

УДК 004.89

Научная статья

DOI: 10.35330/1991-6639-2025-27-2-11-22

EDN: EWHPZV

Построение модели машинного обучения для прогнозирования мошеннических транзакций

А. Ф. Константинов✉, Л. П. Дьяконова

Российский экономический университет имени Г. В. Плеханова
115054, Россия, Москва, Стремянный переулок, 36

Аннотация. В статье представлена разработка модели машинного обучения для прогнозирования мошеннических транзакций на примере транзакционных данных банка. Рассмотрены особенности кодирования категориальных переменных, связанные с наличием времени в транзакционных данных, чтобы избежать утечек информации. Проведены эксперименты по применению баггинга (bootstrap aggregating) и созданию дополнительных переменных на основе их вклада в итоговый прогноз с применением Shapley values. Рассмотрены показатели качества модели машинного обучения и проведен их анализ.

Ключевые слова: мошеннические транзакции, catboost, кодирование категориальных переменных, catboost_encoder, target_encoder, bagging, создание переменных, Shapley values

Поступила 16.01.2025, одобрена после рецензирования 27.01.2025, принята к публикации 10.03.2025

Для цитирования. Константинов А. Ф., Дьяконова Л. П. Построение модели машинного обучения для прогнозирования мошеннических транзакций // Известия Кабардино-Балкарского научного центра РАН. 2025. Т. 27. № 2. С. 11–22. DOI: 10.35330/1991-6639-2025-27-2-11-22

MSC: 90C99

Original article

Building a machine learning model for predicting fraudulent transactions

A.F. Konstantinov✉, L.P. Dyakonova

Plekhanov Russian University of Economics
115054, Russia, Moscow, 36 Stremyannyy lane

Abstract. The article presents development of a machine learning model for predicting fraudulent transactions using transactional data from a bank. It discusses the features of encoding categorical variables related to the presence of time in the transactional data to avoid information leakage. Additionally, experiments were conducted on the application of bagging and the creation of additional variables based on their contribution to the final prediction using Shapley values. The quality metrics of the machine learning model are examined and analyzed.

Keywords: fraudulent transactions, catboost, encoding categorical variables, catboost_encoder, target_encoder, bagging, variables creation, Shapley values

Submitted 16.01.2025, approved after reviewing 27.01.2025, accepted for publication 10.03.2025

For citation. Konstantinov A.F., Dyakonova L.P. Building a machine learning model for predicting fraudulent transactions. *News of the Kabardino-Balkarian Scientific Center of RAS*. 2025. Vol. 27. No. 2. Pp. 11–22. DOI: 10.35330/1991-6639-2025-27-2-11-22

ВВЕДЕНИЕ

Финансовые организации ежедневно анализируют, оценивают и минимизируют большое количество рисков, связанных с финансовыми активами и обязательствами организаций. Основной целью управления финансовыми рисками является защита организаций от финансовых потерь, возникающих из-за изменений внешней среды, а также поддержание финансовой устойчивости и увеличение прибыли организаций.

Модели искусственного интеллекта (ИИ) широко применяются в управлении финансовыми рисками. Техники машинного обучения, используемые для управления финансовым риском, приведены в таблице 1 [1].

Таблица 1. Техники машинного обучения, используемые для управления финансовым риском

Table 1. Machine learning techniques for financial risk management application

Метод обучения	Задача обучения	Приложение для управления финансовым риском
Обучение с учителем	Классификация	Поиск мошенничества
		Оптимизация портфолио
		Кредитный скоринг и прогноз банкротства
	Регрессия	Прогноз волатильности
		Анализ чувствительности
		Моделирование претензий
		Резервирование потерь
		Моделирование смертности
Обучение без учителя	Кластеризация	Ценообразование в страховании
		Анализ чувствительности
		Кредитный скоринг и прогноз банкротства
	Определение аномалий	Поиск мошенничества
	Снижение размерности	Андеррайтинг в страховании
		Моделирование смертности
Обучение с подкреплением		Оптимизация портфолио
Обучение с частичным наблюдением		Анализ чувствительности

При управлении финансовыми рисками обычно рассматривают:

- риски, связанные с рынками (анализ чувствительности, оптимизация портфолио, предсказание волатильности);
- кредитные риски (кредитный скоринг, предсказание дефолта или банкротства);
- страхование и демографические риски (моделирование претензий, резервирование потерь, предсказание смертности, андеррайтинг в страховании);
- операционные риски (поиск мошенничества).

В данной работе будут рассматриваться мошеннические транзакции. В общей классификации финансовых рынков они относятся к операционным рискам.

В дополнение к рассмотренным техникам машинного обучения, используемым для управления финансовым риском, необходимо указать на особенность финансовых данных – запрет на разглашение персональной информации о клиентах и их операциях. В своей статье [2] Т. Awosika и соавт. рассматривают применение федеративного обучения [3] при обнаружении финансового мошенничества. При применении федеративного обучения модель обучается на данных, распределенных между разными финансовыми организациями, без необходимости централизованного сбора данных. Вместо передачи данных для обучения в общий сервер модели ИИ обучаются в каждой финансовой организации независимо и передаются непосредственно бинарные файлы моделей ИИ (веса коэффициентов моделей). Далее веса моделей агрегируются и обратно направляются в финансовые организации. Выполняется несколько итераций данного процесса обучения до достижения максимально возможных показателей качества центральной модели ИИ.

В статье [4] А. А. Ali и соавт. проводят исследования, связанные с прогнозированием мошенничества с финансовой отчетностью. Лучшие результаты по сравнению с классическими моделями машинного обучения (логистическая регрессия, деревья решений, машины опорных векторов, AdaBoost и случайный лес) показывают модели на основе бустинга (XGBoost) с проведенными мероприятиями по снижению дисбаланса классов и автоматизированной настройкой гиперпараметров модели. К применению ансамблевых бустинговых моделей следует подходить с особой осторожностью. В основе бустинговых моделей находятся неглубокие деревья решений. В случае сдвига входящих переменных (стоимостные показатели, старение населения) правила разбиения неглубоких деревьев перестают работать, так как, например, автомобиль стоимостью 2 млн руб. в 2020 г. (люксовый европейский автомобиль) существенно отличается от автомобиля за 2 млн руб. в 2025 г. (дешевый китайский кроссовер), и правила разбиения наблюдений по стоимости автомобиля не будут учитывать быстрое изменение в окружающей среде. В связи с этим при использовании бустинговых моделей необходимо проводить мероприятия по предупреждению сдвигов в данных.

Временная структура финансовых данных характерна в том числе для информации о мошеннических транзакциях. Со временем изменяются как способы мошенничества, так и входящие в модель ИИ данные (дрифт входных данных). К. Не и соавт. [5] получили лучшее качество прогнозирования финансовых временных рядов с использованием ансамблевой модели глубокого обучения (Convolutional Neural Network (CNN) – Long Short-Term Memory (LSTM) – AutoRegressive Moving Average (ARMA)) по сравнению с отдельным применением моделей (ARMA, Multi-Layer Perceptron (MLP), LSTM, CNN). Модель CNN-LSTM используется для моделирования данных в пространственно-временной плоскости. Модель ARMA используется для учета автокорреляции в данных. Эти модели объединены в ансамблевой структуре для моделирования смеси линейных и нелинейных характеристик данных в финансовых временных рядах. Таким образом, наблюдается тенденция к использованию ансамблей моделей ИИ, позволяющих решить целый ряд проблем, связанных с особенностями финансовых данных, и достичь лучших результатов.

Цели и задачи исследования. Целью исследования является применение алгоритмов ИИ для обнаружения мошеннических финансовых транзакций и определение моделей, дающих лучшие значения метрик качества. В задачи исследования входит оценка эффективности применения баггинга и создания дополнительных переменных на основе их

вклада в итоговый прогноз с применением Shapley values, оценка и анализ показателей качества классификации.

Методы исследования: анализ эффективности применения дополнительных техник и оценка вклада каждой переменной в итоговый прогноз с применением Shapley values.

ОПИСАНИЕ НАБОРА ДАННЫХ

Набор данных «Transactions Data Bank. Fraud Detection¹ (Данные транзакций банка. Обнаружение мошенничества)» представляет информацию о 1048574 транзакциях банка за период с 01.04.2012 по 31.10.2014. Основная цель – определить, является ли транзакция мошеннической. Набор данных содержит следующие поля:

- дата транзакции (Date);
- номер аккаунта (nameOrig);
- количество денег в транзакции (amount);
- количество денег до транзакции (oldbalanceOrig);
- количество денег после транзакции (newbalanceOrig);
- город, в котором транзакция производится (City);
- тип транзакции (перевод, внесение денег, получение денег) (type);
- тип карты клиента (Card Type);
- цель транзакции (Exp Type);
- пол (Gender);
- метка мошенничества (isFraud).

Доля мошеннических транзакций составляет 16,8 % (175785 из 1048574).

В наборе данных рассмотрено 986 городов Индии. Количество наблюдений равно 1048574 для всех полей.

В таблице 2 приведены базовые статистики категориальных переменных

Таблица 2. Базовые статистики категориальных переменных

Table 2. Basic statistics for the categorical variables

Наименование поля	Уникальных	Первое значение	Частота
Дата (Date)	1326	26-Apr-14	1167
Номер аккаунта (nameOrig)	1048316	C1900095842	2
Город (City)	986	Bengaluru, India	143733
Тип транзакции (type)	5	CASH_OUT	373641
Тип карты клиента (Card Type)	6	Silver	275540
Цель транзакции (Exp Type)	9	Food	220115
Пол (Gender)	2	F	551182

В таблице 3 приведены базовые статистики числовых переменных.

¹Ссылка на набор данных: <https://www.kaggle.com/datasets/qusaybtoush1990/transactions-data-bank-fraud-detection>

Таблица 3. Базовые статистики числовых переменных**Table 3.** Basic statistics for the numerical variables

Наименование поля	Количество наблюдений	Среднее значение	Стандартное отклонение	Минимум	25 % перцентиль	50 % перцентиль	75 % перцентиль	Максимум
Количество денег в транзакции (amount)	1048574	38 028	110 517	0	647	8 263	23 650	10 000 000
Количество денег до транзакции (oldbalanceOrg)	1048574	880 198	2 969 968		4 344	36 539	136 643	38 900 000
Количество денег после транзакции (newbalanceOrig)	1048574	842 171	2 936 373		918	20 552	90 307	38 893 191
Метка мошенничества (isFraud)	1048574	0,168	0,374		0	0	0	1

СХЕМА ВАЛИДАЦИИ

Для того чтобы показатели качества не зависели от единовременного разделения данных на тренировочные и тестовые, проведено разделение данных методом кросс-валидации StratifiedShuffleSplit на тренировочные (50 %) и тестовые (50 %) наборы. Далее для каждого набора тренировочных данных данные повторно 5 раз разделялись на тренировочный (80 %) и валидационный (20 %) набор данных с помощью метода кросс-валидации kFold.

ПРИМЕНЯЕМЫЕ МОДЕЛИ ИИ

В настоящее время наилучшие показатели качества для табличных данных показывают модели бустинга. Бустинг представляет собой метод ансамблевого машинного обучения, который объединяет прогнозы слабых (неглубоких) моделей машинного обучения в одну сильную. При обучении бустинга каждая следующая модель исправляет ошибки предыдущих моделей.

Наиболее популярными алгоритмами бустинга являются CatBoost, XGBoost, LightGBM. При сопоставимых показателях качества библиотека CatBoost (Categorical boosting) специально разработана для работы с категориальными переменными и поддерживает работу с категориальными переменными «из коробки», что упрощает процесс проведения экспериментов. По утверждению разработчиков, CatBoost статистически значительно превосходит Xgboost и LightGBM [6].

Дополнительно необходимо отметить важные особенности кодирования категориальных переменных. Catboost хорошо показывает себя при работе с редкими категориями несбалансированных классов, характерными для финансовых данных. При кодировании категориальных переменных методами One-Hot-encoding² или Label-encoding³ категории с небольшим количеством наблюдений укрупняются и объединяются в большую категорию «Прочее», при этом редкие категории «забываются» алгоритмами.

²One-Hot-encoding создает дополнительные столбцы по количеству уникальных значений категориальных переменных и помещает в него значение 1 или 0 (есть категория в строке или нет). Часто редкие категории объединяются в 1 столбец «Прочее».

³Label-encoding каждой категории присваивается отдельное число. При этом модель ИИ может обнаружить связи там, где их нет (например, 1>2). Также часто редкие категории объединяются в 1 столбец «Прочее».

Метод кодирования категориальных переменных Target_encoder [7] позволяет решить проблему «забывания» редких категорий. Перед началом обучения на тренировочных данных для каждого значения категориальной переменной рассчитывается вероятность положительного класса, и само категориальное значение заменяется на эту вероятность. Если категориальная переменная имеет небольшое число возможных значений (разовые категории), то ее значение заменяется на среднюю вероятность целевой переменной по всей обучающей выборке. Если же категориальная переменная встречается в большом числе наблюдений, то применяется среднее по значению категориальной переменной на данных для обучения, а не по всем наблюдениям, используемым для обучения. Баланс между средним по значению категории и средним по всем тренировочным данным настраивается с помощью параметра сглаживания (smoothing) метода кодирования Target_encoder.

Транзакционные данные имеют следующую особенность: при разделении данных случайным образом на тренировочные и тестовые при расчете среднего по категории в расчет средней вероятности положительного класса попадают данные будущих периодов, что создает утечку, связанную с изменением во времени частоты мошеннических транзакций. В отличие от остальных моделей машинного обучения Catboost имеет встроенный кодировщик категориальных переменных Catboost_encoder, который является вариацией Target_encoder. Достоинством метода кодирования категориальных переменных Catboost_encoder является то, что в нем решена проблема утечки при кодировании категориальных переменных за счет учета времени появления категориальной переменной. При кодировании категориальных переменных значение целевой переменной (частота положительного класса) вычисляется только по предшествующим во времени наблюдениям.

БАЗОВАЯ МОДЕЛЬ (КОД МОДЕЛИ: BASE)

Учитывая описанные выше достоинства модели CatboostClassifier [8] (Categorical boosting Classifier), в нашем исследовании она была выбрана в качестве базовой модели ИИ.

Сама модель CatboostClassifier представляет собой ансамбль решающих деревьев небольшой глубины, причем на каждой последующей итерации модель учится снижать псевдоошибки прогнозов предыдущих итераций деревьев.

При обучении модели были установлены следующие гиперпараметры: число итераций ‘количество итераций’(iterations) = 3000, ‘количество итераций для ранней остановки’(early_stopping_rounds) = 100, ‘набор данных для остановки обучения’(eval_set) = (X_val, y_val). Параметр ‘количество итераций’(iterations) установлен заранее завышенным. Catboost автоматически рассчитывает параметр, регулирующий скорость обучения ‘скорость обучения’(learning_rate) с учетом параметра ‘количество итераций’(iterations) и особенностей набора данных. При этом обучение остановится при отсутствии роста показателя качества на числе итераций, установленном в параметре ‘количество итераций для ранней остановки’(early_stopping_rounds) = 100. Таким образом достигается оптимальное соотношение количества итераций и темпа обучения. Точка (порог отсечения) отнесения к положительному классу – 0,5 по умолчанию.

МОДЕЛЬ БАГГИНГА (КОД МОДЕЛИ: BAGG_TEMP_08)

Баггинг (bootstrap aggregating) [9] – метод ансамблевого обучения, при котором тренировочные данные несколько раз случайным образом разделяются и на каждом наборе данных обучается модель ИИ. Для одного наблюдения формируются прогнозы нескольких моделей, которые объединяются путем усреднения. За счет разного набора данных снижается зависимость от случайного разделения данных на тренировочный и тестовый наборы данных. Схема баггинга приведена на рисунке 1.

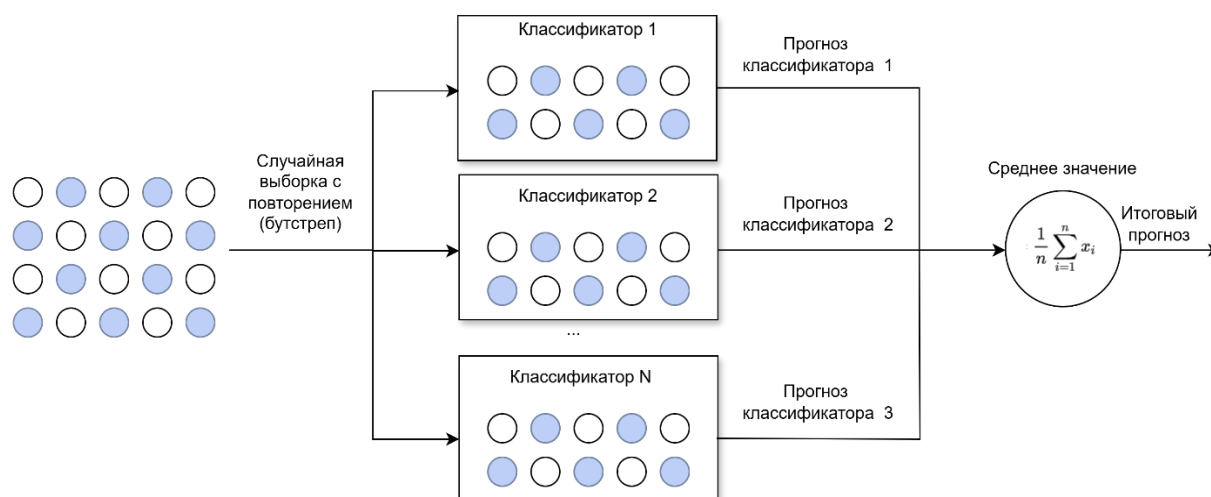


Рис. 1. Схема баггинга

Fig. 1. Bagging scheme

При обучении модели были установлены гиперпараметры, аналогичные базовой модели. Для того чтобы организовать баггинг, дополнительно мы добавили значение параметра «bagging_temperature» [10] = 0.8. По умолчанию этот параметр установлен равным 1.0⁴. Значения параметра «bagging_temperature» изменяются в диапазонах:

- «bagging_temperature» = 0: отменяет баггинг. При этом каждая отдельная модель будет обучаться на всем наборе данных;
- «bagging_temperature» > 0: включает баггинг. Увеличение этого параметра приводит к большей случайности при формировании подвыборок, что может помочь снизить переобучение модели, но также может увеличить вероятность недообучения.

МОДЕЛЬ СОЗДАНИЯ ДОПОЛНИТЕЛЬНЫХ ПЕРЕМЕННЫХ ДЛЯ САМЫХ БОЛЬШИХ ПО ВКЛАДУ SHAPLEY VALUES (КОД МОДЕЛИ: SHAP_COL_INTERACTION)

Shapley values – метод из теории кооперативных игр, который позволяет «честно» определить вклад значения каждой переменной в итоговый прогноз модели ИИ. Данный метод объяснимого искусственного интеллекта был разработан в 1951 году [11] и получил широкое распространение благодаря современной реализации на библиотеке SHAP [12].

Основная гипотеза данного эксперимента: создание новых признаков, основанное на поэлементных операциях над наиболее значимыми исходными переменными, иногда улучшает показатели качества моделей ИИ за счет обогащения признакового пространства, выявления/усиления скрытых взаимосвязей. Использование комбинаций только с переменными с наибольшим вкладом по Shapley values позволяет создавать новые, более релевантные переменные без перебора всех переменных, входящих в модель ИИ, что снижает уровень дополнительного шума.

В данном эксперименте формируется список переменных с наибольшим вкладом Shapley values в прогноз (на валидационных данных) только из числовых переменных с суммарным вкладом в больше 70 % от общего вклада Shapley values в прогноз всех

⁴Необходимо добавить, что параметр «bagging_temperature» может быть использован только при установке в параметре «bootstrap_type» значения, равного «Bayesian». В ходе экспериментов данный параметр не устанавливался дополнительно, но для задач бинарной классификации параметр «bootstrap_type» по умолчанию установлен равным «Bayesian».

числовых переменных. Для определения Shapley values обучаем временную базовую модель с параметрами, аналогичными базовой модели, с помощью встроенного в библиотеку SHAP метода TreeExplainer преобразуем валидационные данные в Shapley values. Далее на основе этого списка переменных с наибольшим вкладом Shapley values был сформирован список комбинаций переменных длиной 2, то есть получились комбинации ((«Столбец1», «Столбец2»), («Столбец1», «Столбец3»), ...). Для расчета новых полей по списку комбинаций переменных использовались методы поэлементных арифметических операций (сложение, вычитание, умножение, деление, остаток от деления). Данный подход включает следующие шаги:

- Разделение данных на данные для обучения (train), настройки (val), тестирования (test).
- Обучение модели CatboostClassifier с параметрами базовой модели на данных для обучения (train).
- Передача обученной модели CatboostClassifier алгоритму SHAP. Обучение алгоритма shap.TreeExplainer(CatboostClassifier).
- Преобразование данных для настройки (val) в Shapley values.
- Расчет суммарных вкладов каждой переменной в итоговый прогноз на данных для настройки (val).
- Сортировка суммарных вкладов переменных в итоговый прогноз от большего к меньшему.
- Отбор переменных с наибольшим вкладом так, чтобы вклад отобранных переменных в итоговый прогноз был больше 70 % от вклада всех переменных.
- Формирование списка комбинаций отобранных переменных длиной 2 (только числовые столбцы) ((«Столбец1», «Столбец2»), («Столбец1», «Столбец3»), ...).
- Формирование новых переменных: применение поэлементных арифметических операций (сложение, вычитание, умножение, деление, остаток от деления) для каждого элемента списка комбинаций отобранных переменных длиной 2 («Столбец1», «Столбец2») : (сложение, вычитание, ...).
- Обучение модели CatboostClassifier с параметрами базовой модели на наборе данных для обучения (train) с включенными новыми переменными.
- Замер показателей качества классификации на данных для тестирования (test).

РЕЗУЛЬТАТЫ ИССЛЕДОВАНИЯ

В таблице 4 приведены рассчитанные на тестовых данных средние метрики оценки эффективности моделей классификации для базовой модели.

Таблица 4. Результаты экспериментов

Table 4. Experiment results

Код эксперимента	Средняя точность (Average precision)	Сбалансированная точность (Balanced accuracy)	Оценка Брайера (Brier score)	F1-score	ROC AUC
base	0,57090	0,64304	0,10078	0,42777	0,85807
bagg_temp_08	0,57090	0,64304	0,10078	0,42777	0,85807
SHAP_col_interaction	0,49836	0,65012	0,10414	0,44715	0,84425

Все дополнительные эксперименты не дали улучшения показателей качества классификации. На рисунке 2 приведены ROC-AUC кривые моделей ИИ.

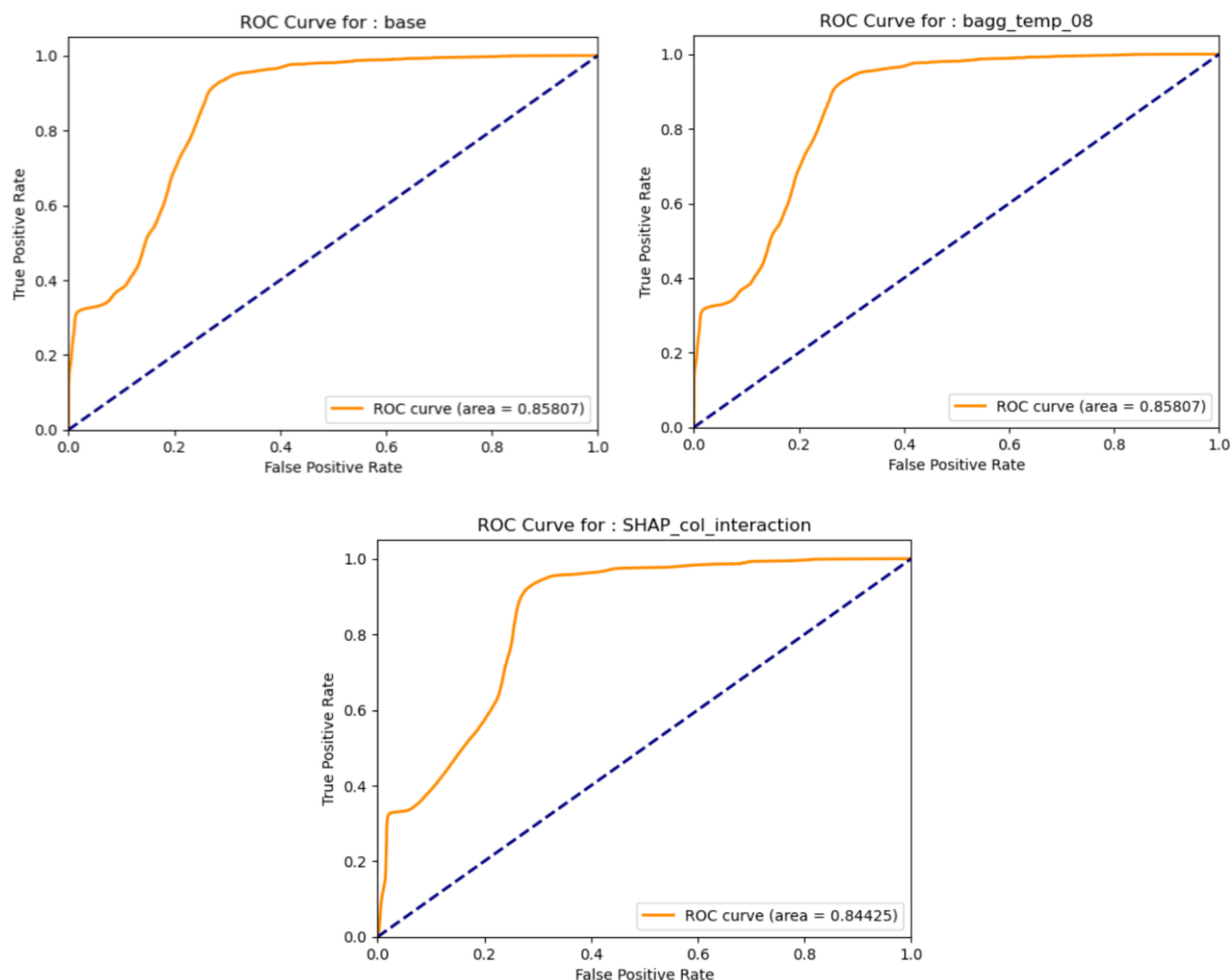


Рис. 2. ROC-AUC кривые моделей ИИ

Fig. 2. ROC-AUC curves of AI models

ОБСУЖДЕНИЕ РЕЗУЛЬТАТОВ

БАЗОВАЯ МОДЕЛЬ

Высокое значение показателя ROC_AUC говорит о хорошей ранжирующей способности модели. В то же время значения метрик, среднее гармоническое значение между точностью и полнотой F1-score и средняя точность довольно низкие. Значение оценки Брайера [13] = 0,10078, отличное от нуля, указывает на то, что модель не очень хорошо откалибрована, то есть прогнозные вероятности модели ИИ отличаются от фактических вероятностей.

Отличие параметров ROC_AUC от F1-score связано с их природой. Параметр ROC_AUC определяет ранжирующую способность модели ИИ и не учитывает долю наблюдений положительного класса в предсказанном положительном классе⁵. Параметр F1-score, рассчи-

⁵Краткая методика построения ROC_AUC: наблюдения ранжируются по прогнозной вероятности положительного класса (мошеннической транзакции) от большего к меньшему. Ось Y – True Positive Rate, Ось X – False Positive Rate. В начало координат ставится точка, и в цикле по ранжированным переменным от большей прогнозной вероятности положительного класса к меньшей: если значение истинно положительное (мошенническая транзакция), то отрезок кривой ROC_AUC делает шаг горизонтально вверх, в противном случае вертикально вправо.

танный как среднегармоническое между точностью и полнотой, учитывает долю правильно классифицированных наблюдений и сильно зависит от порога отнесения к положительному классу (установлен 0,5 по умолчанию). Исходя из этого для увеличения показателя качества F1-score необходимо пересмотреть точку отнесения к положительному классу классификации.

Для дальнейшего улучшения показателей качества модели необходимо провести модификацию настройки базовой модели:

- провести мероприятия по корректровке дисбаланса классов, например, применить метод OverSampling;
- провести настройку гиперпараметров с помощью байесовских методов, например, с помощью одного из лучших отраслевых решений по настройке гиперпараметров – алгоритма Optuna [14].

МОДЕЛЬ БАГГИНГА

По данной модели мы не получили статистически значимого прироста показателей качества. В базовой модели использовался Catboost с установленным по умолчанию параметром «bagging_temperature», который равняется 1.0, что соответствует применению технологии баггинга. В ходе эксперимента мы снизили «bagging_temperature» с установленного по умолчанию значения, равного 1, до 0.8, однако не провели сравнение с показателями качества без баггинга («bagging_temperature» = 0). Из-за того, что разница между 1.0 и 0.8 незначительная, получили статистически не значимое различие между базовой моделью и экспериментом.

При практическом применении на реальных задачах метод баггинга дает прирост показателей качества, поэтому в последующих экспериментах мы продолжим исследования, как лучше его настроить, чтобы получить улучшение показателей качества. В дальнейшем мы предполагаем применить другой метод обучения модели баггинга: используем прямое разделение данных дополнительным методом кросс-валидации kFold, обучим N (N изменяется от 1 до 100) моделей Catboost и усредним их прогнозы. При этом параметр «bagging_temperature» оставим установленным по умолчанию и равным 1.

МОДЕЛЬ СОЗДАНИЯ ДОПОЛНИТЕЛЬНЫХ ПЕРЕМЕННЫХ (КОД МОДЕЛИ: SHAP_COL_INTERACTION)

Для этой модели мы также не получили статистически значимого прироста показателей качества. Видимо, создание и включение столбцов с Shapley values не дает дополнительной информации, так как Shapley values сформированы статистическими методами из тех же данных, что и сами данные, используемые моделью ИИ для обучения. Shapley values хорошо использовать как инструмент объяснимого искусственного интеллекта, однако он не дает прироста качества классификации на финансовых данных. В последующих исследованиях мы рассмотрим гибридные подходы с объединением систем экспертных правил и моделей ИИ.

ЗАКЛЮЧЕНИЕ

В ходе исследования проведено построение модели ИИ на транзакционных данных банка. Рассмотрены особенности кодирования категориальных переменных при наличии в данных временных меток. Проведен анализ рассчитанных на тестовых данных показателей качества модели. Из всей совокупности показателей только метрика ROC_AUC имеет высокое значение, так как не зависит от выбора точки отнесения к положительному классу

классификации. Эксперименты по использованию баггинга и созданию новых переменных на основе вклада *Shapley values* в прогноз не дали статистически значимых результатов, в то же время анализ результатов привел к необходимости проведения дальнейших исследований. В последующем будет проведена корректировка метода баггинга с использованием прямого разделения данных дополнительным методом кросс-валидации *kFold* и обучением *N* моделей *Catboost* с усреднением их прогноза. Также представляет интерес рассмотрение гибридных подходов с объединением систем экспертных правил и моделей ИИ.

СПИСОК ЛИТЕРАТУРЫ / REFERENCES

1. Mashrur A., Luo W., Zaidi N.A., Robles-Kelly A. Machine Learning for Financial Risk Management: A Survey. *IEEE Access*. 2020. Vol. 8. Pp. 203203–203223. DOI: 10.1109/ACCESS.2020.3036322
2. Awosika T., Shukla R.M., Pranggono B. Transparency and Privacy: The Role of Explainable AI and Federated Learning in Financial Fraud Detection. *IEEE Access*. 2024. Vol. 12. Pp. 64551–64560. DOI: 10.1109/ACCESS.2024.3394528
3. McMahan B., Moore E., Ramage D. et al. Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*. 2017. Vol. 54. Pp. 1273–1282. DOI: 10.48550/arXiv.1602.05629
4. Ali A.A., Khedr A.M., El-Bannany M., Kanakkayil S. A Powerful Predicting Model for Financial Statement Fraud Based on Optimized XGBoost Ensemble Learning Technique. *Applied Sciences*. 2023. Vol. 13. No. 4. P. 2272. DOI: 10.3390/app13042272
5. He K., Yang Q., Ji L. et al. Financial Time Series Forecasting with the Deep Learning Ensemble Model. *Mathematics*. 2023. Vol. 11. No. 4. P. 1054. DOI: 10.3390/math11041054
6. Prokhorenkova L., Gusev G., Vorobev A. et al. CatBoost: unbiased boosting with categorical features. *NIPS'18: Proceedings of the 32nd International Conference on Neural Information Processing Systems*. 2018. Pp. 6639–6649. DOI: 10.48550/arXiv.1706.09516
7. Micci-Barreca D. A Preprocessing Scheme for High-Cardinality Categorical Attributes in Classification and Prediction Problems. *ACM SIGKDD Explorations Newsletter*. Vol. 3. No. 1. Pp. 27–32. DOI: 10.1145/507533.507538
8. Dorogush A.V., Ershov V., Gulin A. CatBoost: gradient boosting with categorical features support. *Workshop on ML Systems at NIPS*. 2017. DOI: 10.48550/arXiv.1810.11363
9. Breiman L. Bagging predictors. *Machine Learning*. 1996. Vol. 24. No. 2. Pp. 123–140. DOI: 10.1007/BF00058655
10. Official website Catboost. Common parameters. Точка доступа: https://catboost.ai/en/docs/references/training-parameters/common#bagging_temperature (дата обращения: 10 января 2025)
11. Shapley L. Notes on the *n*-person game, ii: the value of an *n*-person game. 1951.
12. Official website SHAP library. Точка доступа: https://shap.readthedocs.io/en/latest/example_notebooks/tabular_examples/tree_based_models/Catboost%20tutorial.html (дата обращения: 10 января 2025)
13. Brier Glenn W. Verification of forecasts expressed in terms of probability. *Monthly Weather Review*. 1950. Vol. 78. No. 1. Pp. 1–3. Bibcode:1950MWRv...78....1B. DOI: 10.1175/1520-0493(1950)078<0001:VOFEIT>2.0.CO
14. Akiba T., Sano S., Yanase T. et al. Optuna: A Next-generation Hyperparameter Optimization Framework. *KDD '19: Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. Pp. 2623–2631. DOI: 10.1145/3292500.3330701

Финансирование. Исследование проведено без спонсорской поддержки.

Funding. The study was performed without external funding.

Вклад авторов: все авторы сделали эквивалентный вклад в подготовку публикации. Авторы заявляют об отсутствии конфликта интересов.

Contribution of the authors: the authors contributed equally to this article. The authors declare no conflicts of interests.

Информация об авторах

Константинов Алексей Федорович, аспирант кафедры информатики, Российский экономический университет имени Г. В. Плеханова;

115054, Россия, Москва, Стремянный переулок, 36;

konstantinovaf@gmail.com, ORCID: <https://orcid.org/0009-0000-9591-3301>, SPIN-код: 3088-3121

Дьяконова Людмила Павловна, канд. физ.-мат. наук, доцент, кафедра информатики, Российский экономический университет имени Г. В. Плеханова;

115054, Россия, Москва, Стремянный переулок, 36;

Dyakonova.LP@rea.ru, ORCID: <https://orcid.org/0000-0001-5229-8070>, SPIN-код: 2513-8831

Information about the authors

Alexey F. Konstantinov, Post-graduate Student, Department of Informatics, Plekhanov Russian University of Economics;

115054, Russia, Moscow, 36 Stremyanny lane;

konstantinovaf@gmail.com, ORCID: <https://orcid.org/0009-0000-9591-3301>, SPIN-code: 3088-3121

Lyudmila P. Dyakonova, Candidate of Physical and Mathematical Sciences, Associate Professor, Department of Informatics, Plekhanov Russian University of Economics;

115054, Russia, Moscow, 36 Stremyanny lane;

Dyakonova.LP@rea.ru, ORCID: <https://orcid.org/0000-0001-5229-8070>, SPIN-code: 2513-8831