

ЭКОНОМИЧЕСКИЕ НАУКИ / ECONOMICS

Научная статья

УДК 336.717:004.85

<https://doi.org/10.35266/2949-3455-2025-3-1>



Модели машинного обучения как инструмент обнаружения подозрительных банковских транзакций

Ольга Геннадьевна Аркадьева^{1✉}, Александр Васильевич Петров²

¹Чувашский государственный университет имени И. Н. Ульянова, Чебоксары, Россия

²ООО «Гарден СПА Отель», Чебоксары, Россия

Аннотация. В условиях цифровой трансформации банковской системы и расширения доступности финансовых услуг проблема обеспечения безопасности банковских транзакций приобретает стратегическое значение. Актуальность темы обусловлена необходимостью изучения проблемы эффективности методов прогнозирования и предотвращения мошенничества, которые являются важными составляющими комплексной системы экономической безопасности в банковском секторе.

В данном исследовании сделан акцент на формировании моделей обнаружения мошенничества с банковскими картами. В результате исследования предложены конкретные рекомендации по совершенствованию систем антифрод-защиты в банковском секторе с использованием моделей машинного обучения. В рамках этого исследования были разработаны и протестированы модели XGBoost и ANN для выявления мошеннических транзакций, реализованные на языке Python. Эксперименты продемонстрировали их высокую эффективность, а также гибкость и способность адаптироваться к новым данным. Использование этих моделей позволяет банкам оперативно обнаруживать подозрительные операции, что снижает риски потерь и улучшает общую защищенность финансовой системы.

Обеспечение безопасного функционирования системы совершения банковских транзакций требует комплексного подхода, включающего не только внедрение современных аналитических инструментов, но и постоянное обучение и обновление моделей для выявления новых способов мошенничества. Реализация предложенных мер позволит кредитным организациям повысить эффективность защиты от мошенничества, снизить финансовые потери от мошеннических действий и укрепить доверие клиентов.

Ключевые слова: искусственный интеллект, типология моделей, метрики качества моделей, ограничения методов, градиентный бустинг, нейронная сеть

Для цитирования: Аркадьева О. Г., Петров А. В. Модели машинного обучения как инструмент обнаружения подозрительных банковских транзакций // Вестник Сургутского государственного университета. 2025. Т. 13, № 3. С. 8–21. <https://doi.org/10.35266/2949-3455-2025-3-1>.

Original article

Machine learning models as tool for detecting suspicious bank operations

Olga G. Arkadeva^{1✉}, Aleksandr V. Petrov²

¹I. N. Ulianov Chuvash State University, Cheboksary, Russia

²Garden Hotel & SPA, Cheboksary, Russia

Abstract. Given the digital transformation of the banking system and the expanding availability of financial services, the problem of ensuring the security of banking operations becomes strategically important. The importance of the issue lies in studying the effectiveness of fraud prediction and prevention methods, which are key components of a complete economic security system in the banking sector.

Authors focus on the formation of bank card fraud detection models. The research proposes specific recommendations for improving anti-fraud systems in the banking sector using machine learning models. The researchers implemented and tested XGBoost and ANN models in Python to detect fraudulent transactions. The experiments substantiate their high performance and their flexibility and ability to adapt to new data. Using these models enabled banks for quick detection of suspicious operations, which reduces the risk of losses and improves the overall financial system security.

Ensuring the secure functioning of the banking transaction system requires a comprehensive approach that includes not only the introduction of modern analytical tools but also continuous training and updating models to identify new methods of fraud. Implementation of the proposed measures will enable credit organizations to improve the effectiveness of fraud protection, reduce financial losses from fraudulent actions and strengthen client confidence.

Keywords: artificial intelligence, models typology, model quality metrics, limitations of methods, gradient boosting, neural network

For citation: Arkadeva O. G., Petrov A. V. Machine learning models as tool for detecting suspicious bank operations. *Surgut State University Journal*. 2025;13(3):8–21. <https://doi.org/10.35266/2949-3455-2025-3-1>.

ВВЕДЕНИЕ

Цифровизация финансовых услуг открыла клиентам коммерческих банков новые возможности для удобного управления средствами, но параллельно увеличила риск финансовых преступлений, включая мошеннические действия. Масштабы мошеннических операций становятся значительной угрозой для финансовых институтов и их клиентов по всему миру, а экономический ущерб от финансового мошенничества измеряется миллиардами долларов в год. В этих условиях возникает необходимость в развитии высокоэффективных инструментов для защиты от мошенничества.

Методы машинного обучения (далее – МО) и в целом искусственного интеллекта (далее – ИИ) в последние годы стали важными компонентами защиты от мошенничества. Используя автоматизированные алгоритмы, эти технологии позволяют выявлять аномалии в поведении пользователей, обнаруживать подозрительные транзакции и реагировать на потенциальные угрозы в реальном времени. Благодаря этому внедрение ИИ и МО предоставляет мощные инструменты для своевременного и точного выявления мошеннических операций, минимизируя финансовые риски.

Каждый случай мошенничества не только приводит к финансовым потерям, но и негативно сказывается на репутации банков, что может вызвать отток клиентов, поэтому разработка и внедрение эффективных мер

защиты становятся приоритетом для финансового сектора. Управление рисками, мониторинг транзакций в реальном времени и повышение осведомленности клиентов о способах мошенничества – ключевые элементы, способствующие снижению количества преступлений в этой области.

Правильное и эффективное использование инструментов обнаружения и предотвращения мошенничества оказывает положительное влияние на банковскую систему, поскольку позволяет:

- выявлять подозрительную активность и финансовые махинации;
- идентифицировать клиентов и проверять их платежеспособность;
- автоматизировать анализ финансовых операций и обнаруживать аномалии с использованием технических средств;
- защищать аккаунты клиентов от несанкционированного доступа и взлома с помощью биометрических технологий и двухфакторной аутентификации;
- повышать уровень кибербезопасности и обучать персонал;
- сотрудничать с правоохранительными органами и развивать корпоративную культуру.

В данном исследовании сделан акцент на формировании моделей обнаружения мошенничества с банковскими картами, поскольку это одна из наиболее распространенных сфер совершения мошеннических действий [1].

МАТЕРИАЛЫ И МЕТОДЫ

Для сопоставления возможностей различных моделей МО использован метод описания. В ходе подготовки к апробации практической части в программной среде разработки и выполнения программного кода на языке Python в облаке – Google Colab – были импортированы библиотеки pandas, numpy, matplotlib, seaborn, xgboost, tensorflow. Основным прикладным методом исследования выступило формирование моделей машинного обучения XGBoost и нейронной сети ANN и сравнение результатов их работы по метрикам качества. Для проверки гипотезы об эффективности данных моделей машинного обучения в целях обнаружения подозрительных транзакций использовался набор данных creditcard.csv платформы Kaggle.

РЕЗУЛЬТАТЫ И ИХ ОБСУЖДЕНИЕ

В банковском секторе широко распространены технологии ИИ (в особенности МО) [2]. Большая часть прикладных проблем, которые решаются с их помощью, связана с деятельностью вокруг поиска и удовлетворения потребностей клиентов, включая привлечение, удержание, отток, улучшение клиентского опыта и легкости взаимодействия, а также формирование лояльности клиентов [3]. Вторая область применения ИИ (МО), в которой сконцентрировано большое количество научных исследований – автоматизированные методы кредитного скоринга [4, 5], третья – прогнозирование банковских кризисов [6]. В последнее время возрастает научный и практический интерес к проблеме обнаружения и предотвращения мошенничества [7]. В сферу внимания исследователей попадают и разные нефинансовые аспекты изучаемой проблемы – так, Р. К. Rajani и соавт. рассматривали использование людьми масок при пользовании банковскими услугами как угрозу безопасности в период пандемии COVID-19 и вне ее [8]. Исследование регуляторных практик не выявило единообразия в подходах к использованию ИИ в финансовом секторе, однако подтвердило необходимость развития технологий ИИ, в том числе в целях противодействия мошенничеству [9].

МО предлагает более динамичный и адаптивный подход к обнаружению подозрительных операций. Основные преимущества использования МО в задачах борьбы с мошенничеством включают:

- адаптацию к новым паттернам мошенничества. Алгоритмы МО могут обучаться на новых данных и адаптироваться к изменяющимся стратегиям мошенников;
- обнаружение скрытых закономерностей. МО позволяет анализировать большой объем данных и находить сложные взаимосвязи, которые трудно выявить с помощью простых аналитических правил;
- автоматизация и скорость. Алгоритмы могут анализировать транзакции в режиме реального времени, мгновенно выявляя аномалии и отправляя их на проверку;
- снижение уровня ложных срабатываний. За счет обучения на исторических данных и их анализа МО-системы способны более точно выявлять реальные случаи мошенничества, минимизируя ложные тревоги.

Эти подходы классифицируются на основе их ключевых тактик, которые включают контролируемое обучение, неконтролируемое обучение и обучение с подкреплением [10, 11]. Для выявления мошенничества используются несколько типов моделей машинного обучения, среди которых наиболее распространены следующие:

1. Модели на основе супервизорного обучения. Супервизорное обучение предполагает использование размеченных данных, где транзакции уже помечены как мошеннические или нет. Алгоритм обучается на таких данных, создавая модели, способные классифицировать новые транзакции. К распространенным алгоритмам этого типа относятся логистическая регрессия, деревья решений и случайный лес (Random Forest), метод опорных векторов (SVM).

2. Модели без учителя (Unsupervised Learning). Методы без учителя используются, когда размеченные данные недоступны или ограничены. Такие алгоритмы, как кластеризация и выявление выбросов, полезны для обнаружения аномалий, не требуя предварительно размеченных данных. К популярным методам относятся алгоритмы кластеризации

(например, K-средних) и методы выявления выбросов (Isolation Forest, алгоритмы на основе плотности).

3. Гибридные методы. Гибридные методы комбинируют преимущества как супервизорного, так и безнадзорного обучения. Например,

после того как алгоритм выявляет аномальные транзакции без участия учителя, супервизорная модель с учетом ранее размеченных данных уточняет классификацию.

Сравнительная характеристика методов ИИ и моделей МО приведена в табл. 1.

Таблица 1

Сравнительная таблица методов ИИ и моделей МО

Тип модели МО / направление ИИ	Описание	Примеры методов	Преимущества	Ограничения
Супервизорные модели МО	Модели обучаются на размеченных данных, где транзакции помечены как мошеннические или немошеннические	Логистическая регрессия, деревья решений, случайный лес, метод опорных векторов (SVM)	Высокая точность на размеченных данных, возможность точной классификации	Зависимость от объема и качества размеченных данных
Модели МО без учителя	Модели используют неразмеченные данные для выявления аномалий	Кластеризация (например, метод K-средних), Isolation Forest, методы на основе плотности (DBSCAN, HDBSCAN, OPTICS)	Не требует размеченных данных, что полезно для поиска неизвестных угроз	Менее точные результаты, возможны ложные срабатывания
Гибридные методы МО	Совмещение супервизорного и безнадзорного подходов для повышения эффективности	Сочетание кластеризации и классификации	Высокая адаптивность, сочетание преимуществ двух подходов	Более сложная реализация и настройка
Обработка естественного языка (NLP)	Анализ текстовой и речевой информации для выявления признаков мошенничества	NLP для анализа текстов и речи, классификация	Автоматизация анализа, высокая точность на текстовых данных	Зависит от объема обучающей выборки языковых моделей
Глубокое обучение для сложных данных	Использование глубоких нейронных сетей для анализа многомерных и сложных данных	Глубокие нейронные сети (DNN)	Высокая эффективность для сложных данных, адаптация к новым данным	Высокая вычислительная стоимость, необходимость обработки больших данных
Адаптивные алгоритмы для персонализации и самообучения	Алгоритмы обучаются и корректируются на основе новых данных	LR FTRL-Proximal (Google)	Адаптация к изменениям, точность классификации	Сложность настройки и управления
Распознавание биометрических данных	Использование биометрических данных для идентификации клиентов	Распознавание лиц, голоса, отпечатков пальцев	Высокая надежность и безопасность	Ограничения по точности, зависимость от качества данных
Прогнозирование и профилактика мошенничества	Прогнозирование мошеннических действий на основе анализа поведения и истории клиента	Предсказательные модели (ARIMA, GARCH)	Возможность предотвращения мошенничества до его совершения	Сложность интерпретации, необходимость в большом объеме данных

Примечание: составлено авторами.

ИИ в банковской сфере применяется для повышения точности и адаптивности. ИИ имеет важное значение для эффективности действий финансовых организаций, направленных на предупреждение случаев мошенничества в банковской сфере. Также ИИ используется для повышения качества и актуализации систем обнаружения кибератак. Применение ИИ улучшает процессы идентификации совершаемых операций, снижает различные финансовые риски и укрепляет безопасность банков. Кроме того, технологии ИИ постоянно адаптируются к новым методам и формам мошенничества, оперативно маркируют подозрительные операции, совершаемые мошенниками в целях обмана клиентов.

ИИ (в частности, МО) применяется по следующим основным направлениям:

1. МО для обнаружения аномалий, анализа транзакций. Алгоритмы МО способны анализировать значительный объем банковских транзакций, они служат для выявления аномальных операций в режиме реального времени, распознают закономерности в поведении своих клиентов и определяют случаи нетипичного поведения клиентов, которые могут указывать на мошеннические действия. Модели МО могут адаптироваться к появлению новых данных и распознавать новые схемы, взятые на вооружение мошенниками.

2. Обработка естественного языка в целях анализа текстовой информации. Обработка естественного языка (NLP) способна анализировать текстовую информацию в режиме реального времени. Это касается заявок на кредиты, сообщений от клиентов с признаками мошенничества.

3. Глубокое обучение (далее – ГО) для анализа данных. Глубокие нейронные сети (DNN) способны обрабатывать сложные и многомерные данные (изображения, аудиофайлы, числовую информацию), что позволяет обнаруживать случаи мошенничества в многослойных данных. Так, например, ГО применяется для анализа фото- и видеофайлов для проверки подлинности документов.

4. Адаптивные алгоритмы для самообучения и персонализации. С помощью ИИ

внедряются адаптивные алгоритмы. Они способны быстро обучаться, используя новые сведения о мошеннических схемах, и автоматически корректировать свои прогнозы. В условиях постоянно меняющейся среды и совершенствования мошеннических схем это особенно актуально. Адаптивные алгоритмы автоматически подстраиваются под поведение каждого клиента, меняют критерии оценки рисков и их пороговые значения.

5. Биометрическая идентификация для усиления безопасности. В последнее время банки все чаще применяют биометрическую идентификацию (распознавание лиц, голосов, отпечатков пальцев). С помощью биометрии очень точно определяется личность клиента, предотвращается доступ к его банковским системам со стороны посторонних лиц. Применение биометрии вкупе с ИИ обеспечивает качественный уровень защиты. Алгоритмы способны выявлять все подозрительные изменения в биометрических данных и выдавать предупреждения о возможных угрозах безопасности.

Прогностические модели, основанные на ИИ, позволяют не только выявлять мошеннические действия в процессе их совершения, но и предсказывать возможные риски. Анализ исторических данных и поведения клиентов с помощью ИИ помогает определить, когда и какие действия могут привести к мошенничеству, и принимать превентивные меры. Применение ИИ для повышения точности и адаптивности систем борьбы с мошенничеством значительно укрепляет стабильность и безопасность банковской системы, снижая вероятность финансовых потерь. Внедрение технологий ИИ позволяет банкам быстро реагировать на изменения, успешно противостоять новым видам мошенничества и укреплять доверие клиентов к финансовому сектору.

Для создания модели обнаружения мошенничества важно использовать соответствующий набор данных, включающий как нормальные, так и мошеннические транзакции. Как правило, такие наборы данных содержат обезличенные транзакции с информацией о суммах операций, времени их

осуществления, клиентах и других сведениях, которые помогают выявлять аномалии.

В данном исследовании использовался набор данных «creditcard.csv», содержащий транзакции по кредитным картам держателей из стран Европы [12]. Всего за два дня выявлено 492 случая мошенничества из 284 807 транзакций. Положительный класс (мошенничество) составил 0,172 % от указанного количества транзакций. Набор данных несбалансирован, в нем имеются только числовые переменные, которые были получены в результате преобразования методом главных компонент (PCA). В результате были получены признаки V1, V2, ..., V28, являющиеся основными компонентами, в то же время признаки «Время» и «Сумма» этим методом не были преобразованы.

Признак «Время» содержит секунды, которые произошли между доступом и первой транзакцией в наборе данных. Признак «Сумма» содержит общую сумму транзакций, он мо-

жет быть использован для обучения с учетом объема транзакции. Признак «Класс» является переменной ответа. В случае положительного ответа он принимает значение 1, в случае отсутствия мошеннических действий – 0.

Перед началом моделирования был проведен анализ данных, основные этапы анализа включали:

- исследование описания полей;
- анализ дисбаланса классов. Определено, что доля мошеннических транзакций составляет лишь небольшой процент от общего числа. Из-за дисбаланса классов дальнейшие этапы потребуют специальных методов для повышения точности распознавания меньшего класса;
- исследование распределения признаков.

Построены графики распределения транзакционных характеристик, таких как сумма и время, чтобы выявить возможные закономерности или аномалии, отличающие мошеннические операции (рис. 1, 2).



Рис. 1. Распределение времени операций в наборе данных

Примечание: составлено авторами.

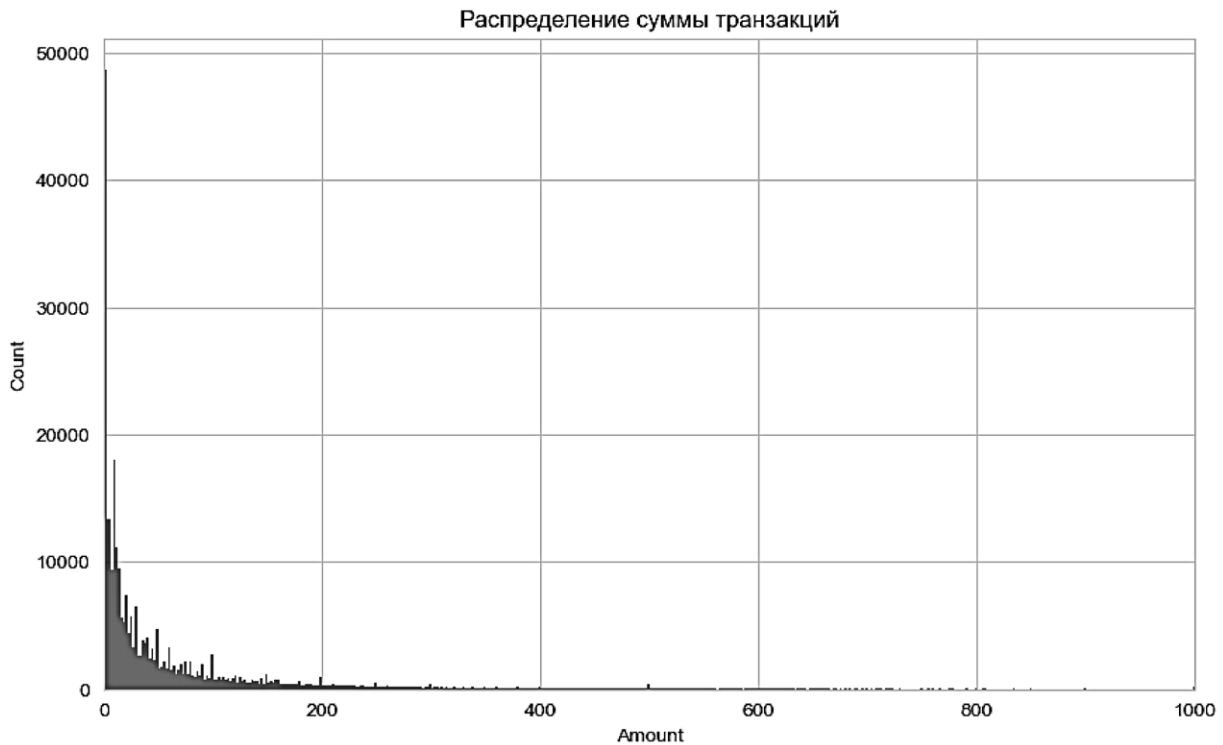


Рис. 2. Распределение суммы операций в наборе данных

Примечание: составлено авторами.

Традиционные аналитические действия не позволяют установить связь времени транзакции и мошеннических действий, несмотря на наличие явно выраженных экстремумов распределения транзакций по времени. Анализ распределения транзакций по сумме также непоказателен, что требует применения моделей МО для более глубокого анализа.

Проводя анализ распределений, можно получить представление об искажении этих признаков; кроме того, определяется дальнейшее распределение других признаков и определение методов, которые могут сделать распределение менее искаженными.

Для выявления возможных связей между признаками была построена корреляционная матрица. Это может помочь модели по распознаванию мошеннических транзакций (рис. 3). Самые высокие корреляции получены от: время и V3 (-0,42), сумма и V2 (-0,53), сумма и V4 (0,4).

Несмотря на то что эти корреляции являются достаточно высокими, связь не настолько выраженная, чтобы считать эти переменные подверженными риску мультиколлинеарности.

Ни один из компонентов PCA V1–V28, как показывает матрица корреляции, не имеет никакой выраженной корреляции друг с другом. Однако класс имеет некоторую форму положительной и отрицательной корреляции с компонентами V и не имеет корреляции с временем операций и их суммами.

Полученные результаты анализа помогают уточнить гипотезы о паттернах, характеризующих мошеннические операции, и определить оптимальный подход к обработке данных, но сами по себе не могут служить базой для формирования четких критериев мошеннических действий.

Предварительная обработка данных в наборе состоит из нескольких этапов, необходимых для подготовки данных в целях обучения модели для дальнейшего обнаружения мошеннических транзакций.

Работа включает следующие этапы:

- нормализация признаков;
- разделение данных на признаки и целевую переменную;
- разделение на обучающую и тестовую выборки;

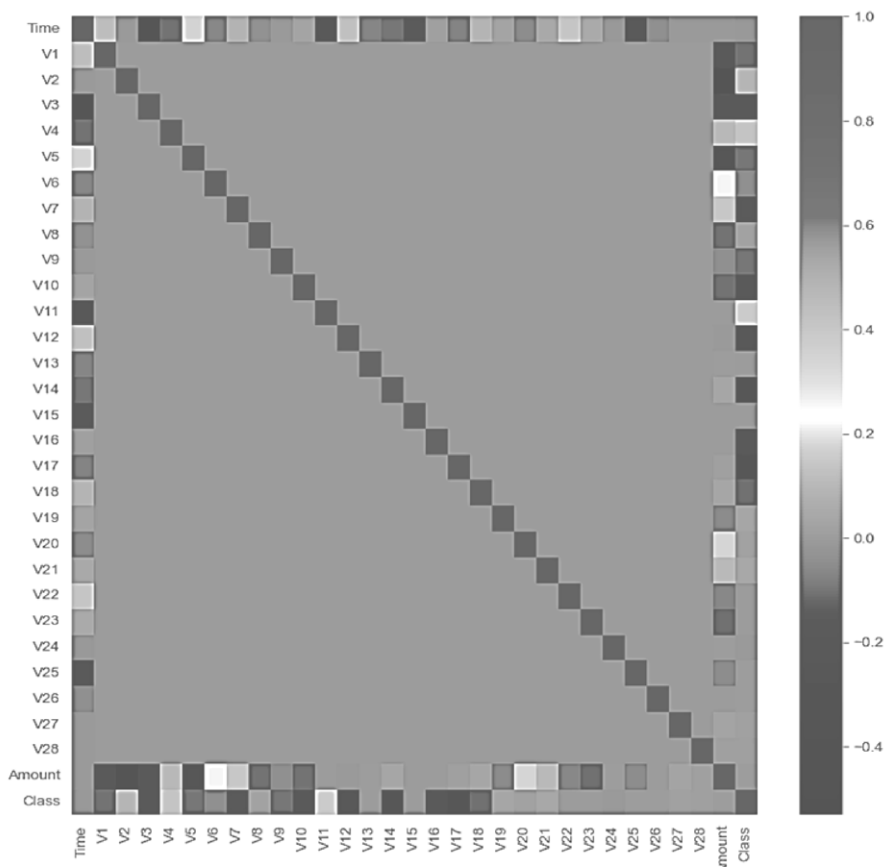


Рис. 3. Корреляционная матрица по набору данных

Примечание: составлено авторами.

- расчет весов классов;
- вывод формы массивов данных.

XGBoost (Extreme Gradient Boosting) представляет собой алгоритм градиентного бустинга, показывающий высокую эффективность в задачах классификации. Это особенно хорошо проявляется на дисбалансных данных, таких как наборы данных с мошенническими транзакциями. Популярность алгоритма объясняется высокой точностью и способностью эффективно работать с большими объемами данных, хотя и за счет потерь в скорости вычислений.

Этапами построения модели являются:

- инициализация и обучение модели. Для тонкой настройки модели XGB Classifier можно дополнительно оптимизировать гиперпараметры `learning_rate`, `max_depth`, `estimators` вместо применения параметров по умолчанию;
- прогнозирование и оценка точности. Обучение модели происходит на тренировоч-

ных данных `X_train` и `y_train`. После окончания обучения модель может прогнозировать значения как на тренировочной, так и на тестовой выборках.

Искусственная нейронная сеть (ANN) использует библиотеку Tensor Flow (модуль keras). ANN обрабатывает сложные и нелинейные зависимости в данных, адаптируется к признакам для более успешного выполнения задач по обнаружению мошенничества.

Структура модели:

- модель содержит несколько полносвязных слоев (Dense). Каждый из них состоит из 256 нейронов с функцией активации ReLU, что помогает моделировать нелинейные зависимости;
- batch normalization применяется после каждого слоя, стабилизируя и ускоряя процесс обучения;
- dropout с вероятностью 0,3 помогает исключить переобучение, случайным образом «отключая» нейроны во время обучения;

– выходной слой с функцией активации `sigmoid` и одним нейроном используется для бинарной классификации (мошенническая операция или нет).

Модель компилируется с использованием оптимизатора `Adam` и низкой скоростью обучения (`learning_rate=1e-4`), что способствует более точной настройке параметров. В качестве функции потерь используется `binary_crossentropy`, так как задача – бинарная классификация. В список метрик включены показатели `precision` и `recall`, а также `True/FalsePositives` и `True/FalseNegatives` для глубокого анализа результатов.

Обучение модели:

– параметр `class_weight`: Веса классов `{0: w_p, 1: w_n}` используются для борьбы с дисбалансом классов, так как мошеннических транзакций в выборке меньше;

– `CallbackModelCheckpoint`: используется для сохранения модели после каждой эпохи, чтобы в случае успешного результата можно было выбрать лучшую версию;

– обучение проводится на 300 эпохах и с большим размером батча (2048), что снижает вариативность обновлений весов и стабилизирует обучение;

– валидационные данные помогают контролировать, как модель обучается и не переобучается ли она.

Визуализируются метрики на каждом этапе обучения (например, `Loss`, `FalseNegatives`, `Precision`, `Recall`). Эти графики позволяют оценить, как модель улучшала свои показатели, и увидеть, на какой эпохе она достигла наилучшего результата (рис. 4).

Для оценки точности моделей `XGBoost` и искусственной нейронной сети `ANN` использовались метрики классификации, такие как точность (accuracy), точность классификации (`precision`), полнота (`recall`), `F1`-меры (рис. 5), а также матрицы ошибок (рис. 6). Эти метрики помогают оценить, насколько хорошо модели справляются с задачей выявления мошеннических транзакций на тренировочных и тестовых данных.

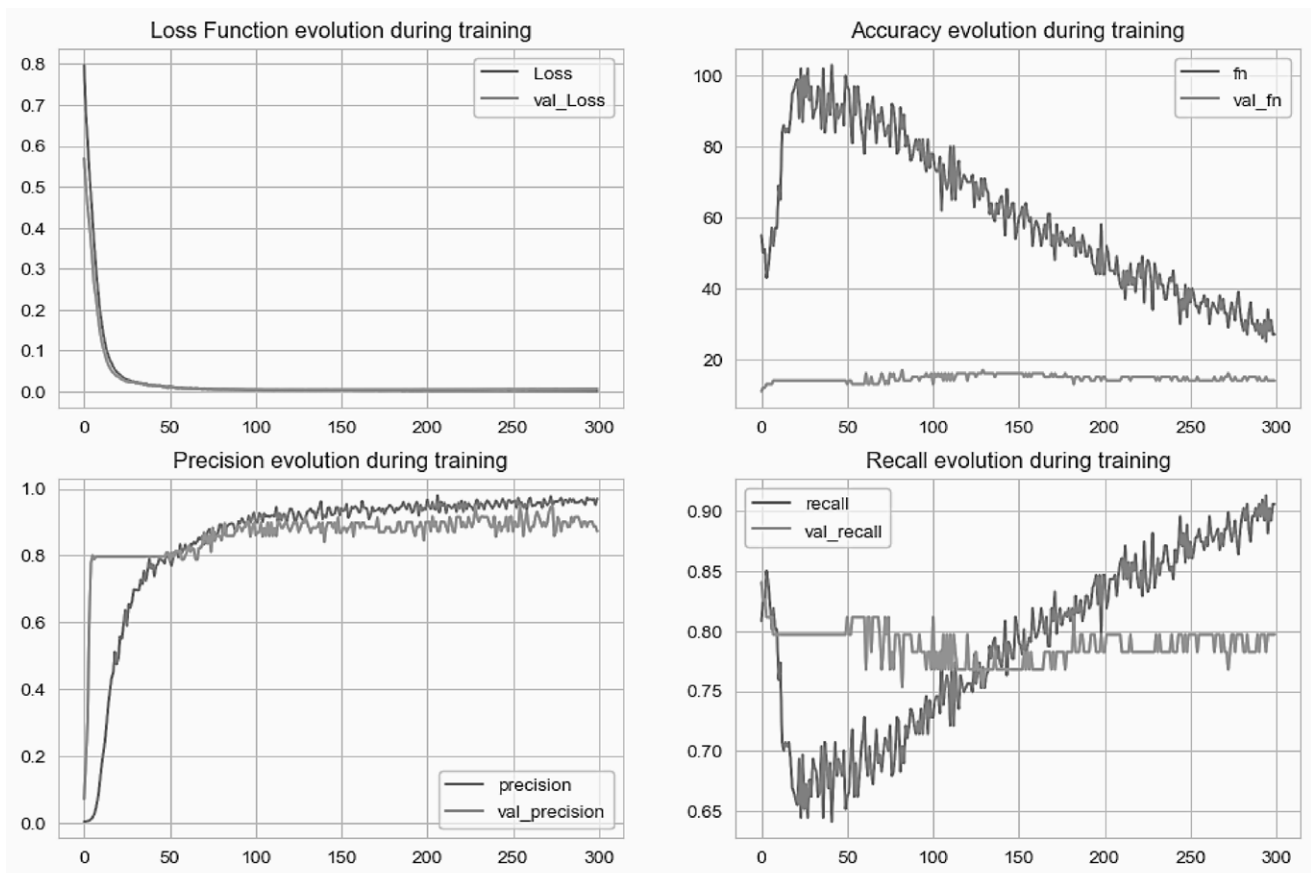


Рис. 4. Обучение XGBoost-модели

Примечание: составлено авторами.

Результаты на обученных данных:
=====

Точность прогнозирования: 100.00%

Результаты классификации:

	0	1	accuracy	macro avg	weighted avg
precision	1.00	1.00	1.00	1.00	1.00
recall	1.00	1.00	1.00	1.00	1.00
f1-score	1.00	1.00	1.00	1.00	1.00
support	159204.00	287.00	1.00	159491.00	159491.00

Матрица ошибок:
[[159204 0]
[0 287]]

Результаты на тестовых данных:
=====

Accuracy Score: 99.96%

Результаты классификации:

	0	1	accuracy	macro avg	weighted avg
precision	1.00	0.95	1.00	0.97	1.00
recall	1.00	0.82	1.00	0.91	1.00
f1-score	1.00	0.88	1.00	0.94	1.00
support	85307.00	136.00	1.00	85443.00	85443.00

Матрица ошибок:
[[85301 6]
[25 111]]

Рис. 5. Метрики качества XGBoost-модели

Примечание: составлено авторами.

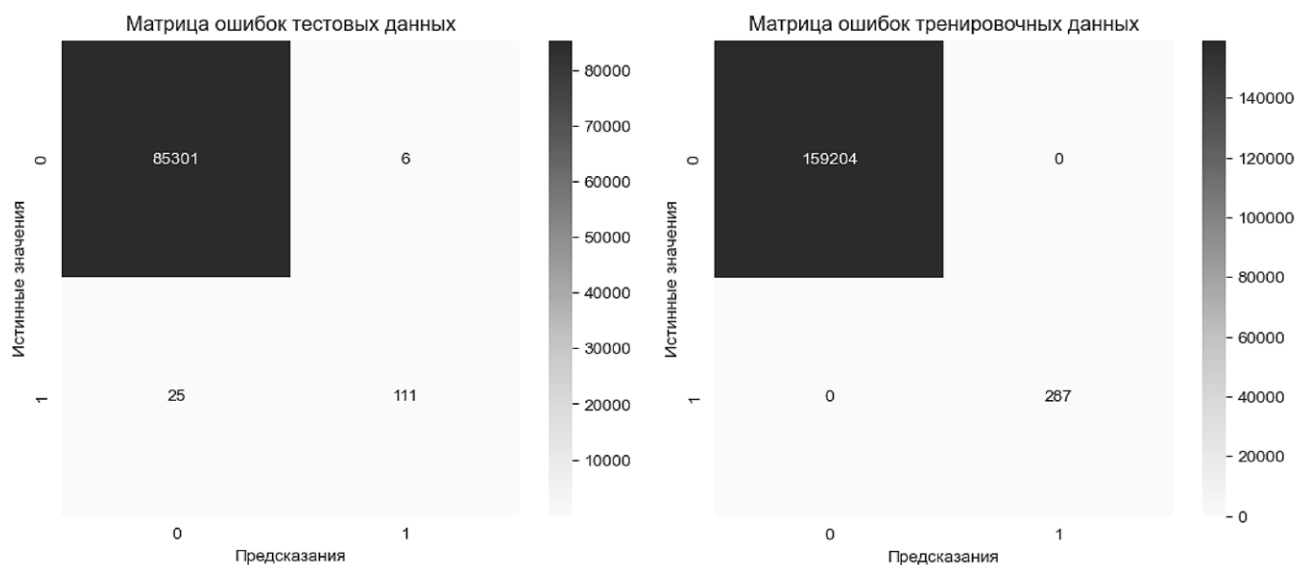


Рис. 6. Матрицы ошибок XGBoost-модели

Примечание: составлено авторами.

Матрица ошибок на тренировочной выборке не имеет ошибок (все мошеннические и немошеннические транзакции были правильно классифицированы). Точность модели на тестовой выборке составила 99,96 %, что также свидетельствует о высокой способности модели обучаться на новых данных. Модель достигла F1-score = 0,88 для класса

«мошенничество» и 1,00 для не мошеннического класса, что говорит о хорошей точности, но несколько сниженной полноте для класса «мошенничество» (метрика recall для этого класса составила 0,82).

Матрица ошибок на тестовой выборке показывает, что модель допустила 6 ошибок при классификации немошеннических

транзакций и 25 ошибок для мошеннических, что является приемлемым уровнем ошибок, учитывая общий дисбаланс данных.

Оценка модели ANN представлена на рис. 7, 8.

Точность на тренировочной выборке для ANN составила 99,99 %, что немного ниже по сравнению с XGBoost, но все еще свидетельствует о высокой точности. Модель также показала высокий уровень F1-score для обоих

```

4985/4985 ————— 5s 946us/step
2671/2671 ————— 3s 945us/step
Результаты на обученных данных:
=====
Точность прогнозирования: 99.99%

Результаты классификации:
      0      1  accuracy  macro avg  weighted avg
precision    1.00    1.00      1.00      1.00      1.00
recall       1.00    0.95      1.00      0.97      1.00
f1-score     1.00    0.97      1.00      0.99      1.00
support    159204.00  287.00      1.00  159491.00  159491.00

Матрица ошибок:
[[159203    1]
 [   15   272]]

Результаты на тестовых данных:
=====
Точность прогнозирования: 99.96%

Результаты классификации:
      0      1  accuracy  macro avg  weighted avg
precision    1.00    0.92      1.00      0.96      1.00
recall       1.00    0.81      1.00      0.90      1.00
f1-score     1.00    0.86      1.00      0.93      1.00
support    85307.00  136.00      1.00  85443.00  85443.00

Матрица ошибок:
[[85298     9]
 [   26   110]]
    
```

Рис. 7. Метрики качества ANN-модели

Примечание: составлено авторами.

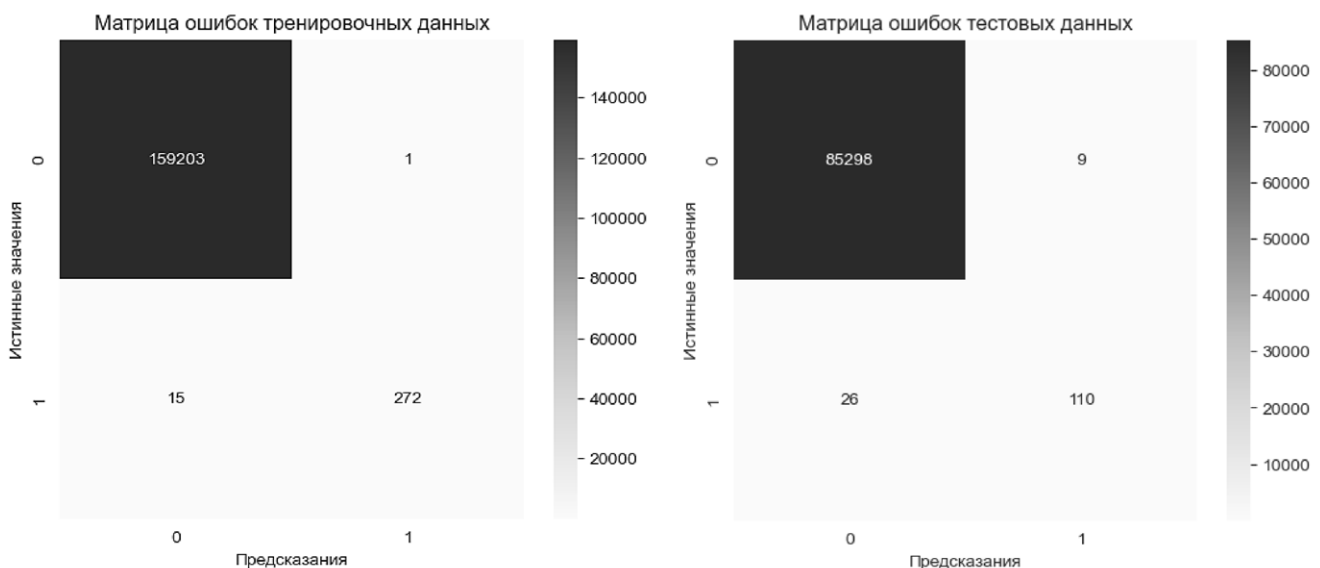


Рис. 8. Матрицы ошибок ANN-модели

Примечание: составлено авторами.

классов, однако для класса «мошенничество» F1-score составил 0,97, что отражает небольшие ошибки в классификации. В матрице ошибок на тренировочных данных видно, что модель допустила 1 ошибку для не мошеннического класса и 15 ошибок для мошеннического, что допустимо и свидетельствует о высокой обучаемости сети. Точность на тестовой выборке для ANN составила 99,96 %, что аналогично точности XGBoost. Метрика F1-score для класса «мошенничество» составила 0,86, что немного ниже, чем у XGBoost, в то время как для не мошеннического класса она сохранила значение 1,00. Матрица ошибок показывает 9 ошибок для не мошеннического класса и 26 ошибок для мошеннического. Хотя модель ANN также справляется с задачей, ее precision, F1-score и recall-метрики немного уступают модели XGBoost (табл. 2).

Обе модели продемонстрировали достаточно хорошие результаты на задаче обнаружения мошенничества с высокой точностью классификации. Однако модель XGBoost показала чуть более высокую эффективность в плане метрик F1-score и общего числа ошибок. Это позволяет сделать вывод, что XGBoost является предпочтительной моделью для решения данной задачи. Однако в случае использования любой из моделей остаются проблемы, связанные с такими явлениями, как дрейф концепций и дисбаланс классов, которые ставят задачи усовершенствования моделирования МО [13].

ЗАКЛЮЧЕНИЕ

Банковское мошенничество – серьезная проблема, с которой сталкиваются финан-

совые институты по всему миру. Технологический прогресс и рост онлайн-активности добросовестных экономических субъектов делают преступников более изобретательными и агрессивными, создавая новые вызовы для банков и их клиентов. Для обеспечения безопасности финансовых операций, защиты персональных данных и поддержания доверия клиентов к банковским услугам необходим комплексный подход к предотвращению мошеннических действий с использованием моделей машинного обучения для анализа банковских транзакций.

По причине того что риск потери репутации для банков зачастую превышает риски финансовых потерь, некоторые кредитные организации не раскрывают сведения о произошедших случаях мошенничества. Хотя эксперты в области банковской безопасности хорошо осведомлены о способах хищения средств, методах расследования и масштабах краж, внутренние инциденты, как правило, остаются поводами для внутренних расследований и последующих действий своими силами или с привлечением правоохранительных органов, но не банковского сообщества в целом. Между тем проблема мошенничества при совершении банковских транзакций для российского банковского сектора исключительно актуальна. Основные предпосылки роста такого мошенничества включают увеличение количества совершаемых финансовых операций, доступность персональных данных граждан, развитие интернет-торговли и технологических инноваций. В связи с этим банки и другие финансовые институты должны совершенствовать методы защиты данных, активно работать

Таблица 2

Сравнительная таблица метрик

Модель	Данные	Точность (%)	F1-score Fraud	F1-score Non-fraud	Ошибка Fraud	Ошибки Non-fraud
XGBoost	Тренировочные	100	1,00	1,00	0	0
	Тестовые	99,96	0,88	1,00	25	6
ANN	Тренировочные	99,99	0,97	1,00	15	1
	Тестовые	99,96	0,86	1,00	26	9

Примечание: составлено авторами.

в направлении повышения осведомленности клиентов о мерах безопасности, сотрудничать с правоохранительными органами и развивать собственную корпоративную культуру. При этих условиях грамотное и высококвалифицированное использование инструментов обнаружения и предотвращения мошенничества

на основе технологий МО позволит банкам точно и безошибочно выявлять подозрительную активность, точнее идентифицировать клиентов, эффективнее проводить аудит и контролировать свои процессы и операции, а также защищать аккаунты клиентов от несанкционированного доступа и взлома.

Список источников

1. Hernandez Aros L., Bustamante Molano L. X., Gutierrez-Portela F. et al. Financial fraud detection through the application of machine learning techniques: A literature review // *Humanities and Social Sciences Communications*. 2024. Vol. 11. <https://doi.org/10.1057/s41599-024-03606-0>.
2. Kalyani S., Gupta N. Is artificial intelligence and machine learning changing the ways of banking: A systematic literature review and meta analysis // *Discover Artificial Intelligence*. 2023. Vol. 3. <https://doi.org/10.1007/s44163-023-00094-0>.
3. Курносова В. В. Технологии искусственного интеллекта в банкинге // *Экономика и менеджмент инновационных технологий*. 2022. № 8. URL: <https://ekonomika.snauka.ru/2022/08/23330> (дата обращения: 06.02.2025).
4. Трусова А. Ю., Ильина А. И. Методология внедрения машинного обучения в банковской сфере // *Вестник Самарского университета. Экономика и управление*. 2023. Т. 14, № 4. С. 186–201. <https://doi.org/10.18287/2542-0461-2023-14-4-186-201>.
5. Шебалков М. П. Современные методы машинного обучения в банковской отрасли // *Экономика и управление: проблемы, решения*. 2022. Т. 3, № 9. С. 9–14. <https://doi.org/10.36871/ek.up.p.r.2022.09.03.002>.
6. Виноградов А. С. Использование машинного обучения в финансовом прогнозировании в банках // *Актуальные вопросы современной экономики*. 2022. № 5.
7. Аркадьева О. Г. Использование методов машинного обучения для отслеживания социальной инженерии в банковских транзакциях // *Oeconomia et Jus*. 2024. № 4. С. 1–14. <https://doi.org/10.47026/2499-9636-2024-4-1-14>.
8. Rajani P. K., Khaparde A., Bendre V. et al. Fraud detection and prevention by face recognition with and without mask for banking application // *Multimedia Tools and Applications*. 2025. Vol. 84. P. 781–804. <https://doi.org/10.1007/s11042-024-19021-1>.
9. Аркадьева О. Г., Березина Н. В. Формирование модели государственного регулирования развития технологий искусственного интеллекта в финансовом секторе // *Oeconomia et Jus*. 2023. № 4. С. 12–21. <https://doi.org/10.47026/2499-9636-2023-4-12-21>.
10. Лихоузов К. И. Применение задач машинного обучения на платформе распределенных вычислений больших данных в банковской сфере // *Цифровая трансформация управления: проблемы*

References

1. Hernandez Aros L., Bustamante Molano L. X., Gutierrez-Portela F. et al. Financial fraud detection through the application of machine learning techniques: A literature review. *Humanities and Social Sciences Communications*. 2024;11. <https://doi.org/10.1057/s41599-024-03606-0>.
2. Kalyani S., Gupta N. Is artificial intelligence and machine learning changing the ways of banking: A systematic literature review and meta analysis. *Discover Artificial Intelligence*. 2023;3. <https://doi.org/10.1007/s44163-023-00094-0>.
3. Kurnosova V. V. Artificial intelligence technologies in banking. *Economics and Management of Innovative Technologies*. 2022;(8). URL: <https://ekonomika.snauka.ru/2022/08/23330> (accessed: 06.02.2025). (In Russ.).
4. Trusova A. Y., Ilyina A. I. Methodology for the implementation of machine learning in the banking sector. *Bulletin of Samara University. Economics and Management*. 2023;14(4):186–201. <https://doi.org/10.18287/2542-0461-2023-14-4-186-201>. (In Russ.).
5. Shebalkov M. P. Modern methods of machine learning in the banking industry. *Ekonomika i upravlenie: problemy, resheniya*. 2022;3(9):9–14. <https://doi.org/10.36871/ek.up.p.r.2022.09.03.002>. (In Russ.).
6. Vinogradov A. S. Usage of machine learning in financial forecasting in banks. *Actual Issues of the Modern Economy*. 2022;(5). (In Russ.).
7. Arkadeva O. G. The use of machine learning techniques to track social engineering in banking transactions. *Oeconomia et Jus*. 2024;(4):1–14. <https://doi.org/10.47026/2499-9636-2024-4-1-14>. (In Russ.).
8. Rajani P. K., Khaparde A., Bendre V. et al. Fraud detection and prevention by face recognition with and without mask for banking application. *Multimedia Tools and Applications*. 2025;84:781–804. <https://doi.org/10.1007/s11042-024-19021-1>.
9. Arkadeva O. G., Berezina N. V. Forming the model of state regulation for the development of artificial intelligence technologies in the financial sector. *Oeconomia et Jus*. 2023;(4):12–21. <https://doi.org/10.47026/2499-9636-2023-4-12-21>. (In Russ.).
10. Likhovozov K. I. Primenenie zadach mashinnogo obucheniya na platforme raspredelennykh vychesleniy bolshikh dannykh v bankovskoy sfere. In: *Proceedings of 5th All-Russian Research-to-Practice Conference "Tsifrovaya transformatsiya upravleniya: problemy i resheniya"*, May 11, 2023, Moscow. Moscow:

- и решения : материалы V Всерос. науч.-практич. конф., 11 мая 2023 г., г. Москва. М. : Государственный университет управления, 2023. С. 127–129.
11. Gandhar A., Gupta K., Pandey A. K. et al. Fraud detection using machine learning and deep learning // *SN Computer Science*. 2024. Vol. 5. <https://doi.org/10.1007/s42979-024-02772-x>.
 12. Credit Card Fraud Detection. URL: <https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud/data> (дата обращения: 06.02.2025).
 13. Wang S., Chen B. Credit card attrition: An overview of machine learning and deep learning techniques // *Informatics. Economics. Management*. 2023. Vol. 2, no. 4. P. 134–144. <https://doi.org/10.47813/2782-5280-2023-2-4-0134-0144>.
- State University of Management; 2023. p. 127–129. (In Russ.).
 11. Gandhar A., Gupta K., Pandey A. K. et al. Fraud detection using machine learning and deep learning. *SN Computer Science*. 2024;5. <https://doi.org/10.1007/s42979-024-02772-x>.
 12. Credit Card Fraud Detection. URL: <https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud/data> (accessed: 06.02.2025).
 13. Wang S., Chen B. Credit card attrition: An overview of machine learning and deep learning techniques. *Informatics. Economics. Management*. 2023;2(4):134–144. <https://doi.org/10.47813/2782-5280-2023-2-4-0134-0144>.

Информация об авторах

О. Г. Аркадьева – кандидат экономических наук, доцент;

knedlix@yandex.ru✉

А. В. Петров – интернет-маркетолог;
shemerka588@yandex.ru

About the authors

O. G. Arkadeva – Candidate of Sciences (Economics), Docent;

knedlix@yandex.ru✉

A. V. Petrov – Internet Marketing Specialist;
shemerka588@yandex.ru