

Wikipedia-based Datasets in Russian Information Retrieval Benchmark RusBEIR

Grigory Kovalev¹, Mikhail Tikhomirov¹, Pavel Mamaev¹, Olga Babina², and Natalia Loukachevitch¹

¹ Lomonosov Moscow State University, Russia

² South Ural State University, Chelyabinsk, Russia

Abstract. In this paper, we present a novel series of Russian information retrieval datasets constructed from the “Did you know...” section of Russian Wikipedia. Our datasets support a range of retrieval tasks, including fact-checking, retrieval-augmented generation, and full-document retrieval, by leveraging interesting facts and their referenced Wikipedia articles annotated at the sentence level with graded relevance. We describe the methodology for dataset creation that enables the expansion of existing Russian Information Retrieval (IR) resources. Through extensive experiments, we extend the RusBEIR research [8] by comparing lexical retrieval models, such as BM25, with state-of-the-art neural architectures fine-tuned for Russian, as well as multilingual models. Results of our experiments show that lexical methods tend to outperform neural models on full-document retrieval, while neural approaches better capture lexical semantics in shorter texts, such as in fact-checking or fine-grained retrieval. Using our newly created datasets, we also analyze the impact of document length on retrieval performance and demonstrate that combining retrieval with neural reranking consistently improves results. Our contribution expands the resources available for Russian information retrieval research and highlights the importance of accurate evaluation of retrieval models to achieve optimal performance. All datasets are publicly available at HuggingFace³. To facilitate reproducibility and future research, we also release the full implementation on GitHub⁴.

Keywords: Information retrieval · Wikipedia · Neural models · Lexical models

1 Introduction

The field of information retrieval (IR) has become a cornerstone of modern technology, enabling users to efficiently navigate the vast digital information landscape. Central to advancements in this field are high-quality datasets, which serve as benchmarks for evaluating retrieval systems, training machine learning models, and exploring contextual relationships within text. While English-language datasets such as BEIR [17], TREC⁵, and others have been instrumental in driving progress in IR, resources for other languages - particularly those with complex linguistic structures and rich cultural

³ <https://huggingface.co/collections/msu-rcc-lair/rusbeir-datasets-6720fb076978ab6a77f4f64c>

⁴ <https://github.com/kaengreg/rusBeIR>

⁵ <https://trec.nist.gov>

contexts - remain underdeveloped. The Russian language remains underrepresented in information retrieval (IR) research, with most existing evaluation datasets for Russian lacking in scale, diversity, and often relying on translated or multilingual sources. To address this gap, the creation of the RusBEIR benchmark [8] includes efforts to gather as many native Russian IR datasets as possible, as well as to develop new ones to expand and enrich the available resources for Russian IR.

Wikipedia, as a collaboratively edited, multilingual knowledge repository, offers a compelling foundation for constructing comprehensive information retrieval datasets. This idea is supported by popular datasets like Natural Questions (NQ) [5] and MIRACL [23], which are based on Wikipedia, along with several other examples. The Russian Wikipedia alone contains over 2 million articles, encompassing a broad range of topics, high-quality content, and structured metadata. However, its potential for IR research has been underutilized.

Из новых материалов

Знаете ли вы?

- **Мясной рулет** *(на илл.)* назван в честь **принцессы**, которая должна была стать императрицей.
- В состав **бальзамов Дон Кихота** и Д'Артаньяна входили одинаковые ингредиенты, но применяли их по-разному.
- Для своего **нового проекта создатель** «Ковбоя Бибопа» объединился с режиссёром «Джона Уика».
- Имя **англосаксонского короля**, возможно, произошло от слова «бог».
- Действующий **министр обороны Австрии** *(на илл.)* — первая женщина на этом посту.
- **Родоначальник** французского **хайку** начал с трёхстиший о **Первой мировой войне**.
- «Индианаполис Мотор Спидвей» находится не в Индианаполисе, а в **Спидуэе**.
- «**Духовный наследник**» *Bloodborne* стал эксклюзивом **новой консоли Nintendo**.
- Семья купила **королевскую смотровую башню** *(на илл.)* на берегу реки и в течение трёх поколений содержала в ней ресторан.
- **Французский лётчик** «Нормандии — Неман» получил награды трёх стран.
- Взяв немецкий псевдоним, **английскому композитору** пришлось доказывать, что он — это он.
- **Чешский доктор химии** готовил диссертацию в парижской лаборатории под руководством **двух нобелевских лауреатов**, позже работал в Лейпциге у **третьего**, а свою пражскую лабораторию предоставил **будущему четвёртому**.



Fig. 1. "Did you know..." section in Russian Wikipedia

In this paper, we suggest to utilize for information-retrieval evaluation the Wikipedia's "Did you know..." section (fig. 1). The section features interesting facts extracted from Wikipedia articles, with each fact including hyperlinks to the referenced articles. This section of Wikipedia, previously utilized primarily by scholars for academic research

on the nature of information interest and data attractiveness [13, 19, 10], has now been expanded to serve a broader audience.

Furthermore, the adaptable design of the corpus allows for the creation of diverse dataset sizes and configurations. For example, when the documents comprise complete articles, the dataset can function as a standard for full-document retrieval; conversely, if the corpus consists of individual sentences or their combinations, it is ideally suited for retrieval-augmented generation (RAG) tasks. A test sample of these datasets has already been incorporated into RusBEIR [8], an information-retrieval benchmark for Russian, but the datasets were not properly described. This paper examines the methodology behind the development of these datasets, explores their potential real-world applications, and extends the RusBEIR dataset collection while evaluating the impact of corpus size on model performance.

2 Related work

Throughout the time, Wikipedia has been a cornerstone for creating influential Natural Language Processing (NLP) datasets that drive progress in various language tasks.

SQuAD 2.0 [14] is a reading comprehension dataset with over 150,000 questions, including 50,000+ unanswerable ones, based on Wikipedia articles. Queries are sentence-length questions; answers are text spans or “unanswerable”. Crowdworkers created questions to test answer extraction and unanswerable question detection.

WikiQA [21] is a dataset for open-domain question answering with 3,047 Bing query log questions linked to Wikipedia pages. Contains 29,258 candidate answer sentences, 1,473 labeled correct. Queries are short phrases; documents are Wikipedia summary sentences. Data collected via query sampling and crowdsourced labeling.

Natural Questions (NQ) [9] contains approximately 323,000 real Google Search queries paired with Wikipedia pages. Annotators mark long and short answers or note their absence. Queries are user searches; documents are Wikipedia passages. Data reflects real-world question answering challenges.

WikiReading [6] is a dataset with 18M instances for predicting Wikidata values from Wikipedia articles. Queries are Wikidata properties (short phrases); documents are full articles. Data aligns Wikipedia text with Wikidata triples (entity, property, value).

DBpedia [1] is a structured knowledge base comprising over 100 million RDF triples from Wikipedia infoboxes, categories, and more. Queries are infobox properties (short phrases); documents are Wikipedia articles. Data extracted via automated parsing and ontology mapping.

FEVER [18] is a fact-checking dataset with 185,445 claims labeled SUPPORTED, REFUTED, or NOT ENOUGH INFO, verified against Wikipedia introductions. Queries are sentence-length claims; documents are introductory paragraphs. Claims created by modifying Wikipedia sentences.

As mentioned above, Wikipedia serves as a valuable source of data that enables the creation of variety of datasets. This resource has also been utilized for the Russian language. Below are examples of datasets developed using Russian Wikipedia data:

RuBQ [15] is a Russian knowledge base question answering (KBQA) dataset with 2,910 questions over Wikidata, including SPARQL queries. Queries are Russian ques-

Dataset	Queries	Queries Source	Corpus	Data Collection Method	Relevance Level
SQuAD2.0	Questions	Generated by crowdworkers	Wikipedia articles	Crowdworkers verify questions	2
WikiQA	Search queries	Real user search queries	Wikipedia summary sentences	Bing query logs, crowdsourced labeling	2
Natural Questions	Search queries	Real user search queries	Wikipedia pages	Google search queries, annotated answers	2
WikiReading	Wikidata properties	Structured Wikidata facts	Wikipedia articles	Automated Wikidata-Wikipedia alignment	2
DBpedia	Infobox properties	Structured Wikipedia infoboxes	Wikipedia infoboxes	Automated infobox parsing to RDF	-
FEVER	Claims	Mutated Wikipedia sentences	Wikipedia intros	Sentence mutation, annotator verification	2
RuBQ	Questions in Russian	Real user search queries	Wikipedia paragraphs	Search suggestions, annotated Wikidata links	2
WikiOmnia	Generated questions in Russian	Auto-generated QA pairs	Wikipedia summaries	Fully automated generative pipeline	2
SberQuAD	Questions in Russian	Generated by crowdworkers	Russian Wikipedia paragraphs	Crowdsourced annotation (Toloka)	2
wikifacts-articles	Assertions about facts	Authored "Did you know" Wikipedia facts	Full articles	Manual sentence-level annotation	3
wikifacts-para			Paragraphs		
wikifacts-sents			Sentences		
wikifacts-window			Sliding sentence chunks (2–6 sent.)		

Table 1. Summary of Wikipedia-based datasets. The newly introduced Wikipedia Interesting Facts series are highlighted.

tions from search suggestions; documents are Wikipedia paragraphs. Data collected via crowdsourcing and in-house annotation.

WikiOmnia [12] is a generative QA dataset comprising over 3.5 million verified question-answer pairs from Russian Wikipedia. Queries are auto-generated, sentence-length questions; documents are summary sections. Data created via automated pipeline using models like ruGPT-3 XL.

SberQuAD [3] is a Russian reading comprehension dataset with 50K+ paragraph-question-answer triples derived from Russian Wikipedia. Queries are sentence-length questions; answers are text spans without explicit start positions. Created following the SQuAD procedure, it shares a similar structure and task format but differs by having only one correct answer per question and no answer start indices. Data have been collected via crowdsourcing on Toloka.

All these datasets demonstrate the remarkable adaptability of Wikipedia as a data source, demonstrating its capacity to support a wide range of natural language processing tasks - from reading comprehension and question answering to semantic retrieval and fact-checking. However, there are two main problems with existing Wikipedia-based datasets. First, most of these datasets are primarily available in English. Second, the method of creating such datasets often relies on crowdworkers generating queries, as

previously mentioned. In this work, we introduce a new method for utilizing Wikipedia data to develop datasets for Russian.

3 Wikipedia interesting facts

We introduce a Wikipedia-based dataset, wikifacts-sents, derived from the “Did you know...” section (further “interesting facts”) and the corresponding articles referenced in these facts. This dataset is designed to support fact-checking tasks and experiments with retrieval-augmented contexts, with the following key features:

- Queries represent assertions about interesting facts, meaning the dataset can be categorized as fact-checking.
- Wikipedia facts are authored by Wikipedia contributors, eliminating the need for separate fact creation.
- These facts are specially crafted by modifying initial assertions, often by combining information from multiple sources; this transformation can be described as a series of operations applied to the original Wikipedia content.
- Each fact can have several supporting sentences: some sentences can contain full confirmation of the initial fact, while others may only confirm parts of the fact.
- Interesting facts are framed using descriptive mentions (such as “sportsman”, “city”, or “book”) but include references to specific Wikipedia articles and titles, enabling identification of particular persons, objects, or events. This structure allows for evaluating LLMs’ knowledge of the facts and experimenting with retrieval-augmented contexts.
- The primary dataset is sentence-based, with facts confirmed at the sentence level, but its relevance markup enables fact verification in larger text units, such as paragraphs, articles, or text chunks of varying lengths, by aggregating relevant sentences. This flexibility allowed us to develop derived datasets with documents of different sizes.

В Антарктиде взлётно-посадочные полосы для тяжёлых самолётов. имеют голубой цвет (на илл.).
[hypelink to wikipedia article](#)

Fig. 2. Examples of interesting facts from Russian Wikipedia

We conducted the annotation of sentences extracted from the Wikipedia articles, related to facts. In total, 55 university students were asked to identify relevant sentences. Each student was provided with a set of facts accompanied by referenced articles. The task was to read the articles associated with each fact and mark the sentences that appeared relevant, selecting from the predefined scoring options. A three-level relevance scoring system was employed: a score of 2 indicates that a sentence contains the complete information from the fact, a score of 1 shows that it includes only partial information, and a score of 0 denotes an irrelevant sentence.

A unique feature of our dataset is that crowdworkers are not required to create queries; their only task is to read the corresponding articles and mark relevant sentences. This is a much simpler task, making the work cheaper and allowing more data to be annotated for the same cost. Modern Large Language Models (LLMs) can also be used for the annotation process [11]. This can furthermore ease the process of creating similar datasets, making our dataset series even more actual and allowing to use the same methodology for the creation of similar wikipedia-based datasets in other languages.

At the end of the annotation process, we ended up with the wikifacts-sents dataset, containing 1,500,778 documents (sentences) and 5,433 queries, which serves as the foundation for creating multiple datasets derived from the “Did you know...” section.

4 Wikipedia interesting facts series of datasets

Based on the annotated wikifacts-sents data, additional seven wikifacts datasets have been produced. These datasets share a common set of queries but diverge in their corpus structures, providing a versatile platform for evaluation various retrieval approaches:

- **wikifacts-articles** dataset contains full articles referenced in the facts. Articles were systematically marked as relevant if at least one sentence within them was relevant to the corresponding fact. When multiple sentences showed different levels of relevance, the highest relevance score was chosen. This dataset features the longest documents across the series, making it an ideal resource for comprehensive full-document retrieval evaluation, particularly suited for testing the robustness of retrieval systems across extensive text corpora.
- **wikifacts-para** consists of paragraphs extracted from articles by splitting text at double newlines (“\n\n”). The documents in this variant are notably shorter than those in wikifacts-articles, yet they retain sufficient depth to serve as a valuable tool for evaluating full-text retrieval systems. This intermediate length allows for a balanced assessment of retrieval accuracy and efficiency.
- **wikifacts-window** is a specialized subset derived from the wikifacts-sents dataset, comprising sliding sentence chunks that range from 2 to 6 sentences. This configuration is specifically designed to evaluate Retrieval-Augmented Generation (RAG) models, which benefit from contextual text segments. Multiple versions of this dataset were generated with varying chunk sizes to assess the influence of text length on model performance.

To comply with the BEIR [17] and RusBEIR[8] formats, all datasets and their relevance annotations are provided as separate datasets on HuggingFace⁶. The corpora and queries are stored in JSONL format, while the relevance judgments (qrels) are available in TSV files within datasets that have a qrels suffix matching the names of the corresponding corpus and queries.

It should be noted that the preliminary wikifacts data (540 facts, marked with a v0 suffix), constructed in the same manner, have already been included in the RusBEIR

⁶ <https://huggingface.co/collections/msu-ccc-lair/rusbeir-datasets-6720fb076978ab6a77f4f64c>

Task (↓)	Dataset (↓)	Relevancy	Dev	Corpus	Avg. Word Lengths (D/Q)
Information-Retrieval	wikifacts-articles	3-level	5,433	12,848	2059.1 / 11.4
Fact Checking	wikifacts-para	3-level	5,433	176,364	150 / 11.4
Information-Retrieval	wikifacts-sents	3-level	5,433	1,500,778	17.6 / 11.4
Fact Checking	wikifacts-window_2	3-level	5,433	1,500,776	35.3 / 11.4
Fact Checking	wikifacts-window_3	3-level	5,433	1,500,774	52.9 / 11.4
Fact Checking	wikifacts-window_4	3-level	5,433	1,500,772	70.5 / 11.4
Fact Checking	wikifacts-window_5	3-level	5,433	1,500,770	88.1 / 11.4
Fact Checking	wikifacts-window_6	3-level	5,433	1,500,768	105.8 / 11.4

Table 2. Overview of datasets.

benchmark [8]. Now we introduce a complete version of the Wikipedia Interesting Facts dataset series, which currently comprises more than 5,000 facts.

To demonstrate that our datasets are not merely another collection of Wikipedia-based resources, but rather valuable benchmarks for information retrieval, we conduct an extensive series of experiments that highlight their practical utility. Building on this, we further extend the RusBEIR benchmark [8] by evaluating both previously used and newly developed models on these datasets. This broadens the range of assessed approaches and yields a more comprehensive leaderboard for Russian information retrieval.

5 Models

To thoroughly assess the effectiveness of information retrieval models using the Wikipedia Interesting Facts series of datasets, we employed traditional methods such as BM25, neural models based on bi-encoder and cross-encoder architectures, and a combination of traditional and neural approaches.

5.1 Preprocessing for BM25 model

The main baseline was established using the BM25 lexical model implemented in the Elasticsearch engine⁷, with the language analyzer disabled to avoid stemming, which is known to be less suitable for the Russian language. We specially preprocessed data to be used as input for BM25.

The text preprocessing method consists of the following steps:

1. **Lowercasing:** Converting all text to lowercase to ensure uniformity and eliminate case sensitivity.
2. **Punctuation and Special Character Removal:** Using regex to remove non-alphanumeric characters, leaving only letters, digits, and spaces to reduce noise.

⁷ <https://www.elastic.co/>

3. **Space Normalization:** Removing extra spaces and trimming leading or trailing whitespace.
4. **Tokenization:** Splitting text into individual words for processing.
5. **Lemmatization:** Using the PyMorphy3 Russian morphological tool [7] to convert words into their dictionary forms, reducing data dimensionality while preserving word meaning. This approach is particularly effective for the Russian language due to its rich morphology, as it avoids the inaccuracies that stemming introduces by truncating words without context.
6. **Stop Word Removal:** excluding overly frequent words that contribute little to the text content using the default stopwords list provided by the NLTK package⁸ [5], augmented with two Russian pronouns: “which” and “such”.

5.2 Neural baseline models

In our experiments, we used neural models based on bi-encoder and cross-encoder architectures. Bi-encoders were employed as dense retrievers: they generate vector representations for both queries and documents, and their similarity is calculated using cosine similarity. Cross-encoders were used as rerankers, where each query is processed jointly with a document to estimate the likelihood of the document’s relevance to the query. Rerankers are typically applied to the top-k documents retrieved by either neural or lexical models to refine the results. In general, combined retrieval approaches tend to outperform individual models.

The following pre-trained bi-encoders were used as dense retrievers:

LaBSE bi-encoder [4] was pre-trained with a translation ranking task. This allows to find sentence paraphrases in a single language or different languages.⁹

Multilingual E5 in three variants: large¹⁰, base¹¹ and small¹² [20]. The multilingual E5 model was developed using a vast multilingual dataset, employing a weakly supervised contrastive learning approach with InfoNCE loss. It was further refined on high-quality, labeled multilingual datasets specifically for retrieval tasks.

BGE-M3 model¹³ [2] was pre-trained on a large multilingual and cross-lingual unsupervised data, and subsequently fine-tuned on multilingual retrieval datasets using a custom loss function based on the InfoNCE loss function.

USER-BGE-M3¹⁴ is a sentence-transformer model designed to generate embeddings optimized for Russian. The model is initialized from the en-ru-BGE-M3 model¹⁵, a shrunk version of the BGE-M3 model, and then trained on the Russian datasets.

⁸ <https://www.nltk.org>

⁹ <https://huggingface.co/sentence-transformers/LaBSE>

¹⁰ <https://huggingface.co/intfloat/multilingual-e5-large>

¹¹ <https://huggingface.co/intfloat/multilingual-e5-base>

¹² <https://huggingface.co/intfloat/multilingual-e5-small>

¹³ <https://huggingface.co/BAAI/bge-m3>

¹⁴ <https://huggingface.co/deepvk/USER-bge-m3>

¹⁵ https://huggingface.co/TatonkaHF/bge-m3_en_ru

ru-en-RoSBERTa¹⁶ [16] is initialized from ruRoBERTa model [24]¹⁷ and then RoSBERTa embeddings were fine-tuned on Russian and English datasets.

FRIDA model¹⁸ is a large-scale, fine-tuned text embedding model that builds on the denoising approach originally introduced by T5, utilizing the encoder component of the FRED-T5 model[24]. Trained on a bilingual Russian-English corpus, it received additional fine-tuning to enhance its capabilities in tasks such as information retrieval, classification, and paraphrasing, with a focus on the Russian language.

Model	Based on	Parameters	Dim	Max input
Multilingual-E5-small	Multilingual-MiniLM	118M	384	512
Multilingual-E5-base	XLNet-RoBERTa-base	278M	768	512
Multilingual-E5-large	XLNet-RoBERTa-large	560M	1024	512
BGE-M3	BGE-M3	568M	1024	8192
USER-BGE-M3	BGE-M3	359M	1024	8192
LaBSE	LaBSE	471M	768	256
RoSBERTa	SBERT	404M	1024	512
FRIDA	FRED-T5	823M	1536	512
bge-reranker-v2-m3	BGE-M3	568M	1024	8192

Table 3. Model Specifications and Details

As shown in Table 3, different models support different maximum sequence lengths, which define their context windows. In our approach, we load both the model and tokenizer from HuggingFace, and set the `maxlen` parameter according to the model specification. Internally, the tokenizer produces at most `maxlen` tokens; if the input text is longer, it is truncated to this limit. As a result, the model processes only the first `maxlen` tokens, which are then used for subsequent computations.

Additionally, for state-of-the-art neural models such as mE5, BGE, and FRIDA, as well as the lexical BM25, we decided to use a reranker derived from one of these models to evaluate the effectiveness of such combinations. As demonstrated in [22], the `bge-reranker-v2-m3`¹⁹ achieves strong performance compared to other rerankers. Consequently, we selected it as the reranking model for our study.

To enable a fair and robust comparison of retrieval models, we adopt the Normalized Discounted Cumulative Gain at rank 10 (NDCG@10) as the primary evaluation metric. NDCG@10 is widely regarded in information retrieval for its ability to assess not only the presence of relevant documents among the top retrieved results, but also their ranking quality

¹⁶ <https://huggingface.co/ai-forever/ru-en-RoSBERTa>

¹⁷ <https://huggingface.co/ai-forever/ruRoberta-large>

¹⁸ <https://huggingface.co/ai-forever/FRIDA>

¹⁹ <https://huggingface.co/BAAI/bge-reranker-v2-m3>

6 Models' Performance on Wikipedia Interesting Facts

The series of Wikipedia Interesting Facts datasets have been designed in a manner that enables a thorough assessment of model performance across diverse scenarios, including full-document retrieval, fact-checking, and precise fine-grained retrieval. The underlying concept is to evaluate models within the same domain while varying document size and structure, thereby uncovering valuable insights.

Experiments with the selected models revealed important insights into the conditions under which neural models surpass lexical approaches, and vice versa. These findings offer a basis for understanding how model performance relates to specific document characteristics.

Dataset (↓)	Lexical	Dense								Re-ranking			
	BM25	mE5-large	mE5-base	mE5-small	BGE-M3	USER-BGE-M3	RoSBERTa	LaBSE	FRIDA	BM25+BGE	E5+BGE	BGE+BGE	FRIDA+BGE
wikifacts-sents	13.58	16.70	13.82	9.59	17.30	16.46	17.94	9.49	<u>20.79</u>	16.31	17.79	17.65	22.10
wikifacts-para	<u>44.57</u>	28.81	28.48	17.51	37.70	40.29	36.60	11.27	41.34	51.13	39.77	46.89	48.82
wikifacts-articles	<u>74.95</u>	60.62	56.21	51.69	64.11	66.02	59.80	27.25	60.54	74.97	70.39	71.22	69.31
Avg	<u>44.37</u>	35.38	32.84	26.26	39.70	40.92	38.11	16.00	40.89	47.47	42.65	45.25	46.74

Table 4. Performance comparison across different models. The best results for each dataset are in bold; the results of the best single models are underlined.

Analyzing the results of lexical and neural models on the Wikipedia Interesting Facts datasets (presented in Table 4), we observed that the lexical model BM25 performs increasingly well as document length increases. Notably, BM25 outperformed neural models on the "wikifacts-articles" and "wikifacts-para" datasets, which contain relatively long documents. This indicates a potential weakness of neural networks when processing longer documents: they may struggle to maintain consistent performance when compared to lexical approaches such as BM25.

7 Impact of Document Length on the Model Performance

Neural models are constrained by limitations such as a maximum input sequence length, which can significantly impact their performance. Evaluations conducted using the extended set of the models from the RusBEIR benchmark [8] revealed that neural models begin to encounter difficulties when processing longer documents, often failing to maintain consistent performance as document size increases.

To address this, we conducted an in-depth investigation into these capabilities, leveraging our newly developed series of wikifacts-window datasets. This dataset series is uniquely suited to provide a clearer understanding of the inherent limitations of neural models, particularly in handling texts of varying lengths. These findings offer critical insights that contribute to a more informed selection of models for specific tasks, ensuring optimal performance based on the nature of the documents involved.

7.1 Lexical models vs Neural models

Drawing on the observations highlighted in Section 6, we conducted a detailed exploration of the relationship between text length and the performance of the lexical model BM25 using our wikifacts-window datasets. This collection of datasets serves as an ideal platform for such an analysis, as it encompasses documents of diverse lengths, ranging from short sentence clusters to larger text chunks, enabling a comprehensive assessment of retrieval performance across different scales.

Dataset ()	Lexical	Dense								Re-ranking			
	BM25	mE5-large	mE5-base	mE5-small	BGE-M3	USER-BGE-M3	RoSBERTa	LaBSE	FRIDA	BM25+BGE	E5+BGE	BGE+BGE	FRIDA+BGE
wikifacts-sents	13.58	16.70	13.82	9.59	17.30	16.46	17.94	9.49	<u>20.79</u>	16.31	17.79	17.65	22.10
wikifacts-window_2	21.27	24.98	20.98	15.70	25.56	24.90	25.03	12.38	<u>28.84</u>	26.80	29.48	29.79	31.95
wikifacts-window_3	24.94	27.30	23.97	18.65	27.81	27.97	26.40	12.00	<u>30.21</u>	30.85	32.84	33.13	34.36
wikifacts-window_4	27.48	29.09	25.84	20.93	29.03	29.59	27.05	11.23	<u>30.71</u>	33.37	34.76	34.44	35.38
wikifacts-window_5	29.43	30.10	26.82	22.11	29.84	30.60	27.52	10.76	<u>31.22</u>	35.32	36.27	35.48	36.05
wikifacts-window_6	31.54	31.42	28.11	23.34	31.15	32.08	28.39	10.88	<u>32.09</u>	37.66	37.84	36.91	37.33

Table 5. Performance comparison across different models and Wikipedia Interesting Facts datasets. The best results for each dataset are in bold; the results of the best single models are underlined.

Table 5 shows that BM25 performs competitively with neural retrievers on longer texts, nearly matching their performance on the wikifacts-window_6 dataset, where it is only 1.7 percentage points behind the best single model, FRIDA. Notably, BM25 outperforms mE5-large and BGE-M3 by 1.6 percentage points on this dataset, underscoring the robustness of lexical models in extended contexts. In contrast, on shorter datasets like wikifacts-window_2 and wikifacts-window_3, which have an average word count of 44, BM25 falls behind top neural retrievers by more than 15 percentage points, due to neural models’ superior ability to capture semantic nuances in brief texts. However, as document length exceeds 80 words, the performance gap narrows significantly, reflecting the input length limitations of most neural models. Notably, FRIDA and USER-BGE-M3, both non-multilingual models, consistently outperform other neural and lexical models across all the longest datasets. Their superior performance likely comes from fine-tuning on the target language, which allows for a deeper understanding that improves retrieval accuracy across both short and long texts. These results indicate that although lexical models like BM25 are efficient for longer documents, fine-tuned, language-specific neural models such as FRIDA and USER-BGE-M3 deliver consistently strong results.

7.2 BGE max-length experiments

As previously mentioned in Table 3, BGE is a unique model compared to others presented, as it has the ability to process up to 8192 tokens, allowing it to perform well on datasets with a high average word count. Therefore we chose to explore how the model’s performance on the Wikipedia Interesting Facts datasets varies under different maximum length constraints. To mitigate domain-specific biases, we also incorporated datasets from RusBEIR that feature the longest texts across diverse domains.

Table 6 shows the results of experiments conducted with different maximum-length parameters of the input.

Model (→)	Max-length / batch-size									
	512 / 64		1024 / 64		2048 / 64		4096 / 32		8192 / 8	
Dataset (↓)	BGE-M3	USER-BGE-M3	BGE-M3	USER-BGE-M3	BGE-M3	USER-BGE-M3	BGE-M3	USER-BGE-M3	BGE-M3	USER-BGE-M3
wikifacts-articles	61.33	63.63	65.29	63.42	64.11	66.02	62.49	65.15	59.72	63.47
wikifacts-para	37.0	39.56	37.73	40.29	37.7	40.29	37.68	40.29	37.66	40.28
wikifacts-sents	17.32	16.39	17.32	16.4	17.3	16.46	17.3	16.38	17.26	16.33
wikifacts-window_2	25.58	24.95	25.58	24.95	25.56	24.9	25.58	24.95	25.57	24.9
wikifacts-window_3	27.83	27.97	27.83	27.98	27.81	27.97	27.82	27.98	27.83	27.96
wikifacts-window_4	29.04	29.57	29.03	29.59	29.03	29.59	29.03	29.59	29.02	29.57
wikifacts-window_5	29.84	30.6	29.83	30.6	29.84	30.6	29.82	30.6	29.84	30.61
wikifacts-window_6	31.15	32.08	31.15	32.08	31.15	32.09	31.15	32.08	31.15	32.08
rus-NFCorpus	30.69	30.33	30.83	30.26	30.86	30.28	30.78	30.21	30.82	30.28
rus-SciFact	61.84	58.33	62.32	58.21	62.42	58.25	62.33	58.25	62.35	58.25
rus-ArguAna	50.68	46.65	50.71	46.58	50.75	46.52	50.72	46.54	50.73	46.54
Ria-News	83.01	83.34	83.00	83.36	82.99	83.52	83.01	83.37	83.02	83.36

Table 6. Performance comparison across BGE-M3 models with different max-length

We observe a consistent improvement in quality as the max-length increases from 512 to 2048. In some cases, models with max-lengths set to 4096 or 8192 achieve slightly better performance than those with 2048. However, these gains are not substantial enough to justify the additional computational cost and complexity associated with processing longer sequences.

Based on results it was determined that the BGE setup with a max-length of 2048 is the optimal choice, striking a balance between quality and efficiency. It provides strong performance across datasets while maintaining manageable resource requirements, making it a practical and effective configuration.

8 Comparison to RusBEIR results

The test sample from the Wikipedia Interesting Facts datasets was incorporated into RusBEIR, and since that time, the corpus size has expanded tenfold, significantly increasing the complexity of information retrieval tasks for neural models. This comparison - evaluating the same task with a substantially larger corpus - provides valuable insights into the challenges and nuances of employing neural models for information retrieval tasks, particularly in terms of scalability and performance under varying data volumes.

The results in Table 7 demonstrate similar trends when comparing the performances of lexical, neural, and re-ranking models on the test version of Wikipedia Interesting Facts datasets (marked by v0 suffix) and on the extended version of the datasets.

Analyzing the complete version of the Wikipedia Interesting Facts datasets we can note, that the substantial corpus expansion significantly impacts all models, with performance reductions across the board due to increased retrieval complexity. While BM25 maintains strong performance on longer textual data, such as wikifacts-articles (74.95) and wikifacts-para (44.57), its relative effectiveness diminishes on shorter, sliding-window

Dataset (↓)	Lexical	Dense								Re-ranking			
	BM25	mE5-large	mE5-base	mE5-small	BGE-M3	USER-BGE-M3	RoSBERTa	LaBSE	FRIDA	BM25+BGE	E5+BGE	BGE+BGE	FRIDA+BGE
wikifacts-articles_v0	<u>84.28</u>	66.09	63.04	67.86	74.50	79.41	74.13	45.17	75.47	85.25	83.06	83.91	83.07
wikifacts-para_v0	<u>61.31</u>	50.15	49.51	34.71	54.55	57.53	50.66	14.78	55.17	66.61	59.95	63.76	63.85
wikifacts-sents_v0	33.64	35.90	30.75	22.57	37.59	34.90	40.59	25.79	<u>46.53</u>	39.96	38.53	39.20	49.41
wikifacts-window_2_v0	38.52	40.98	37.54	28.10	42.71	41.13	42.01	24.36	<u>47.23</u>	46.80	48.26	48.74	52.15
wikifacts-window_3_v0	41.93	44.84	42.38	32.44	44.80	44.34	42.02	22.17	<u>47.77</u>	50.23	51.49	52.30	53.30
wikifacts-window_4_v0	45.41	47.07	44.66	35.56	46.47	46.65	41.95	21.08	<u>47.94</u>	53.84	54.24	54.35	54.27
wikifacts-window_5_v0	48.72	49.93	47.18	37.70	48.53	49.22	43.83	21.13	49.13	57.45	57.38	57.02	55.87
wikifacts-window_6_v0	<u>51.88</u>	51.87	49.26	39.96	49.69	50.93	44.28	20.97	49.58	59.92	59.58	58.63	57.26
wikifacts-articles	<u>74.95</u>	60.62	56.21	51.69	64.11	66.02	59.80	27.25	60.54	74.97	70.39	71.22	69.31
wikifacts-para	<u>44.57</u>	28.81	28.48	17.51	37.70	40.29	36.60	11.27	41.34	51.13	39.77	46.89	48.82
wikifacts-sents	13.58	16.70	13.82	9.59	17.30	16.46	17.94	9.49	<u>20.79</u>	16.31	17.79	17.65	22.10
wikifacts-window_2	21.27	24.98	20.98	15.70	25.56	24.90	25.03	12.38	<u>28.84</u>	26.80	29.48	29.79	31.95
wikifacts-window_3	24.94	27.30	23.97	18.65	27.81	27.97	26.40	12.00	<u>30.21</u>	30.85	32.84	33.13	34.36
wikifacts-window_4	27.48	29.09	25.84	20.93	29.03	29.59	27.05	11.23	<u>30.71</u>	33.37	34.76	34.44	35.38
wikifacts-window_5	29.43	30.10	26.82	22.11	29.84	30.60	27.52	10.76	<u>31.22</u>	35.32	36.27	35.48	36.05
wikifacts-window_6	31.54	31.42	28.11	23.34	31.15	32.08	28.39	10.88	<u>32.09</u>	37.66	37.84	36.91	37.33

Table 7. Performance comparison across different models and datasets. The best results for each dataset are in bold; the results of the best single models are underlined.

datasets, losing dominance on datasets like wikifacts-window_5 and wikifacts-window_6. It should be noted that on wikifacts-window_6, BM25 still outperforms mE5-large and BGE-M3, although the performance difference notably decreases, highlighting increased competition from neural retrieval models within increased dataset size.

Fine-tuned neural models specifically trained on Russian language data, such as FRIDA and USER-BGE-M3, consistently outperform multilingual neural models including the mE5 series and BGE-M3 starting from wikifacts-window_3. FRIDA achieves the highest single-model scores across nearly all datasets, securing the leading position in individual neural model comparisons.

The combination of lexical or neural models with neural re-ranking demonstrates substantial improvements over individual retrieval approaches. Initially, the BM25+BGE combination delivered state-of-the-art results, particularly on longer texts, effectively combining lexical recall with semantic precision. However, a clear shift is observed in the expanded version of datasets, where the leadership of BM25+BGE notably declines for longer sliding-window datasets. Specifically, on wikifacts-window_5 and wikifacts-window_6, the combination mE5-large with the BGE-reranker emerges as the top-performing approach, outperforming the previously leading BM25+BGE. This shift indicates the enhanced ability of dense retrieval combined with neural re-ranking to handle increased complexity and lexical variability in larger datasets.

Overall, these findings underline the necessity of integrating powerful neural rerankers with appropriate retrieval models to optimize performance in large-scale information retrieval tasks.

9 Current leadership in RusBEIR

With the introduction of new challenging datasets into RusBEIR, along with newly developed models, we updated the benchmark leaderboard to examine potential shifts in model rankings.

Model (→)	Lexical	Dense								Re-ranking			
Dataset (↓)	BM25	mE5-large	mE5-base	mE5-small	BGE-M3	USER-BGE-M3	RoSBERTa	LaBSE	FRIDA	BM25+BGE	E5+BGE	BGE+BGE	FRIDA+BGE
rus-NFCorpus	<u>32.33</u>	30.96	26.90	26.79	30.86	30.28	27.24	18.53	29.40	34.83	33.18	32.46	32.38
rus-ArguAna	41.49	49.06	39.40	39.59	<u>50.75</u>	46.52	49.38	25.52	41.90	52.91	54.01	53.87	52.24
rus-SciFact	<u>65.60</u>	63.49	63.46	60.46	62.42	58.25	53.90	29.07	63.43	70.40	71.34	69.64	70.09
rus-SCIDOCs	13.99	13.47	12.09	10.60	<u>15.04</u>	14.46	14.43	8.17	12.78	15.31	15.98	16.21	15.33
rus-TREC-COVID	62.47	76.38	74.45	73.95	62.66	55.07	68.43	23.04	<u>82.42</u>	73.46	83.11	77.66	85.97
rus-FIQA	22.60	34.71	30.44	25.74	<u>38.16</u>	37.09	32.62	7.19	36.50	29.37	39.19	40.21	38.36
rus-Quora	61.34	80.18	78.62	74.99	<u>80.28</u>	79.87	68.73	72.39	74.54	71.01	81.85	81.59	80.38
rus-CQADupstack	24.49	32.08	28.43	27.31	<u>32.57</u>	31.40	27.85	21.36	27.69	28.86	33.78	33.90	31.44
rus-Touche	<u>30.59</u>	25.88	22.88	23.48	28.06	24.36	24.11	8.46	30.44	32.72	29.46	31.43	32.21
rus-MMARCO	15.25	<u>34.04</u>	30.27	29.07	29.51	27.92	20.16	9.06	33.55	24.12	36.95	34.52	36.33
rus-MIRACL	25.13	66.99	61.41	58.52	70.50	67.23	53.11	15.70	<u>71.91</u>	41.51	75.90	76.44	76.36
rus-XQuAD	96.19	<u>97.33</u>	95.84	95.66	95.97	95.63	93.90	69.77	93.49	98.85	98.97	98.97	98.24
rus-XQuAD-Sentences	82.36	<u>88.84</u>	86.37	85.41	86.91	85.42	83.20	75.33	86.89	89.93	92.08	91.69	91.60
rus-TyDi QA	35.80	59.41	55.91	55.23	58.34	57.86	52.06	28.05	<u>59.84</u>	50.12	66.20	65.78	65.84
SherQuad-retrieval	68.19	67.11	65.13	61.03	<u>68.26</u>	67.03	63.59	37.54	62.68	70.34	69.41	68.21	65.34
ruSciBench-retrieval	36.69	50.81	45.74	42.93	<u>55.85</u>	53.58	44.89	17.93	52.92	49.93	65.33	69.05	66.44
ru-facts	92.56	93.65	93.55	93.06	93.91	93.77	93.66	93.10	93.74	92.72	92.87	92.87	92.90
RuBQ	37.33	<u>74.11</u>	69.63	68.60	71.26	70.00	66.81	30.59	73.12	56.90	77.03	76.00	75.71
Ria-News	64.63	80.67	70.24	70.00	82.99	<u>83.52</u>	78.85	61.57	82.92	78.12	86.22	86.85	86.73
wikifacts-articles	<u>74.95</u>	60.62	56.21	51.69	64.11	66.02	59.80	27.25	60.54	74.97	70.39	71.22	69.31
wikifacts-para	<u>44.57</u>	28.81	28.48	17.51	37.70	40.29	36.60	11.27	41.34	51.13	39.77	46.89	48.82
wikifacts-sents	13.58	16.70	13.82	9.59	17.30	16.46	17.94	9.49	<u>20.79</u>	16.31	17.79	17.65	22.10
wikifacts-window_2	21.27	24.98	20.98	15.70	25.56	24.90	25.03	12.38	<u>28.84</u>	26.80	29.48	29.79	31.95
wikifacts-window_3	24.94	27.30	23.97	18.65	27.81	27.97	26.40	12.00	<u>30.21</u>	30.85	32.84	33.13	34.36
wikifacts-window_4	27.48	29.09	25.84	20.93	29.03	29.59	27.05	11.23	<u>30.71</u>	33.37	34.76	34.44	35.38
wikifacts-window_5	29.43	30.10	26.82	22.11	29.84	30.60	27.52	10.76	<u>31.22</u>	35.32	36.27	35.48	36.05
wikifacts-window_6	31.54	31.42	28.11	23.34	31.15	32.08	28.39	10.88	<u>32.09</u>	37.66	37.84	36.91	37.33
Avg.	43.58	50.67	47.22	44.52	50.99	49.90	46.88	28.06	<u>51.33</u>	50.66	55.63	55.66	55.89

Table 8. Performance comparison across different models and datasets. The best results for each dataset are in bold; the results of the best single models are underlined.

From Table 8, we confirm that newly added FRIDA continues to lead overall, consistently demonstrating strong performance both as a standalone model and when combined with reranking strategies.

As mentioned previously in Section 8, BM25 provides superior results on datasets with longer texts. The combination of BM25 and the BGE reranker performs comparably to single dense retrievers such as mE5-large and BGE-M3, which rank as the second and third best single models on our leaderboard. This further underscores that BM25 remains a robust and competitive model for information retrieval, while neural models excel in the context of shorter documents due to their ability to capture lexical semantics and provide deeper understanding.

Hybrid approaches that combine dense retrievers with rerankers after the retrieval stage show superior search quality, improving standalone results by an average of 9 percentage points, making them the best choice in terms of quality.

10 Conclusion

In this paper, we introduced and thoroughly analyzed the Wikipedia Interesting Facts series of datasets, specifically designed for Russian Information Retrieval tasks. Utilizing Wikipedia’s “Did you know...” section, we created diverse datasets tailored for various retrieval tasks, ranging from full-document retrieval to fact-checking and fine-grained retrieval. These new datasets exhibited unique flexibility through varying document structures and sizes, enabling a comprehensive evaluation of different retrieval approaches.

Our experiments provided insights into model performance, revealing that lexical models such as BM25 remain competitive and robust, especially in full-document retrieval scenarios. However, neural models demonstrated superior ability to capture meaning within shorter text units. Language-specific neural models fine-tuned for Russian, such as FRIDA and USER-BGE-M3, outperformed multilingual variants on datasets featuring longer texts, highlighting the importance of accurately selecting retrievers for specific tasks. Despite the superior performance of Russian-specific models on wiki-facts variants with larger sliding windows, their overall performance, as demonstrated in RusBEIR and Table 8, remains below that of multilingual variants.

Moreover, combining retrieval methods with neural reranking significantly enhanced performance, particularly in larger datasets. While the BM25+BGE combination initially established baselines for longer texts, further experiments on expanded datasets indicated that combinations involving dense neural retrievers, such as mE5-large with the BGE-reranker, surpassed previous leading models. For shorter datasets, the combination of FRIDA with the BGE-reranker demonstrated superior performance, surpassing mE5-large+BGE and BGE-M3+BGE. This underscores the importance of integrating neural rerankers to improve retrieval effectiveness.

Our novel dataset creation strategy not only expands the RusBEIR collection - facilitating more precise evaluation of information retrieval models - but also lays the groundwork for developing multilingual variants of these datasets. Our experiments indicate that evaluating models within the same domain under various scenarios can highlight critical nuances and ease the selection of optimal retrieval approaches.

Acknowledgments

The study was supported by the grant of the Russian Science Foundation No. 25-21-00206²⁰.

The research was carried out using the MSU-270 supercomputer of Lomonosov Moscow State University.

References

- [1] Sören Auer et al. “Dbpedia: A nucleus for a web of open data”. In: *international semantic web conference*. Springer. 2007, pp. 722–735.
- [2] Jianlv Chen et al. “Bge m3-embedding: Multi-lingual, multi-functionality, multi-granularity text embeddings through self-knowledge distillation”. In: *arXiv preprint arXiv:2402.03216* (2024).
- [3] Pavel Efimov et al. “Sberquad–russian reading comprehension dataset: Description and analysis”. In: *Experimental IR Meets Multilinguality, Multimodality, and Interaction: 11th International Conference of the CLEF Association, CLEF 2020, Thessaloniki, Greece, September 22–25, 2020, Proceedings 11*. Springer. 2020, pp. 3–15.

²⁰ <https://rscf.ru/project/25-21-00206/>

- [4] Fangxiaoyu Feng et al. “Language-agnostic BERT Sentence Embedding”. In: *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 2022, pp. 878–891.
- [5] Nitin Hardeniya et al. *Natural language processing: python and NLTK*. Packt Publishing Ltd, 2016.
- [6] Daniel Hewlett et al. “Wikireading: A novel large-scale language understanding task over wikipedia”. In: *arXiv preprint arXiv:1608.03542* (2016).
- [7] Mikhail Korobov. “Morphological analyzer and generator for Russian and Ukrainian languages”. In: *Analysis of Images, Social Networks and Texts: 4th International Conference, AIST 2015, Yekaterinburg, Russia, April 9–11, 2015, Revised Selected Papers 4*. Springer. 2015, pp. 320–332.
- [8] Grigory Kovalev et al. “Building Russian Benchmark for Evaluation of Information Retrieval Models”. In: *Proceedings of the International Conference “Dialogue*. Vol. 2025. 2025.
- [9] Tom Kwiatkowski et al. “Natural questions: a benchmark for question answering research”. In: *Transactions of the Association for Computational Linguistics 7* (2019), pp. 453–466.
- [10] Jingun Kwon et al. “Hierarchical trivia fact extraction from Wikipedia articles”. In: *Proceedings of the 28th International Conference on Computational Linguistics*. 2020, pp. 4825–4834.
- [11] Shengjie Ma et al. “Leveraging large language models for relevance judgments in legal case retrieval”. In: *arXiv preprint arXiv:2403.18405* (2024).
- [12] Dina Pisarevskaya and Tatiana Shavrina. “Wikiomnina: generative qa corpus on the whole russian wikipedia”. In: *arXiv preprint arXiv:2204.08009* (2022).
- [13] Abhay Prakash et al. “Did you know? mining interesting trivia for entities from wikipedia”. In: *Proceedings of the 24th International Conference on Artificial Intelligence*. 2015, pp. 3164–3170.
- [14] Pranav Rajpurkar, Robin Jia, and Percy Liang. “Know What You Don’t Know: Unanswerable Questions for SQuAD”. In: *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*. Association for Computational Linguistics. 2018.
- [15] Ivan Rybin et al. “RuBQ 2.0: an innovated Russian question answering dataset”. In: *The Semantic Web: 18th International Conference, ESWC 2021, Virtual Event, June 6–10, 2021, Proceedings 18*. Springer. 2021, pp. 532–547.
- [16] Artem Snegirev et al. “The Russian-focused embedders’ exploration: ruMTEB benchmark and Russian embedding model design”. In: *CoRR* (2024).
- [17] Nandan Thakur et al. “BEIR: A Heterogeneous Benchmark for Zero-shot Evaluation of Information Retrieval Models”. In: *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*. 2021. URL: <https://openreview.net/forum?id=wCu6T5xFjeJ>.
- [18] James Thorne et al. “FEVER: a Large-scale Dataset for Fact Extraction and VERification”. In: *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*. 2018, pp. 809–819.

- [19] David Tsurel et al. “Fun facts: Automatic trivia fact extraction from wikipedia”. In: *Proceedings of the Tenth ACM International Conference on Web Search and Data Mining*. 2017, pp. 345–354.
- [20] Liang Wang et al. “Multilingual e5 text embeddings: A technical report”. In: *arXiv preprint arXiv:2402.05672* (2024).
- [21] Yi Yang, Wen-tau Yih, and Christopher Meek. “Wikiqa: A challenge dataset for open-domain question answering”. In: *Proceedings of the 2015 conference on empirical methods in natural language processing*. 2015, pp. 2013–2018.
- [22] Xin Zhang et al. “mGTE: Generalized Long-Context Text Representation and Reranking Models for Multilingual Text Retrieval”. In: *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*. 2024, pp. 1393–1412.
- [23] Xinyu Zhang et al. “Miracl: A multilingual retrieval dataset covering 18 diverse languages”. In: *Transactions of the Association for Computational Linguistics* 11 (2023), pp. 1114–1131.
- [24] Dmitry Zmitrovich et al. “A Family of Pretrained Transformer Language Models for Russian”. In: *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*. 2024, pp. 507–524.

11 Response to Reviewers

11.1 Review 1

“Recommendation: please clarify new contributions of the paper carefully compared to the previous work by the same authors (DIALOGUE-2025 conference), revise the claimed results so that they do not repeat the prior work, provide more detailed steps to reproduce the results, revise the title to highlight the main result of the paper” — Thank you for your comments. We have revised the paper to clearly highlight the new contributions.

11.2 Review 2

1. “In the abstract, the authors claim that the paper presents "a novel series of Russian information retrieval datasets" and also that "we compare lexical retrieval models, such as BM25, with state-of-the-art neural architectures fine-tuned for the Russian language as well as multilingual models". It seems to me that the paper presents basically a comparison of models, the datasets are collateral and are not in the goals of the paper.” — The primary goal of our research was to create a novel series of datasets valuable for information retrieval tasks. As described in Section 4, these datasets can be applied in multiple scenarios. To demonstrate their usefulness, we conducted an extensive set of experiments showing how model performance varies under different conditions.
2. “The fact is that I did not find in the paper (it may be my mistake)where it is shown anything about "semantic nuances". It your claim is true (and I am sure it is) it is a

relevant result in your paper, but you need to indicate it in the results section and make it clear and explicit.” — The text has been revised to remove ambiguity. The phrase “semantic nuances” was intended to describe the ability of neural models to capture semantic relationships in text.

3. “Finally, could you comment on the fact that the specific choice of datasets to be analyzed may influence your results? Taking into account that you tailored your dataset, do you think that your results are agnostic to the dataset or, on the contrary, they may have biased the results?” — Although our datasets are derived from Wikipedia, their wide topical coverage reduces the risk of bias. Furthermore, as shown in Section 9, the same performance trends between lexical and neural models are observed in the full RusBEIR benchmark, which includes multiple datasets across diverse domains. This consistency supports the robustness of our findings beyond the Wikipedia domain.

11.3 Review 3

1. “Section 3. The human annotation process (wikifacts-sents) should be described in more detail. Namely, it should provide a number of human annotators per sentence and an agreement score.” — We added more details to the description of the annotation process. “In total, 55 university students were asked to identify relevant sentences. Each student was provided with a set of facts accompanied by referenced articles. The task was to read the articles associated with each fact and mark the sentences that appeared relevant, selecting from the predefined scoring options.”
2. “Section 6. All the considered models (except BGE M3) have a relatively small context window (256-512 tokens). On the one hand, many of them outperform BM25 at the sentence level and underperform BM25 at the passage or article level. On the other hand, Section 4 teaches that the score of an article is the maximum score of the article’s sentences. It would be helpful to have a more detailed description of applying the neural-network-based models with small contexts to large texts (articles and passages) to tackle that seeming inconsistency.” — We expanded the description of how neural models with small context windows are applied. “As shown in Table 3, different models support different maximum sequence lengths, which define their context windows. In our approach, we load both the model and tokenizer from HuggingFace, and set the `maxlen` parameter according to the model specification. Internally, the tokenizer produces at most `maxlen` tokens; if the input text is longer, it is truncated to this limit. As a result, the model processes only the first `maxlen` tokens, which are then used for subsequent computations.”