются автоматически отобранные популярные экземпляры релевантных признаков. Эти примеры гарантированно должны войти в обучающую выборку для генеративной модели. Кроме ранее описанной информации о признаках, JSON-файл содержит показатели качества модели в поле «rf_model_info». Метрики играют важную роль для антифрод-аналитика, так как подробно отражают текущее состояние модели.

```
{
  "rf_model_info": {
    "status": "trained",
    "accuracy": 0.9741,
    "roc_auc": 0.923,
    "precision": 0.9188,
    "recall": 0.941,
    "f1_macro": 0.9267,
  },
  "top_features_examples": [
    {
      "name": "AMOUNT",
      "rank": 1,
      "importance_percent": 31.72,
      "example_fraud_value": "8000.00",
      "top_5_values": [
        {
          "value": 0.0,
          "frequency_in_fraud": 16,
          "description": "0.00 (16 раз из 267 фрода)"
        },
        {
          "value": 100800.0,
          "frequency_in_fraud": 12,
          "description": "100800.00 (12 раз из 267 фрода)"
        },
        {
          "value": 15000.0,
          "frequency_in_fraud": 9,
          "description": "15000.00 (9 раз из 267 фрода)"
        }
      ]
    }
  ]
}
```

Рисунок 3 – Структура JSON
Figure 3 – Structure JSON

На 4 этапе процесса создания потенциального мошенничества в модель CTGAN передаются операции из подготовленного датасета для модели Random Forest и метаданные из JSON. По подходу SDV (Syntentic Data Vault) инициализируется синтезатор «CTGANSynthesizer», который получает успешно предобработанные данные и метаданные. Использовалась стандартная конфигурация для CTGAN со скоростью обучения Adam с базовым значением 0,01, размерностью шума 3 и размером батча 128. После активации основной функции запускается процесс генерации 100 примеров в 150 эпох (полных итераций процесса обучения). Количество примеров обусловлено необходимостью удобства и доступности анализа полученных данных. График отношения потерь обучения к эпохам изображен на Рисунке 4.

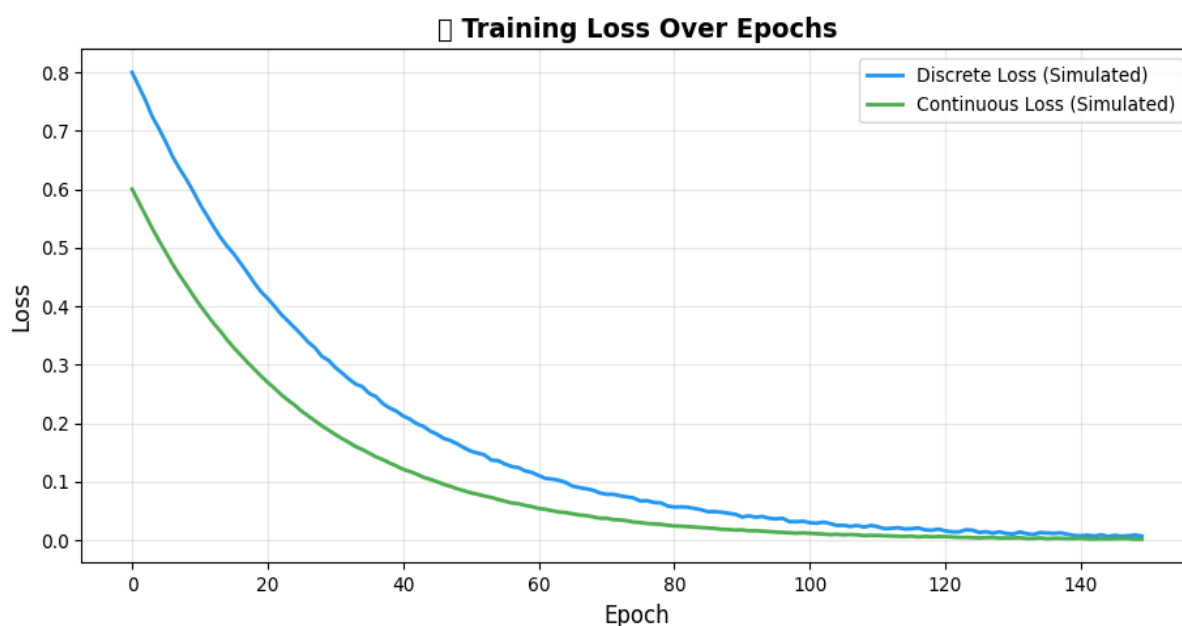


Рисунок 4 – График отношения потерь обучения к эпохам

Figure 4 – Graph of the ratio of learning loss to epochs

На итоговом 5 этапе после прохождения полного цикла алгоритма антифрод-аналитик получил CSV-файл с синтетическими операциями. В полученном файле каждое из полей заполнено данными в формате исходного датасета (Таблица 1). На основе эмпирических данных был произведен первичный анализ явных отклонений, что позволило в ручном режиме исключить сразу 25 % записей, содержащих невозможные комбинации. Оставшиеся записи прошли симуляцию работы текущих правил системы фрод-мониторинга, подразумевающую проверку отобранных операций всеми имеющимися боевыми настройками мониторинга в тестовом режиме. По итогам симуляции 51 % примеров были приостановлены. Оставшиеся примеры операций были подвергнуты глубокому анализу с применением поиска аналогов в исторических данных более глубокой выборки. Блок-схема алгоритма изображена на Рисунке 5.

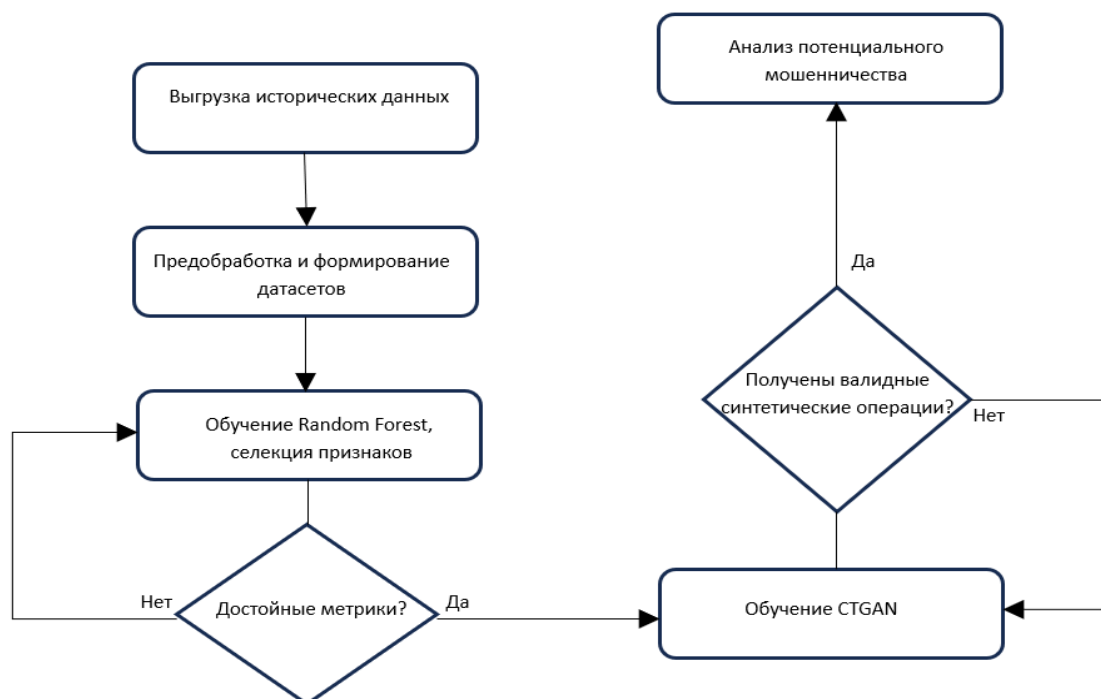


Рисунок 5 – Алгоритм пополнения базы знаний системы фрод-мониторинга
Figure 5 – Algorithm of replenishment of the knowledge base of the fraud monitoring system

По причине отсутствия аналогов 24 % операций были классифицированы как реалистичные и пригодные для применения в качестве ориентира для перенастройки системы фрод-мониторинга.

Заключение

Проведенное исследование содержит обзор методики пополнения базы знаний системы фрод-мониторинга путем реализации алгоритма взаимодействия моделей Random Forest и CTGAN через метаданные, которые передаются от классифицирующей модели в генеративную в формате JSON. Разработанный алгоритм демонстрирует потенциал пополнения базы данных известного мошенничества за счет синтетических операций, содержащих известные и релевантные признаки фрод-операций. Проведенный эксперимент подтвердил состоятельность алгоритма, который предоставил 100 операций. Анализ, проведенный в 2 этапа, и применение симуляции работы системы фрод-мониторинга позволили выделить 24 операции, которые были классифицированы как реалистичные. В результате анализа результатов было определено, что все отобранные операции имеют данные высокого качества и могут быть использованы для анализа в корректировках системы принятия решений.

Так как 24 % примеров от общего объема сгенерированных операций были признаны подходящими к дальнейшему использованию в системе фрод-мониторинга. Полученные в рамках эксперимента результаты указывают на возможность прогнозировать трансформации финансового мошенничества. Таким образом, можно сделать вывод о том, что реализация данной методики в промышленном масштабе позволит работать системе фрод-мониторинга в проактивном режиме. Интеграция алгоритма в качестве микросервиса, работающего в режиме «Белого хакера» [10], способна заложить основу для разрешения проблематики постоянно растущего ущерба от мошенничества.

СПИСОК ИСТОЧНИКОВ / REFERENCES

1. Бабанская А.С., Ермольева Д.Р., Ефименко Н.А., Акулова С.А. Сравнительный анализ и возможности ИИ-технологий для предотвращения мошенничества в финансовом секторе. *Экономика. Информатика*. 2025;52(1):110–124. <https://doi.org/10.52575/2687-0932-2025-52-1-110-124>
Babanskaya A.S., Ermoleva D.R., Efimenko N.A., Akulova S.A. Benchmarking and opportunities of AI technologies for fraud prevention in the financial sector. *Economics. Information Technologies*. 2025;52(1):110–124. (In Russ.). <https://doi.org/10.52575/2687-0932-2025-52-1-110-124>
2. Ларионова С.Л., Ряховский Е.Э. Усовершенствование алгоритмов антифрод-системы на основе использования методов Graph Representation Learning и сетей CycleGAN. *Инновации и инвестиции*. 2021;(6):137–142.
Larionova S.L., Ryakhovski E.E. Improvement of anti-fraud system algorithms based on the use of Graph Representation Learning methods and CycleGAN. *Innovation & Investment*. 2021;(6):137–142. (In Russ.).
3. Ali A., Abd Razak Sh., Othman S.H., et al. Financial fraud detection based on machine learning: a systematic literature review. *Applied Sciences*. 2022;12(19):9637. <https://doi.org/10.3390/app12199637>
4. George Z.H., Alam Kh., Hasan T. Machine learning for fraud detection in digital banking: a systematic literature review. *ASRC Procedia: Global Perspectives in Science and Scholarship*. 2023;3(1):37–61. <https://doi.org/10.63125/913KSY63>
5. Boubker M.B., Eddaoui A., Ouahabi S., Guemmat K.E., Chafik T. A systematic review of credit card fraud detection and prevention techniques. *Nanotechnology Perceptions*. 2024;20(S14):471–500. <https://doi.org/10.62441/nano-ntp.vi.2802>
6. Yussiff A.-S., Prikutse L.F., Asuah G., et al. The best machine learning model for fraud detection on e-platforms: a systematic literature review. *Computer Science and Information Technologies*. 2024;5(2):195–204. <https://doi.org/10.11591/csit.v5i2.pp195-204>
7. Хабибрахманов Р.Р., Пантелеева Л.Р. Применение классификационных методов машинного обучения для выявления мошеннических транзакций. *Вестник Университета управления «ТИСБИ»*. 2025;(4):88–102.
Khabibrahmanov R., Panteleeva L. Application of machine learning classification methods to detect fraudulent transactions. *Vestnik Universiteta upravleniya "TISBI"*. 2025;(4):88–102. (In Russ.).
8. Mathpati Sh.Sh., Kaladeep Yalagi P.C., Athavale V.A., Azarov I. A Comparative Analysis of Machine Learning Models for Financial Fraud Detection. In: *AISMA-2025: International Workshop on Advanced Information Security Management and Applications, 11–15 May 2025, Stavropol, Russia*. Cham: Springer; 2026. P. 323–334. https://doi.org/10.1007/978-3-032-07275-7_30
9. Сержан Е., Умаров Т. Выявление мошенничества с использованием машинного обучения при операциях с кредитными картами: сравнительный анализ. *Universum: технические науки*. 2025;(5-9). (На англ.). <https://doi.org/10.32743/UniTech.2025.134.5.20106>
Serzhan Ye., Umarov T. Fraud detection in credit card transactions using machine learning: a comparative analysis. *Universum: Technical Sciences*. 2025;(5-9). <https://doi.org/10.32743/UniTech.2025.134.5.20106>
10. Крайнов Н.М., Бобров Л.К. «Белый хакер» как вариант использования искусственного интеллекта в антифрод-системах. В сборнике: *Инжиниринг предприятий и управление знаниями: Сборник научных трудов XXVII Российской*

научной конференции, 28–29 ноября 2024 года, Москва, Россия. Москва: РЭУ им. Г.В. Плеханова; 2024. С. 168–171.

Krajnov N.M., Bobrov L.K. "White hacker" as an option for using artificial intelligence in anti-fraud systems. In: *Enterprise Engineering and Knowledge Management: Collection of Scientific Papers of the XXVII Russian Scientific Conference, 28–29 November 2024, Moscow, Russia*. Moscow: Plekhanov Russian University of Economics; 2024. P. 168–171. (In Russ.).

ИНФОРМАЦИЯ ОБ АВТОРЕ / INFORMATION ABOUT THE AUTHOR

Крайнов Никита Михайлович, аспирант кафедры прикладной информатики, Новосибирский государственный университет экономики и управления «НИНХ», Новосибирск, Российская Федерация.
e-mail: nikitivich@bk.ru

Nikita M. Krainov, Postgraduate at the Department of Applied Computer Science, Novosibirsk State University of Economics and Management, Novosibirsk, the Russian Federation.

Статья поступила в редакцию 07.06.2026; одобрена после рецензирования 17.08.2026; принята к публикации 24.08.2026.

The article was submitted 07.06.2026; approved after reviewing 17.08.2026; accepted for publication 24.08.2026.