

ИНФОРМАТИКА И ИНФОРМАЦИОННЫЕ ПРОЦЕССЫ/INFORMATICS AND INFORMATION PROCESSES

DOI: <https://doi.org/10.60797/itech.2026.11.3>

EDN: AWZZXS

УСТОЙЧИВОСТЬ АНСАМБЛЕВЫХ МОДЕЛЕЙ МАШИННОГО ОБУЧЕНИЯ К ИСКАЖЕНИЯМ
ТАБЛИЧНЫХ ДАННЫХ

Научная статья

Джаббаров М.А.^{1,*}, Банных Н.А.²^{1,2} Московский Политехнический Университет, Москва, Российская Федерация

* Корреспондирующий автор (dzhabbarovmanuchehr[at]gmail.com)

Предложена: 25.03.2026; Принята: 13.07.2026; Опубликовано: 14.07.2026

Аннотация

Ансамблевые модели машинного обучения широко применяются в задачах анализа табличных данных благодаря высокой точности и гибкости, однако их поведение в условиях искажённых и неполных данных остаётся предметом активных исследований. В данной обзорной статье рассматриваются современные подходы к оценке устойчивости ансамблевых моделей машинного обучения к различным типам искажений данных.

В работе систематизированы основные виды искажений, характерные для практических задач, включая шум в признаках, пропуски значений и ошибки в метках классов. Проанализированы экспериментальные протоколы оценки устойчивости, применяемые в литературе, а также используемые метрики, позволяющие количественно оценивать деградацию качества моделей при увеличении уровня искажений. Особое внимание уделено сравнению подходов, основанных на бутстреп-агрегации и бустинге, с точки зрения их чувствительности к различным видам искажений.

Показано, что ансамблевые методы в целом демонстрируют повышенную устойчивость по сравнению с одиночными моделями, однако степень этой устойчивости существенно зависит от используемого алгоритма ансамблирования, характеристик данных и характера искажений. Отмечаются ограничения существующих методов оценки устойчивости и подчёркивается необходимость разработки унифицированных экспериментальных протоколов. Результаты обзора могут быть использованы при выборе и настройке моделей машинного обучения для построения надёжных аналитических систем в условиях неопределённости и неполноты данных.

Ключевые слова: ансамблевые модели, устойчивость моделей, шум в данных, пропуски данных, ошибки разметки, табличные данные, машинное обучение.

RESILIENCE OF ENSEMBLE MACHINE LEARNING MODELS TO DISTORTION OF TABULAR DATA

Research article

Dzhabborov M.A.^{1,*}, Bannikh N.A.²^{1,2} Moscow Polytech University, Moscow, Russian Federation

* Corresponding author (dzhabbarovmanuchehr[at]gmail.com)

Suggested: 25.03.2026; Accepted: 13.07.2026; Published: 14.07.2026

Abstract

Ensemble machine learning models are widely used in tabular data analysis tasks due to their high accuracy and flexibility, but their behavior in conditions of distorted and incomplete data remains the subject of active research. This review article examines modern approaches to assessing the resilience of ensemble machine learning models to various types of data distortion.

The paper systematizes the main types of distortions typical for practical tasks, including noise in features, missing values, and errors in class labels. Experimental stability assessment protocols used in the literature are analyzed, as well as the metrics used to quantify the degradation of model quality with increasing distortion levels. Particular attention is paid to comparing bootstrap aggregation and boosting approaches in terms of their sensitivity to various types of distortion.

It is shown that ensemble methods generally demonstrate increased stability compared to single models, however, the degree of this stability significantly depends on the used algorithm of the ensemble, the characteristics of the data and the nature of the distortion. The limitations of existing methods for assessing sustainability are noted and the need to develop unified experimental protocols is emphasized. The results of the review can be used in selecting and configuring machine learning models to build reliable analytical systems in conditions of uncertainty and incompleteness of data.

Keywords: ensemble models, model stability, data noise, data omissions, markup errors, tabular data, machine learning.

Введение

В последние годы устойчивость моделей машинного обучения к искажениям данных рассматривается как один из ключевых факторов надёжности аналитических систем [1], [2]. Рост объёмов и гетерогенности данных приводит к увеличению доли шумных наблюдений, пропусков и ошибок разметки, что существенно влияет на качество прогнозирования и интерпретируемость результатов [3], [4].

Особенно чувствительными к искажениям являются задачи анализа табличных данных, где классические методы обучения часто демонстрируют заметную деградацию качества [1], [3]. В этой связи ансамблевые модели, такие как

случайный лес и градиентный бустинг, рассматриваются как перспективный инструмент повышения устойчивости за счёт агрегирования базовых алгоритмов и использования механизмов регуляризации [1], [2], [5].

Современные исследования показывают, что модификации функций потерь и стратегий взвешивания позволяют повышать устойчивость ансамблей к шуму в признаках и метках без существенного роста вычислительной сложности [1], [2]. Однако подходы к оценке и повышению устойчивости ансамблевых моделей применительно к табличным данным с учётом различных типов искажений до настоящего времени не получили систематизированного изложения.

Целью настоящей статьи является анализ и систематизация методов оценки устойчивости ансамблевых моделей машинного обучения к искажениям данных. В работе рассматриваются основные типы искажений, экспериментальные протоколы и метрики устойчивости, а также обсуждаются преимущества и ограничения существующих подходов.

Методология исследования

Настоящая статья выполнена в соответствии с принципами систематического анализа литературы в области машинного обучения.

Поиск публикаций осуществлялся в базах Scopus, Web of Science и Google Scholar с использованием ключевых слов: ensemble learning, model robustness, noise in data, label noise, missing data, tabular data. Рассматривались преимущественно рецензируемые журнальные статьи и материалы международных конференций.

Отбор работ проводился по критериям: наличие анализа устойчивости моделей; использование ансамблевых методов; рассмотрение шума в признаках, пропусков данных или ошибок разметки. Публикации без количественной оценки устойчивости или не относящиеся к табличным данным исключались.

Анализ включал классификацию типов искажений, экспериментальных протоколов и метрик устойчивости. Метаанализ не проводился из-за различий в датасетах и моделях; вместо этого применялся качественный аналитический подход для выявления общих тенденций и ограничений.

Типы искажений данных в задачах машинного обучения

В задачах машинного обучения качество и устойчивость моделей во многом определяются корректностью и полнотой исходных данных. В реальных прикладных сценариях данные нередко содержат различные искажения, обусловленные ошибками измерений, неполнотой сбора информации и неточностями разметки. Современные исследования и нормативные документы по качеству данных подчёркивают, что подобные искажения являются систематическим, а не исключительным явлением в аналитических системах [3], [4]. Несмотря на отсутствие единой общепринятой классификации, в большинстве работ выделяются три наиболее распространённых и практически значимых типа искажений: шум в признаках, пропуски значений и ошибки в метках классов.

3.1 Шум в признаках

Шум в признаках представляет собой случайные или систематические искажения значений входных переменных, возникающие вследствие ошибок измерений, округлений или влияния внешних факторов. В современных исследованиях по устойчивости ансамблевых моделей шум в признаках, как правило, моделируется путём добавления аддитивных возмущений к числовым признакам либо случайного искажения категориальных переменных [1], [3]. Показано, что наличие шума приводит к росту дисперсии моделей и снижению их обобщающей способности.

Ансамблевые методы, основанные на усреднении предсказаний, в ряде случаев демонстрируют более высокую устойчивость к зашумленным признакам по сравнению с одиночными моделями, однако степень этой устойчивости существенно зависит от структуры шума и используемых механизмов регуляризации [1], [2].

3.2 Пропуски значений

Пропуски данных являются одним из наиболее распространённых искажений в табличных наборах данных и рассматриваются как важный фактор, влияющий на устойчивость моделей машинного обучения. В прикладных исследованиях подчёркивается, что пропуски могут возникать как случайным образом, так и в результате систематических ошибок сбора данных, что приводит к смещению оценок и ухудшению качества прогнозирования [3], [4].

Для ансамблевых моделей на основе деревьев решений пропуски значений часто оказываются менее критичными по сравнению с линейными моделями, поскольку такие методы допускают использование встроенных механизмов обработки отсутствующих значений или предварительной импутации данных [2], [5]. Тем не менее при увеличении доли пропусков устойчивость ансамблей также существенно снижается, что подчёркивает необходимость их явного учёта при анализе надёжности моделей.

3.3 Ошибки в метках классов

Ошибки в метках классов возникают в результате человеческого фактора, автоматической разметки или неоднозначности классов и оказывают прямое влияние на процесс обучения моделей. Современные исследования показывают, что шум в метках является одним из наиболее критичных типов искажений для ансамблевых методов, особенно для алгоритмов бустинга, чувствительных к ошибочным наблюдениям [1], [2].

В ряде работ предлагаются модификации функций потерь и стратегий взвешивания обучающих примеров, направленные на повышение устойчивости ансамблей к ошибочной разметке, что позволяет частично компенсировать негативное влияние шума в метках без существенного усложнения моделей [1], [5].

Ансамблевые модели машинного обучения и их устойчивость

Ансамблевые методы машинного обучения представляют собой класс алгоритмов, основанных на комбинировании базовых моделей с целью повышения качества прогнозирования и устойчивости результатов. Основная идея ансамблевого обучения заключается в снижении дисперсии, смещения или их комбинации за счёт агрегирования предсказаний отдельных моделей, что отражено в фундаментальных работах по теории машинного

обучения [1], [2]. Благодаря этим свойствам ансамблевые методы широко применяются в задачах анализа табличных данных и рассматриваются как базовый инструмент построения прикладных аналитических систем [3], [4].

С точки зрения устойчивости к искажениям данных ансамблевые модели представляют особый интерес, поскольку их архитектура позволяет частично компенсировать влияние шума, пропусков и ошибок разметки. В литературе устойчивость ансамблей обычно анализируется в рамках двух парадигм: бутстреп-агрегации и бустинга [1], [2].

4.1 Методы бутстреп-агрегации

Методы бутстреп-агрегации (bagging) основаны на обучении множества базовых моделей на подвыборках исходного датасета, сформированных с возвращением, с последующим усреднением их предсказаний. Классическим представителем данного подхода является случайный лес [1].

Случайные леса активно исследуются с точки зрения устойчивости к искажениям данных. За счёт случайного выбора подвыборок и подмножеств признаков при построении деревьев решений данный метод демонстрирует пониженную чувствительность к шуму в признаках и выбросам, что подтверждается сравнительными исследованиями [3], [4]. Дополнительным фактором устойчивости является способность деревьев работать с пропущенными значениями, что важно для табличных данных [5].

Показано, что при шуме в метках методы бутстреп-агрегации нередко обеспечивают более стабильное качество по сравнению с одиночными моделями за счёт ослабления влияния ошибочной разметки [4]. Однако при высоком уровне систематического шума устойчивость случайных лесов существенно снижается, что указывает на ограниченность данного подхода [3].

4.2 Методы бустинга

Методы бустинга представляют собой ансамблевые алгоритмы, в которых базовые модели обучаются последовательно, а каждая последующая фокусируется на ошибках предыдущих. Наиболее известным представителем является градиентный бустинг и его современные реализации, широко применяемые для анализа табличных данных [2], [3].

С точки зрения устойчивости к шуму в признаках бустинг демонстрирует неоднозначное поведение: использование неглубоких деревьев и регуляризации частично компенсирует случайные искажения, однако последовательный характер обучения делает модель чувствительной к выбросам, которые могут усиливаться на последующих итерациях [2].

Особое внимание уделяется чувствительности бустинга к ошибкам в метках классов. Показано, что даже небольшой процент ошибочной разметки способен существенно ухудшать качество моделей за счёт переобучения на шумных примерах [3], [6]. В связи с этим предлагаются модификации функций потерь и методы регуляризации для повышения устойчивости бустинга к шуму в метках [7].

4.3 Ансамблевые модели и табличные данные

Табличные данные характеризуются гетерогенностью признаков, наличием категориальных переменных и относительно небольшим числом наблюдений по сравнению с размерностью признакового пространства. В таких условиях ансамблевые методы на основе деревьев решений демонстрируют более высокую эффективность и устойчивость по сравнению с линейными моделями и нейронными сетями, что подтверждается сравнительными исследованиями [3].

Отсутствие единых стандартов оценки устойчивости ансамблевых моделей затрудняет сопоставление результатов различных работ и подчёркивает необходимость разработки унифицированных экспериментальных протоколов для анализа устойчивости моделей на табличных данных [5], [7].

Таким образом, ансамблевые модели обладают свойствами, способствующими повышению устойчивости к искажениям данных, однако её уровень существенно зависит от парадигмы ансамблирования, характеристик данных и типа искажений, что определяет направления дальнейшего анализа.

Методы оценки устойчивости моделей машинного обучения

Оценка устойчивости моделей машинного обучения направлена на анализ их поведения при наличии искажений входных данных и отличается от стандартной оценки качества на «чистых» тестовых выборках. В большинстве современных исследований устойчивость анализируется с использованием экспериментальных протоколов, основанных на контролируемом внесении искажений в данные и последующей оценке деградации качества моделей [1], [4].

5.1 Экспериментальные протоколы

Наиболее распространённым экспериментальным подходом является инъекция шума в признаки, при которой к числовым переменным добавляется случайный шум, а категориальные признаки подвергаются искажениям. Такой метод применяется для анализа чувствительности моделей к ошибкам измерений и выбросам и описывается в базовых работах по машинному обучению [1], [4]. Устойчивость оценивается по изменению показателей качества при росте уровня шума.

Для анализа влияния неполноты данных используется моделирование пропусков значений, при котором пропуски искусственно вводятся в выборку и оценивается влияние их наличия и стратегий обработки на качество моделей. Отмечается, что данный подход позволяет анализировать как устойчивость алгоритмов, так и чувствительность к методам импутации [5].

Отдельную группу составляют методы оценки устойчивости к ошибкам в метках классов, где часть целевых значений заменяется на ошибочные. Показано, что шум разметки является одним из наиболее критичных факторов, существенно влияющих на обучение ансамблевых моделей [6].

5.2 Метрики устойчивости

Для количественной оценки устойчивости используются стандартные метрики классификации, такие как ассигасу и F1-мера, вычисляемые при различных уровнях искажений. Дополнительно анализируется относительное снижение качества по сравнению с базовым уровнем на чистых данных, что упрощает сравнение моделей и экспериментальных настроек [3], [4].

В ряде работ применяются интегральные показатели, учитывающие поведение моделей на всём диапазоне искажений, например агрегирование метрик по нескольким сценариям деградации. Они дают обобщённую оценку устойчивости, однако обладают меньшей интерпретируемостью и зависят от выбранного протокола [1].

5.3 Ограничения подходов

Несмотря на широкое применение, экспериментальные методы оценки устойчивости имеют ряд ограничений. Различия в способах моделирования искажений, выборе уровней шума и используемых метрик затрудняют сопоставление результатов различных исследований. Кроме того, большинство работ рассматривает влияние отдельных типов искажений изолированно, что не всегда соответствует реальным условиям эксплуатации аналитических систем [5], [7].

5.4 Сравнение методов оценки устойчивости

На основании проведённого анализа в таблице 1 представлено сравнение основных методов оценки устойчивости моделей машинного обучения с указанием типов искажений, используемых метрик, а также преимуществ и ограничений каждого подхода.

Таблица 1 - Основные методы оценки устойчивости моделей машинного обучения

DOI: <https://doi.org/10.60797/itech.2026.11.3.1>

Метод	Тип искажения	Метрики	Преимущества	Ограничения
Инъекция шума в признаки	Шум, выбросы	Accuracy, F1	Простота, воспроизводимость	Ограниченная реалистичность
Моделирование пропусков	Пропуски данных	Accuracy, F1	Учёт механизмов пропусков	Зависимость от импутации
Шум в метках классов	Ошибки разметки	Accuracy, F1	Выявляет критическую чувствительность	Сложность моделирования
Относительная деградация	Все типы	Δ Accuracy, Δ F1	Удобство сравнения	Потеря абсолютных значений
Интегральные показатели	Все типы	AUC деградации	Комплексная оценка	Меньшая интерпретируемость

Таким образом, существующие методы оценки устойчивости позволяют проводить сравнительный анализ моделей машинного обучения в условиях искажённых данных, однако требуют аккуратного выбора экспериментального протокола и метрик. Представленная систематизация служит основой для обсуждения подходов к повышению устойчивости ансамблевых моделей.

Заключение

В настоящей статье проведён анализ и систематизация современных исследований, посвящённых устойчивости ансамблевых моделей машинного обучения к искажениям табличных данных. Рассмотрены основные типы искажений, характерные для прикладных задач, включая шум в признаках, пропуски значений и ошибки в метках классов, а также методы их моделирования в экспериментальных исследованиях.

Показано, что ансамблевые модели, в частности методы бутстреп-агрегации и градиентного бустинга, в целом демонстрируют более высокую устойчивость по сравнению с одиночными моделями, однако степень этой устойчивости существенно зависит от используемой парадигмы ансамблирования и характера искажений данных. Наиболее критичным фактором, согласно результатам рассмотренных работ, является наличие ошибок в метках классов, способное приводить к значительной деградации качества моделей.

Отмечено, что в настоящее время отсутствуют единые стандарты оценки устойчивости моделей машинного обучения, что затрудняет сопоставление результатов различных исследований и обобщение полученных выводов. В связи с этим перспективными направлениями дальнейших исследований являются разработка унифицированных экспериментальных протоколов, анализ совместного воздействия различных типов искажений, а также интеграция методов повышения устойчивости в практические пайплайны анализа табличных данных.

Полученные обобщения могут быть использованы при выборе и настройке ансамблевых моделей машинного обучения для построения надёжных аналитических систем, функционирующих в условиях неопределённости и неполноты данных.

**Конфликт интересов**

Не указан.

Рецензия

Все статьи проходят рецензирование. Но рецензент или автор статьи предпочли не публиковать рецензию к этой статье в открытом доступе. Рецензия может быть предоставлена компетентным органам по запросу.

Conflict of Interest

None declared.

Review

All articles are peer-reviewed. But the reviewer or the author of the article chose not to publish a review of this article in the public domain. The review can be provided to the competent authorities upon request.

Список литературы / References

1. Luo J. Robust-GBDT: GBDT with Nonconvex Loss for Tabular Classification in the Presence of Label Noise and Class Imbalance / J. Luo, Y. Quan, S. Xu // arXiv preprint. — 2023.
2. Sharma A.V. Learning Robust XGBoost Ensembles for Regression Tasks / A.V. Sharma, P. Kouvaros, A. Lomuscio // Proc. of the 41st Conference on Uncertainty in Artificial Intelligence, PMLR. — 2025. — № 286. — P. 3809–3825.
3. Kireev K. Transferable Adversarial Robustness for Categorical Data via Universal Robust Embeddings / K. Kireev, M. Andriushchenko, C. Troncoso [et al.] // arXiv preprint. — 2023.
4. Kasneci G. Enriching Tabular Data with Contextual LLM Embeddings: A Comprehensive Ablation Study for Ensemble Classifiers / G. Kasneci, E. Kasneci // arXiv preprint. — 2024.
5. Sierra-Fernández J.M. Robustness of Machine Learning and Deep Learning Models for Power Quality Disturbance Classification / J.M. Sierra-Fernández [et al.] // Applied Sciences. — 2025. — № 15(19). — P. 10602.
6. Zühlke MM. TCR: topologically consistent reweighting for XGBoost in regression tasks / M.M. Zühlke, D. Kudenko // Mach Learn. — 2025. — Vol. 114. — Art. 108. — DOI: 10.1007/s10994-024-06704-x
7. Журавлев В.В. Обзор различных алгоритмов машинного обучения в задачах обнаружения и исправления ошибок данных / В.В. Журавлев // Universum: технические науки. — 2024. — № 8(125). — DOI: 10.32743/UniTech.2024.125.8.18127.

Список литературы на английском языке / References in English

1. Luo J. Robust-GBDT: GBDT with Nonconvex Loss for Tabular Classification in the Presence of Label Noise and Class Imbalance / J. Luo, Y. Quan, S. Xu // arXiv preprint. — 2023.
2. Sharma A.V. Learning Robust XGBoost Ensembles for Regression Tasks / A.V. Sharma, P. Kouvaros, A. Lomuscio // Proc. of the 41st Conference on Uncertainty in Artificial Intelligence, PMLR. — 2025. — № 286. — P. 3809–3825.
3. Kireev K. Transferable Adversarial Robustness for Categorical Data via Universal Robust Embeddings / K. Kireev, M. Andriushchenko, C. Troncoso [et al.] // arXiv preprint. — 2023.
4. Kasneci G. Enriching Tabular Data with Contextual LLM Embeddings: A Comprehensive Ablation Study for Ensemble Classifiers / G. Kasneci, E. Kasneci // arXiv preprint. — 2024.
5. Sierra-Fernández J.M. Robustness of Machine Learning and Deep Learning Models for Power Quality Disturbance Classification / J.M. Sierra-Fernández [et al.] // Applied Sciences. — 2025. — № 15(19). — P. 10602.
6. Zühlke MM. TCR: topologically consistent reweighting for XGBoost in regression tasks / M.M. Zühlke, D. Kudenko // Mach Learn. — 2025. — Vol. 114. — Art. 108. — DOI: 10.1007/s10994-024-06704-x
7. Zhuravlev V.V. Obzor razlichnikh algoritmov mashinnogo obucheniya v zadachakh obnaruzheniya i ispravleniya oshibok dannikh [Review of Various Machine Learning Algorithms in Data Error Detection and Correction Tasks] / V.V. Zhuravlev // Universum: tekhnicheskie nauki [Universum: Technical Sciences]. — 2024. — № 8(125). — DOI: 10.32743/UniTech.2024.125.8.18127. [in Russian]