

ПРИМЕНЕНИЕ МЕТОДОВ МАШИННОГО ОБУЧЕНИЯ ДЛЯ ПОВЫШЕНИЯ КАЧЕСТВА ТЕСТОВ

Р. Р. Минюков¹ [0009-0008-1264-7107], М. М. Абрамский² [0000-0003-3063-8948]

^{1, 2}Казанский (Приволжский) федеральный университет, г. Казань, 420008, Россия

¹ramil.minyukov@gmail.com, ²mabramsk@kpfu.ru

Аннотация

Работа посвящена применению методов машинного обучения для повышения качества тестов. Проведен обзор предметной области и реализованы два метода повышения качества: поиск похожих вопросов и оценка качества дистракторов. Первый включает тестирование пяти моделей трансформеров для получения векторного представления текста и шесть алгоритмов кластеризации. Второй метод основан на использовании тех же моделей трансформеров совместно с тремя алгоритмами классификации. Результаты экспериментов показали высокую эффективность предложенных решений при решении обеих задач.

Ключевые слова: анализ тестовых вопросов, дистракторы, машинное обучение, прохождение тестов, тесты, повышение качества тестов.

ВВЕДЕНИЕ

Рядом исследователей поднимается вопрос выбора валидной подборки вопросов, которые входят в конкретный тест, а также в целом повышения качества тестовых заданий (см., например, [1–3]).

Особое внимание уделяется вопросам автоматического анализа тестовых заданий. Исследования [4, 5] акцентируют внимание на вопросах выбора релевантных вопросов, а также повышения общей надежности и валидности тестов.

В настоящем исследовании была поставлена цель разработки методов и программных средств, направленных на повышение качества тестирования за счет анализа и реализации моделей искусственного интеллекта (ИИ), применяемых при автоматизированном исправлении заданий.

Статья организована следующим образом:

- в первом разделе проведен анализ предметной области и дан обзор существующих моделей и подходов к применению моделей и технологий ИИ в оценке заданий;
- во втором разделе рассмотрены и реализованы методы поиска похожих вопросов на основе их семантической близости;
- в третьем разделе речь рассмотрена реализация метода определения качества дистракторов в заданиях со множественным выбором.

1. СУЩЕСТВУЮЩИЕ ПОДХОДЫ К ПОВЫШЕНИЮ КАЧЕСТВА ТЕСТОВ

На данный момент существует несколько подходов к анализу тестовых заданий. Широко используемый из них основан на статистических показателях [5–7], к которым можно отнести:

- индекс легкости тестового задания;
- среднеквадратичное отклонение;
- индекс дифференциации;
- эффективность дистракторов;
- коэффициент угадывания;
- оценка надежности по методу расщепления на эквивалентные половины.

Статистические параметры не учитывают смысловую часть заданий, поэтому необходимо использовать технологии ИИ для нахождения семантических связей между элементами теста.

Одним из способов повышения качества заданий за счет методов ИИ является поиск похожих вопросов, снижающих эффективность теста, поскольку они могут дублировать проверяемые знания, создавая лишнюю когнитивную нагрузку [8]. Одним из подходов является кластеризация тестовых вопросов, которая обсуждается, например, в работах [9–11].

В [12] затронуты вопросы анализа качества дистракторов. Известные данные показывают, что опытные преподаватели редко создают более двух валидных дистракторов. Это подтверждает необходимость регулярного анализа и удаления неэффективных отвлекающих вариантов.

Еще одним из подходов является автоматический поиск скрытых подсказок, содержащихся в формулировках вопросов или вариантах ответов, при котором один вопрос может содержать информацию для ответа на другой [13]. Автоматизация такого подхода значительно повышает эффективность проверки и повышения качества тестовых заданий.

В результате обзора исследований, посвященных использованию методов и подходов искусственного интеллекта, выделены алгоритмы, которые способствуют повышению качества академических тестов.

2. ПОИСК ПОХОЖИХ ВОПРОСОВ

Для сравнения различных алгоритмов определения похожих заданий требуется предварительная подготовка данных.

Для английского языка известен широко используемый набор данных Quora Question Pairs [14]. Он содержит пары вопросов с метками, указывающими на то, являются ли вопросы дубликатами. Этот набор данных применяется в задачах определения семантической близости и поиска похожих вопросов.

Для русского языка имеется аналогичный по структуре набор данных Quora question pairs russian [15]. Он также содержит пары предложений и бинарную метку схожести.

Для тестирования алгоритмов данные из обоих наборов были предварительно агрегированы в группы фиксированного размера. В рамках проведенного эксперимента каждая группа состояла из 10 вопросов, внутри которых в среднем содержалось от 2 до 6 схожих по смыслу формулировок. Такая структура позволяла оценивать способность моделей находить дубликаты в ограниченном контексте.

Для удобства обработки данных и их подготовки к обучению разработаны вспомогательные классы, обеспечивающие удобное обращение с данными. Эти классы содержат методы, которые возвращают следующие данные:

- уникальные идентификаторы вопросов;
- список вопросов;
- список групп вопросов, схожих между собой.

В рамках проведенного исследования для получения векторного представления вопросов были выбраны модели Sentence Transformers, взятые с платформы Hugging Face [16]:

- all-MiniLM-L6-v2 [17];
- sentence-transformers/all-mpnet-base-v2 [18];
- sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2 [19];
- Alibaba-NLP/gte-multilingual-base [20];
- cointegrated/rubert-tiny2 [21].

Выбор этих моделей обусловлен их различной архитектурой, размерностью и способностью обрабатывать тексты на нескольких языках. Некоторые из них, например sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2 и Alibaba-NLP/gte-multilingual-base, обучены на многоязычных наборах данных и способны сохранять семантические связи вне зависимости от языка входной последовательности. Это особенно полезно для задач на русском и английском языках одновременно.

В качестве алгоритмов и методов кластеризации были выбраны следующие алгоритмы:

- косинусное сходство с пороговыми значениями 0.7, 0.8 и 0.9;
- К-средних (K-Means);
- иерархическая кластеризация (Agglomerative Clustering);
- гауссовы смеси (Gaussian Mixture Models);
- метод Сопору в сочетании с алгоритмом К-средних;
- разделяющий (бисекционный) алгоритм К-средних (Bisecting K-Means).

Для количественной оценки качества были выбраны следующие метрики: Adjusted Rand Index (ARI) и нормализованная взаимная информация (NMI).

Результаты тестирования представлены в табл. 1 и 2. Наилучшие результаты показали следующие комбинации:

- Alibaba-NLP/gte-multilingual-base и косинусное сходство;
- Alibaba-NLP/gte-multilingual-base и К-средних;
- Alibaba-NLP/gte-multilingual-base и иерархическая кластеризация;
- Alibaba-NLP/gte-multilingual-base и гауссовы смеси.

Таблица 1. Результаты тестирования моделей Sentence Transformers (по языку тестирования и порогу схожести)

Модель Sentence Transformers		all-MiniLM-L6-v2		all-mpnet-base-v2		gte-multilingual-base		paraphrase-multilingual MiniLM-L12-v2		rubert-tiny2	
	Язык	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.
Порог схожести	Метрика тестирования										
0.70	Скорректированный индекс Рэнда (ARI)	0.68	0.51	0.72	0.51	0.71	0.64	0.68	0.59	0.44	0.56
0.80	Скорректированный индекс Рэнда (ARI)	0.70	0.49	0.73	0.51	0.71	0.64	0.71	0.59	0.44	0.56
0.90	Скорректированный индекс Рэнда (ARI)	0.67	0.49	0.71	0.52	0.72	0.65	0.69	0.61	0.46	0.57
0.70	Нормализованная взаимная информация (NMI)	0.97	0.95	0.97	0.95	0.97	0.97	0.97	0.96	0.90	0.96
0.80	Нормализованная взаимная информация (NMI)	0.97	0.95	0.98	0.95	0.97	0.97	0.97	0.97	0.90	0.96
0.90	Нормализованная взаимная информация (NMI)	0.97	0.95	0.97	0.95	0.97	0.97	0.97	0.97	0.91	0.96

Таблица. 2. Результаты тестирования моделей Sentence Transformers в зависимости от языка тестирования и алгоритма кластеризации

Модель Sentence Transformers		all-MiniLM-L6-v2		all-mpnet-base-v2		gte-multilingual-base		paraphrase-multilingual-MiniLM-L12-v2		rubert-tiny2	
	Язык	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.	Англ.	Рус.
Скорректированный индекс Рэнда (ARI)	Bisecting K-средних	0.04	0.03	0.03	0.03	0.04	0.03	0.03	0.03	0.04	0.03
	K-средних	0.53	0.42	0.53	0.40	0.53	0.52	0.52	0.48	0.51	0.49
	Сапору + K-средних	0.61	0.02	0.62	0.02	0.34	0.28	0.09	0.20	0.00	0.07
	Гауссовы смеси	0.53	0.42	0.53	0.40	0.53	0.52	0.52	0.48	0.51	0.49
	Иерархическая кластеризация	0.53	0.39	0.52	0.36	0.52	0.50	0.51	0.46	0.50	0.47
Нормализованная взаимная информация (NMI)	Bisecting K-средних	0.41	0.39	0.41	0.38	0.42	0.40	0.40	0.37	0.40	0.40
	K-средних	0.92	0.89	0.92	0.87	0.92	0.92	0.92	0.90	0.92	0.92
	Сапору + K-средних	0.95	0.36	0.95	0.40	0.87	0.85	0.96	0.96	0.01	0.66
	Гауссовы смеси	0.92	0.89	0.92	0.87	0.92	0.92	0.92	0.90	0.92	0.92
	Иерархическая кластеризация	0.92	0.85	0.92	0.79	0.92	0.92	0.91	0.89	0.92	0.91

3. РЕАЛИЗАЦИЯ МЕТОДА ОЦЕНКИ ДИСТРАКТОРОВ

Для реализации метода, предназначенного для определения эффективности дистракторов, был выбран набор данных «Многоязычное измерение массового многозадачного понимания языка» (Multilingual Measuring Massive Multitask Language Understanding, MMLU) [22]. Его выбор обусловлен следующими факторами.

Во-первых, этот набор содержит более 50 тем, что обеспечивает разнообразие вопросов и их тематик. Во-вторых, набор поддерживает мультиязычность: вопросы и ответы переведены сразу на несколько языков, в том числе на русский. Это позволяет обучать модель сразу на нескольких языках. В-третьих, MMLU имеет удобную структуру заданий: каждый вопрос оформлен в виде множественного выбора с одним правильным ответом и несколькими дистракторами.

Для создания метода было необходимо сгенерировать отвлекающие варианты ответов: был разработан специализированный алгоритм их автоматизированного построения, применимый для русского и английского языков. Отличался только процесс выбора синонимов, где использовались различные лексические базы слов.

Общий алгоритм заключается в следующем. На первом этапе определялся тип правильного ответа: одиночное слово, числовое значение или короткое выражение. В случае, когда правильным ответом является слово, для генерации использованы его синонимы, случайные числовые значения, а также отдельные слова, случайно выбранные из вопроса. Если ответ представляет собой число, то в его качестве подбираются значения, гораздо большие или меньшие исходного значения, либо случайные ответы из общего набора ответов. В ситуации, когда правильным ответом является выражение, дистракторы формируются за счет выбора случайных ответов из общего набора.

Различие в построении отвлекающих вопросов между русским и английским языками заключается исключительно в этапе поиска синонимов. Для русского языка синонимы извлекались из лексической базы `wiki_ru_wordnet` [23], тогда как для английского использовалась база `WordNet`, доступная через библиотеку `NLTK` [24]. Все остальные этапы являются идентичными для разных языков, что обеспечивает единообразие подхода при работе с многоязычными данными.

Для подготовки данных к обучению были использованы пять различных моделей получения векторного представления слов, аналогичные поиску похожих вопросов. Каждая модель применялась для вычисления косинусного сходства для пар «вопрос – дистрактор» и «правильный ответ – дистрактор». Полученные метрики позволяют количественно оценить, насколько семантически близки элементы задания.

Для каждого дистрактора рассчитывались следующие показатели:

- косинусное сходство между векторными представлениями дистрактора и правильного ответа;
- косинусное сходство между векторными представлениями дистрактора и формулировкой вопроса;
- длина дистрактора, измеряемая количеством слов;
- разность в длине между дистрактором и правильным ответом.

На этапе обучения были протестированы три модели бинарной классификации: логистическая регрессия (Logistic Regression), метод опорных векторов (Support Vector Classifier, SVC) и метод случайного леса (Random Forest). Выбор моделей обусловлен их широкой распространенностью в задачах бинарной классификации.

Использование названных моделей обосновано их применением в аналогичных задачах анализа качества дистракторов. Например, логистическая регрессия и метод случайного леса применялись в работе [25] в качестве базового классификатора и метода ранжирования потенциальных дистракторов.

Для всех моделей зафиксированы следующие параметры:

- Logistic Regression – параметры по умолчанию;
- SVC – использовались StandardScaler и gamma, равная auto;
- Random Forest – параметр n_estimators, равный 100.

Для оценки качества моделей использованы следующие метрики:

- Accuracy – доля правильно классифицированных примеров;
- Macro-F1 – среднее гармоническое значение precision и recall по всем классам без учета их веса;
- Weighted-F1 – взвешенное среднее F1-score, учитывающее несбалансированность классов.

Результаты тестирования моделей представлены на рис. 1–3.

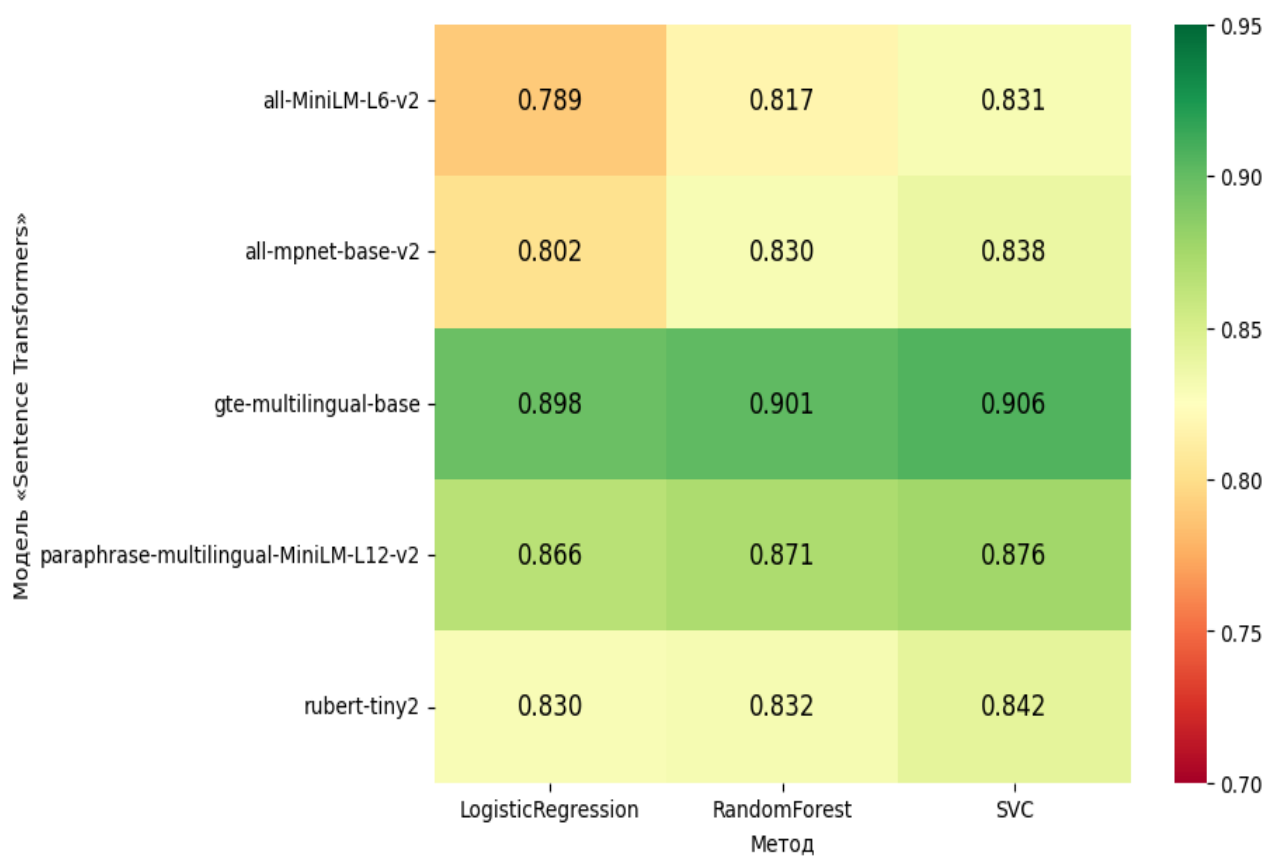


Рис. 1. Сравнение моделей Sentence Transformers по метрике Accuracy для методов классификации

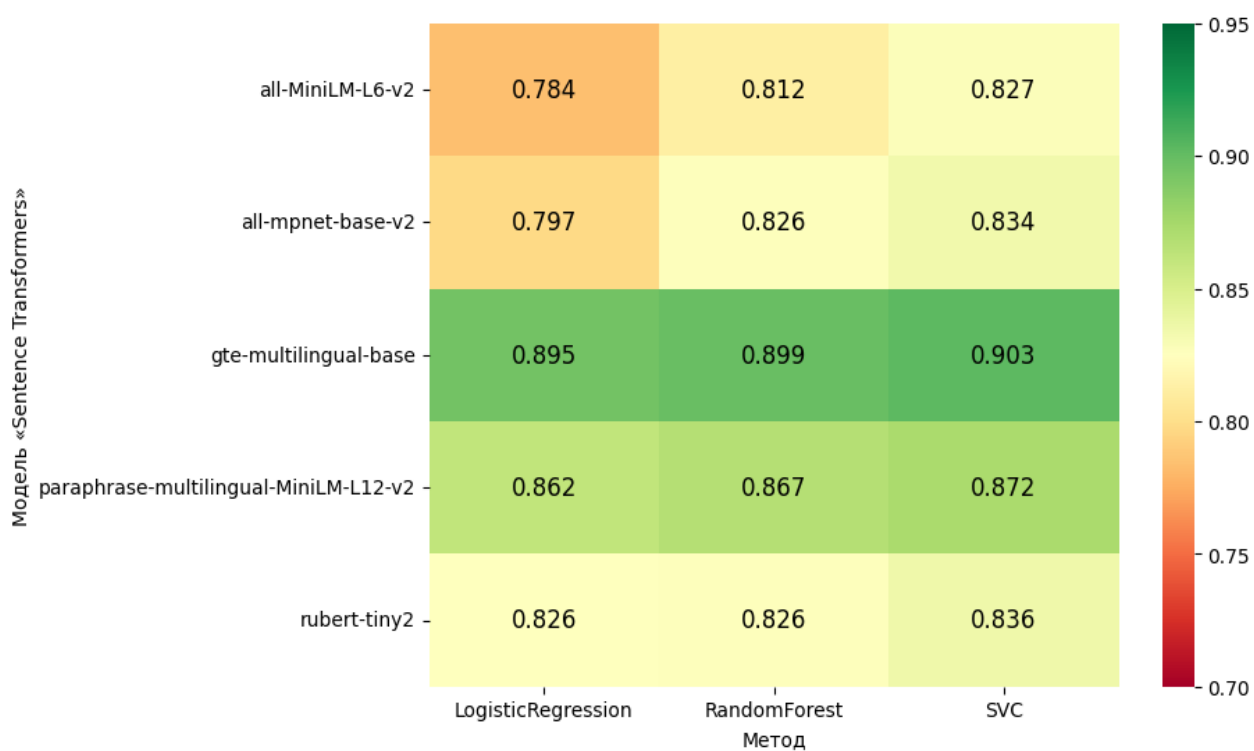


Рис. 2. Сравнение моделей Sentence Transformers по метрике Macro-F1 для методов классификации

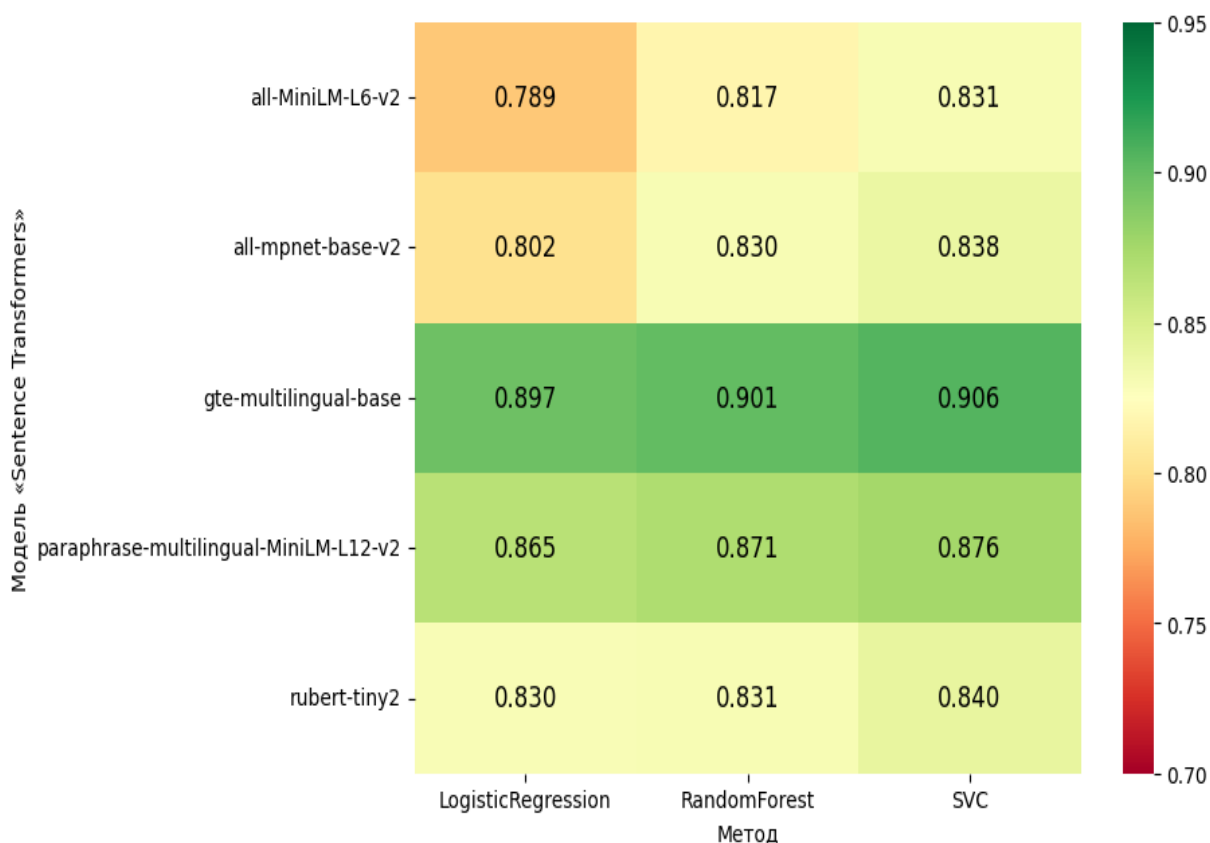


Рис. 3. Сравнение моделей Sentence Transformers по метрике Weighted-F1 для методов классификации

Результаты показали, что наилучших показателей удалось достичь, используя модель опорных векторов в сочетании с моделью получения эмбеддингов Alibaba-NLP/gte-multilingual-base.

ЗАКЛЮЧЕНИЕ

Таким образом, в ходе экспериментов установлено, что сочетания моделей Sentence Transformers и алгоритмов кластеризации способны эффективно находить семантические связи в формулировках вопросов, а сочетание Sentence Transformers и метода опорных векторов – определять качество дистракторов. Это позволяет сократить количество ручных проверок и автоматизировать процесс повышения качества тестов.

Дальнейшая работа направлена на разработку метода определения подсказок в формулировках вопросов и вариантах ответов.

СПИСОК ЛИТЕРАТУРЫ

1. Чельшкова М.Б. Теория и практика конструирования педагогических тестов. М.: Логос, 2002. Т. 432. С. 5.
2. Hrich N., Azekri M., Khaldi M. An ai educational tool for detecting redundancy in distractors and items within multiple-choice tests // INTED2024 Proceedings. IATED, 2024. P. 6454–6458.
3. Аванесов В.С. Теория и практика педагогических измерений. ЦТ и МКО УГТУ-УПИ, 2005.
4. Brusilovsky P., Miller P. Web-based testing for distance education. In: P. DeBra and J. Leggett (Eds.) Proceedings of WebNet'99, World Conference of the WWW and Internet, Honolulu, HI, Oct. 24–30, 1999. AACE, P. 149–154.
5. Толстобров А.П., Коржик И.А. Возможности анализа и повышения качества тестовых заданий при использовании сетевой системы управления обучением Moodle // Вестник ВГУ. 2008. Т. 2. С. 100–106.
6. Алпацкая Е.В., Бубнов Н.В., Минченков А.В. Дифференцирующая способность тестовых материалов для оценки качества обучения // Ученые записки университета им. П.Ф. Лесгафта. 2015. № 11 (129). С. 9–14.
7. Анастаси А. Психологическое тестирование. Питер, 2009.
8. Lord F.M. Applications of item response theory to practical testing problems. Routledge, 2012.
9. AlMahmoud R.H., Alian M. The effect of clustering algorithms on question answering // Expert Systems with Applications. 2024. Vol. 243. Article number 122959.
10. Zhang W.N. et al. A topic clustering approach to finding similar questions from large question and answer archives // PloS one. 2014. Vol. 9, No. 3. Article number e71511.
11. Alian M., Al-Naymat G. Questions clustering using canopy-K-means and hierarchical-K-means clustering // International Journal of Information Technology, 2022. Vol. 14, No. 7. P. 3793–3802.
12. Tarrant M., Ware J., Mohammed A. M. An assessment of functioning and non-functioning distractors in multiple-choice questions: a descriptive analysis // BMC medical education. 2009. Vol. 9. P. 1–8.

13. Moore S. et al. An automatic question usability evaluation toolkit // International Conference on Artificial Intelligence in Education. Cham: Springer Nature Switzerland, 2024. P. 31–46.
 14. First Quora Dataset Release: Question Pairs // Quora, 2017.
URL: <https://www.quora.com/q/quoradata/First-Quora-Dataset-Release-Question-Pairs> (дата обращения: 30.04.2025).
 15. Quora Question Pairs Russian // Платформа Kaggle, 2022.
URL: <https://www.kaggle.com/datasets/loopdigga/quora-question-pairs-russian> (дата обращения: 30.04.2025).
 16. Hugging Face – The AI Community [Электронный ресурс].
URL: <https://huggingface.co/> (дата обращения: 30.04.2025).
 17. All-MiniLM-L6-v2 // Платформа Hugging Face.
URL: <https://huggingface.co/sentence-transformers/all-MiniLM-L6-v2> (дата обращения: 30.04.2025).
 18. Sentence-transformers/all-mpnet-base-v2 // Платформа Hugging Face.
URL: <https://huggingface.co/sentence-transformers/all-mpnet-base-v2> (дата обращения: 30.04.2025).
 19. Paraphrase-multilingual-MiniLM-L12-v2 // Платформа Hugging Face.
URL: <https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2> (дата обращения: 30.04.2025).
 20. Gte-multilingual-base // Платформа Hugging Face.
URL: <https://huggingface.co/Alibaba-NLP/gte-multilingual-base> (дата обращения: 30.04.2025).
 21. Rubert-tiny2 // Платформа Hugging Face.
URL: <https://huggingface.co/cointegrated/rubert-tiny2> (дата обращения: 30.04.2025).
 22. Набор данных MMLU // Платформа Hugging Face.
URL: https://huggingface.co/datasets/alexandrainst/m_mmlu (дата обращения: 30.04.2025).
 23. Wiki-ru-wordnet документация //
URL: <https://wiki-ru-wordnet.readthedocs.io/en/latest/> (дата обращения: 30.04.2025).
-

24. Библиотека для обработки естественного языка NLTK // URL: <https://www.nltk.org/> (дата обращения: 30.04.2025).

25. Liang C. et al. Distractor generation for multiple choice questions using learning to rank // Proceedings of the thirteenth workshop on innovative use of NLP for building educational applications. 2018. Article number 28490.

USING MACHINE LEARNING TO ENHANCE TEST QUALITY

R. R. Miniukov¹ [0009-0008-1264-7107], M. M. Abramskiy² [0000-0003-3063-8948]

^{1, 2}Kazan (Volga Region) Federal University, Kazan, 420061, Russia

¹ramil.minyukov@gmail.com, ²mabramsk@kpfu.ru

Abstract

This study focuses on the application of machine learning methods to improve the quality of test items. The research includes a review of the subject area and the implementation of two enhancement methods: similar question retrieval and distractor quality assessment. The first method involves testing five transformer-based models for generating text embeddings and six clustering algorithms. The second method uses the same transformer models in combination with three classification algorithms. Experimental results demonstrated the high effectiveness of the proposed approaches in solving both tasks.

Keywords: *test item analysis, distractors, examination process, assessments, test quality improvement.*

REFERENCES

1. Chelyshkova M.B. Teoriya i praktika konstruirovaniya pedagogicheskikh testov. Moscow: Logos, 2002. 432 s.
 2. Hrich N., Azekri M., Khaldi M. An ai educational tool for detecting redundancy in distractors and items within multiple-choice tests // INTED2024 Proceedings. IATED, 2024. P. 6454–6458.
 3. Avanesov V.S. Teoriya i praktika pedagogicheskikh izmereniy. TsT i MKO UGTU-UI, 2005.
-

4. *Brusilovsky P., Miller P.* Web-based testing for distance education. In: P. DeBra and J. Leggett (Eds.) *Proceedings of WebNet'99, World Conference of the WWW and Internet*, Honolulu, HI, Oct. 24–30, 1999. AACE, P. 149–154.
5. *Tolstobrov A.P., Korzhik I.A.* Vozmozhnosti analiza i povysheniya kachestva testovykh zadaniy pri ispol'zovanii setevoy sistemy upravleniya obucheniem Moodle // *Vestnik VGU*. 2008. T. 2. S. 100–106.
6. *Alpatskaya E.V., Bubnov N.V., Minchenkov A.V.* Differentsiruyushchaya sposobnost' testovykh materialov dlya otsenki kachestva obucheniya // *Uchyonyye zapiski universiteta im. P. F. Lesgafta*. 2015. No. 11 (129). S. 9–14.
7. *Anastaszi A.* Psikhologicheskoe testirovanie. SPb.: Piter, 2007.
8. *Lord F.M.* Applications of item response theory to practical testing problems. Routledge, 2012.
9. *AlMahmoud R.H., Alian M.* The effect of clustering algorithms on question answering // *Expert Systems with Applications*. 2024. Vol. 243. Article number 122959.
10. *Zhang W.N. et al.* A topic clustering approach to finding similar questions from large question and answer archives // *PloS one*. 2014. Vol. 9, No. 3. Article number e71511.
11. *Alian M., Al-Naymat G.* Questions clustering using canopy-K-means and hierarchical-K-means clustering // *International Journal of Information Technology*, 2022. Vol. 14, No. 7. P. 3793–3802.
12. *Tarrant M., Ware J., Mohammed A. M.* An assessment of functioning and non-functioning distractors in multiple-choice questions: a descriptive analysis // *BMC medical education*. 2009. Vol. 9. P. 1–8.
13. *Moore S. et al.* An automatic question usability evaluation toolkit // *International Conference on Artificial Intelligence in Education*. Cham: Springer Nature Switzerland, 2024. P. 31–46.
14. First Quora Dataset Release: Question Pairs // Quora, 2017.
URL: <https://www.quora.com/q/quoradata/First-Quora-Dataset-Release-Question-Pairs> (last access: 30.04.2025).
15. Quora Question Pairs Russian // Kaggle, 2022.
URL: <https://www.kaggle.com/datasets/loopdigga/quora-question-pairs-russian> (last access: 30.04.2025).

16. Hugging Face – The AI Community [Электронный ресурс].
URL: <https://huggingface.co/> (last access: 30.04.2025).
17. All-MiniLM-L6-v2 // Hugging Face.
URL: <https://huggingface.co/sentence-transformers/all-MiniLM-L6-v2> (last access: 30.04.2025).
18. Sentence-transformers/all-mpnet-base-v2 // Платформа Hugging Face
URL: <https://huggingface.co/sentence-transformers/all-mpnet-base-v2> (last access: 30.04.2025).
19. Paraphrase-multilingual-MiniLM-L12-v2 // Hugging Face.
URL: <https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2> (last access: 30.04.2025).
20. Gte-multilingual-base // Hugging Face.
URL: <https://huggingface.co/Alibaba-NLP/gte-multilingual-base> (last access: 30.04.2025).
21. Rubert-tiny2 // Платформа Hugging Face.
URL: <https://huggingface.co/cointegrated/rubert-tiny2> (last access: 30.04.2025).
22. Biblioteka dlya obrabotki estestvennogo yazyka
URL: <https://www.nltk.org/> (last access: 30.04.2025).
23. *Liang C. et al.* Distractor generation for multiple choice questions using learning to rank // Proceedings of the thirteenth workshop on innovative use of NLP for building educational applications, 2018. P. 284–290.

СВЕДЕНИЯ ОБ АВТОРАХ



МИНЮКОВ Рамиль Радикович – студент магистратуры Института информационных технологий и интеллектуальных систем Казанского (Приволжского) федерального университета.

Ramil Radikovich MINIUKOV – Master's student at the Institute of Information Technology and Intelligent Systems, Kazan Federal University.

email: ramil.minyukov@gmail.com

ORCID: 0009-0008-1264-7107



АБРАМСКИЙ Михаил Михайлович – директор Института информационных технологий и интеллектуальных систем Казанского (Приволжского) федерального университета.

Mikhail Mikhailovich ABRAMSKIY – Director at the Institute of Information Technology and Intelligent Systems, Kazan Federal University.

email: mabramsk@kpfu.ru

ORCID: 0000-0003-3063-8948

Материал поступил в редакцию 5 мая 2025 года