

Impact of Tokenization on LLaMa Russian Adaptation

1st Mikhail Tikhomirov

Lomonosov Moscow State University

Moscow, Russia

tikhomirov.mm@gmail.com

2nd Daniil Chernyshev

Lomonosov Moscow State University

Moscow, Russia

chdanorbis@yandex.ru

Abstract—Latest instruction-tuned large language models (LLM) show great results on various tasks, however, they often face performance degradation for non-English input. There is evidence that the reason lies in inefficient tokenization caused by low language representation in pre-training data which hinders the comprehension of non-English instructions, limiting the potential of target language instruction-tuning. In this work we investigate the possibility of addressing the issue with vocabulary substitution in the context of LLaMa Russian language adaptation. We explore three variants of vocabulary adaptation and test their performance on Saiga instruction-tuning and fine-tuning on Russian Super Glue benchmark. The results of automatic evaluation show that vocabulary substitution not only improves the model’s quality in Russian but also accelerates fine-tuning (35%) and inference (up to 60%) while reducing memory consumption. Additional human evaluation of the instruction-tuned models demonstrates that models with Russian-adapted vocabulary generate answers with higher user preference than the original Saiga-LLaMa model.

Index Terms—large language models, tokenization, language adaptation, llama

I. INTRODUCTION

With the introduction of powerful pre-trained language models such as BERT [9], neural models became the basis of Natural Language Processing (NLP) solutions. Thanks to accumulated global contextual knowledge the pre-trained language models exhibit higher task adaptability, requiring less training data and computational resources for fine-tuning while achieving better results. However such models are harder to replicate as the computational costs of the pre-training procedure are several times higher than training the model from scratch on downstream task [3], which prompted the researchers to pursuit the universal model with the best out-of-box performance.

The initial works focused on task-specific pre-training strategies [6]–[8], however, the exploration of knowledge accumulation boundaries led to the discovery of a more flexible alternative, in-context learning [3]. In contrast to conventional fine-tuning approaches, in-context learning adapts to the task by conditioning on task description (instruction) and a set of analogous examples that are provided within the user input which tunes the network attention mechanism to the task relevant knowledge. This capability was shown to appear only in larger language models (over 6B parameters) that have a sufficient network capacity to memorize the context

of various domains and tasks. Latest iteration of in-context learning leverages explicit instruction tuning [4] to eliminate the need for analogous examples reinforcing the role of Large Language Models (LLM) as universal NLP task solvers.

Current state-of-the-art instruction-tuned LLMs such as ChatGPT show outstanding zero-shot performance on various domains, yet their performance significantly degrades when applied to non-English domains. The main reason is the dominance of English data in available pre-training datasets which leads to English bias in both knowledge and morphology. Consequently, developing a language-specific counterpart is much more challenging due to lack of domain coverage which is required for achieving model universality. A more affordable alternative is fine-tuning multilingual LLMs on target-language instruction datasets [2]. Most implementations¹²³ use LLaMa model [5] as the backbone for instruction-based language adaptation, however, being an English-focused model it suffers from low tokenization efficiency for other languages. Beside performance overhead related to excessive number of tokens, inefficient tokenization also results in violation of morphological characteristics, hindering the language-adaptation process [14], [15].

We argue that tokenization optimization is the necessary step for efficient LLM language adaptation. In this work, we explore different embedding optimization variants in context of LLaMa model and compare their performance on fine-tuning and instruction-tuning settings. Following generalized vocabulary substitution method [12], [13] we considered three vocabulary options: Russian BPE, Russian Unigram and the original. Additionally, we analyzed two different algorithms for embedding and LM head initialization. The quality evaluation was done on Russian Super Glue benchmark and side-by-side tests. Our key contributions are:

- We show that Unigram tokenization has higher morphological accuracy than the BPE algorithms, utilized in state-of-the-art models;
- Our benchmark demonstrates that Russian large language models highly benefit from morphologically accurate

¹<https://github.com/avocardio/Zicklein>

²<https://github.com/22-hours/cabrita>

³<https://github.com/IlyaGusev/rulm>

word tokenization, achieving a considerable quality boost with Unigram vocabulary at all evaluation settings;

- Additional Human evaluation performed by 15 annotators reveals that vocabulary substitution significantly improves the instruction-tuning efficiency, with better relevance of generated answers;
- Performance tests indicate that vocabulary substitution substantially boosts the efficiency of resource utilization, accelerating model inference by up to 60% while also reducing the memory consumption.

II. RELATED WORK

Embedding adaptation has been popular technique since the introduction of pre-trained language models. Lakew [11] et al. studied the problem in the scope of cross-lingual neural machine translation transfer learning and showed that adapting input vocabulary to unseen languages could be more efficient than training a full model from scratch. Similarly, Kuratov et al. [12] showed that vocabulary substitution of multilingual BERT and pre-training the updated embeddings significantly improves the model performance on Russian language. A more detailed study of this language optimization approach was done by Rust et al [13], which concluded that the performance improvements generalize to all languages. Later Vries et al. [16] adopted the vocabulary substitution approach for GPT-2 cross-lingual adaptation and achieved comparable results to Italian-specialized counterpart. Zeng et al. [1] improved the vocabulary substitution approach for monolingual models with bilingual lexicon embedding initialization, which allowed to skip the embedding pre-training stage for encoder-only pre-trained models.

III. METHODOLOGY

Following previous works [12], [13] we use generalized vocabulary substitution training strategy to adapt LLaMa model to Russian language. The approach can be summarized in the following steps:

- 1) Build a new tokenization vocabulary on target language corpus,
- 2) Rebuild the embedding and LM head layers to account the new vocabulary size,
- 3) Initialize the updated layers using the old embedding weights in accordance to the overlap between old and new tokenizations,
- 4) Pre-train the updated layers on language modeling task using the target language corpus.

The implementation details of each step will be discussed further.

A. Step 1: Tokenization

We considered BPE and Unigram tokenization algorithms. Both algorithms use subword vocabularies (char n-grams), however, they differ in terms of vocabulary construction strategies. BPE uses bottom-up strategy, merging the subword tokens according to their corpus frequency. In contrast, Unigram is top-down approach that constructs the tokens by

pruning them to maximize the corpus likelihood. For that reason Unigram is more prone to retain elementary language semantic units like roots or stems [10]. We hypothesize that this trait might be beneficial for inflected languages such as Russian.

B. Step 2-3: Embedding initialization

To initialize replaced matrices, new tokens were processed based on the old tokenization, and then the corresponding rows of the matrix were initialized by averaging the vectors of resulting tokens in the original embedding matrix. In particular, in order to implement this procedure correctly, we need to carefully handle tokens without a leading space, since the tokenizer by default treats the first word as having a space before it.

$$\begin{aligned} v_{new}(t_i^n) &= \frac{1}{K} \sum_{j=1}^K v_{raw}(t_j^r); \\ tokenize_{raw}(t_i^n) &= [t_1^r, \dots, t_K^r], \end{aligned} \quad (1)$$

where $tokenize_{raw}$ is original tokenization function, t_i^n - token in new tokenization, t_1^r - token on original tokenization, $v_{raw}(\dots)$ - vector of token in original model.

To initialize LM head, we investigated 2 options: 1) initialization of LM head with a copy of the embedding layer, and 2) initialization by analogy with embeddings, but using the old LM head (we will call such a case with the suffix *hm*).

C. Step 4: Continued target-language pre-training

We assumed that for language adaptation of multilingual LLM's it is enough to train only the embedding layers, since the model had to develop universal language semantics during its pre-training. For this reason, we froze all layers except the embedding and the LM head and trained on a Russian-language text corpus on the causal language modeling task (the pre-training task of the original LLaMa model) with cross-entropy loss. This decision also makes training procedure more affordable, as the number of training parameters are no more than a few percent of the total model size, thus alleviating memory requirements.

IV. EXPERIMENTAL SETUP

A. Training data

For our adaptation experiments we collected documents from the following domains:

- Russian Wikipedia
- Pikabu
- News
- Habrahabr
- Stackoverflow
- Books

The documents were deduplicated using Locality Sensitive Hashing Minhash algorithm. To better model Russian language vocabulary distribution and to avoid grammatically incorrect examples the metadata, links and comment sections of web-pages were dropped. We used UTF-8 normalization to remove non-standard symbols like emoji or logograms (e.g. Chinese

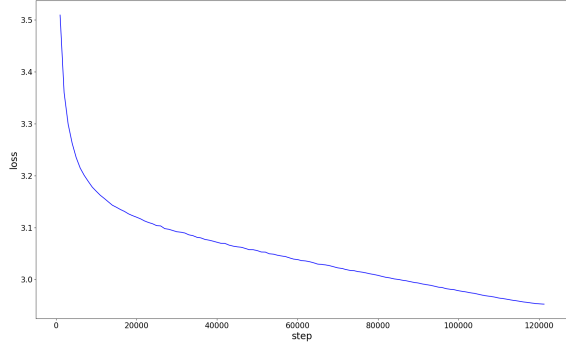


Fig. 1. Validation loss when adapting on Russian language

characters), thus restricting the vocabulary to Cyrillic and Latin languages. Text structure was standardized to singular space characters (i.e. no multi-spaces or multi-newlines). After all filtering procedures we obtained a dataset⁴ of 9 million documents or 3.5 billion words (~ 33 GB).

B. Tokenization

To train the tokenizers we used sentencepiece package⁵ on a subset of 5M documents from our language adaptation dataset. The tokenization characteristics were evaluated on RuMorphs-Words⁶ dataset, which contains morphological labels for more than 1.7M Russian words. To measure morphological accuracy we examined root integrity which is defined as the maximal length of longest common char subsequence among all word tokens normalized by root length.

C. Training settings

Training of new layers on the Russian language modeling dataset took place on 12 Nvidia V100. Since LLaMa-7b isn't natively supported by Volta architecture we had to use DeepSpeed Zero-2 to accelerate the training. The text data was chunked with Block size = 128 without word boundary alignment. We used Adam optimizer with $\text{lr} = 3\text{e-}04$, linear scheduler with warmup of 2000 steps and total batch size = 240. To prevent overfitting the models were trained with early stopping and evaluation on validation set every 0.1 epochs.

Our preliminary experiments indicated that 1 epoch was sufficient for achieving the original level of language coherence. Therefore, to save the costs and to ensure the purity of the experiment all variants of language adaptation were trained for that duration. However, it must be noted that 1 epoch is suboptimal (Figure 1) and the models could be improved by further pre-training.

D. Evaluation protocol

As the baseline we compared against the original LLaMa-7b model and its continued pre-training version (designated with

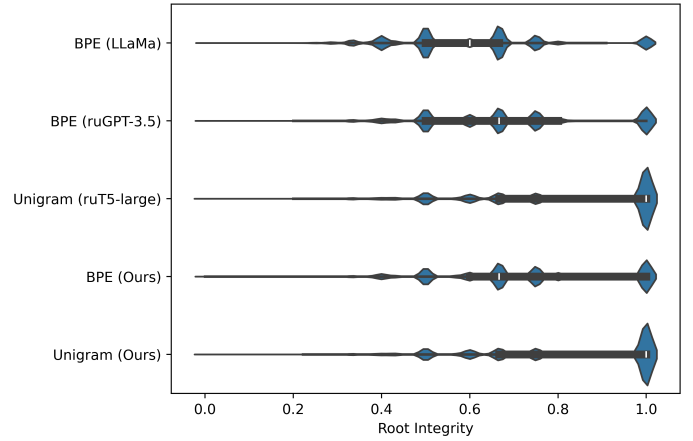


Fig. 2. Average root word preservation rate for different tokenizations

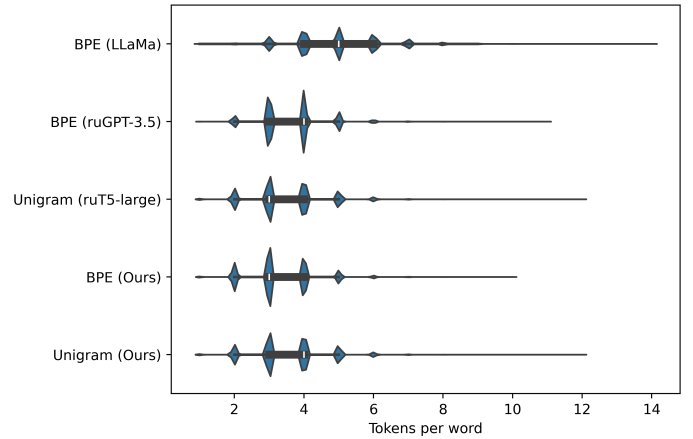


Fig. 3. Average word length in tokens for different tokenizations

the suffix *raw*) that was embedding-tuned without vocabulary substitution on our Russian pre-training corpus. The latter is necessary to measure the net quality increase from the improved tokenization and to disprove the hypothesis that the continued pre-training on its own could be optimal for language adaptation.

To evaluate language modeling performance we used Russian Super Glue (RSG) benchmark [17] which consists of 9 natural language understanding tasks. To maintain consistency with the existing LLaMa submissions we used evaluation protocol from Saiga model github repository⁷. Following the protocol all models were evaluated in two settings: downstream task fine-tuning and instruction-tuned zero-shot.

In the case of zero-shot, the model was first fine-tuned on the Saiga instruction datasets⁸ using the LoRA adapter-tuning approach [18], and then evaluated on RSG tasks through instruction-guided generation.

⁴Link to dataset is hidden for the purpose of blind review

⁵<https://github.com/google/sentencepiece>

⁶<https://cmc-msu-ai.github.io/NLPDatasets/>

⁷<https://github.com/IlyaGusev/rulm>

⁸<https://huggingface.co/collections/IlyaGusev/saiga-datasets-6505d5c30c87331947799d45>

V. RESULTS AND ANALYSIS

A. Tokenization quality

Beside the original LLaMa tokenizer we compared to tokenizers of state-of-the-art Russian Language models^{9,10}. According to Figure 2 conventional BPE algorithm have lower root preservation rate than Unigram regardless of selected vocabulary. However, the token-wise word length distribution (see Figure 3) is more biased to lower token count for BPE tokenizer variant, which suggests superior resource efficiency in terms of memory and processing time. At the same time the difference in median token count values between BPE and Unigram isn't enough to justify the significantly lower morphological qualities of the former. Therefore we proceeded with both tokenization variants in our language adaptation experiments.

B. Russian Super Glue

To evaluate the quality of the adapted models, we used the Russian Super Glue benchmark. The models were assessed in two settings: 1) zero-shot for instruction-tuned versions and 2) fine-tuning on training splits of Russian Super Glue datasets. All fine-tuning (and instruction-tuning) took place using the LoRA approach [18] on Q, V, K, O matrices of attention. In both cases, the answers from the model is obtained by generating text and then automatically searching for key phrases (such as Yes, No) or using the generated text as the entire answer, in the case of question-answering tasks.

In the case of fine-tuning on Russian Super Glue, all training sets of 9 datasets were combined into a single dataset in a dialogue format. This approach may not be optimal in terms of achieving the best quality on Russian Super Glue, but it has been used to replicate past results of LLaMa models. Additionally, it was more important for us not to achieve maximum quality, but to examine the difference when using different weights: original and adapted.

The results of both approaches are presented in Tables I and II. Models based on unigram tokenization perform best on average. BPE tokenization also shows an increase in quality in the zero shot setting, but is inferior to the unigram model. Calibrating the original model for Russian-language data (lama7b_rulm_raw and saiga7b_rulm_raw) also improves the quality of the model, but is inferior to options with a replacing the tokenization.

C. Human evaluation

Since both unigram models showed similar results in automatic evaluation, we performed additional human evaluation in side-by-side comparison setting for Saiga instruction-tuned versions. We prepared 78 questions for knowledge and logic probing and generated answers using the same sampling settings. To confirm the correlation between automatic metrics and real-world task performance we included the generation results of the original Saiga 7b in model comparison set. The

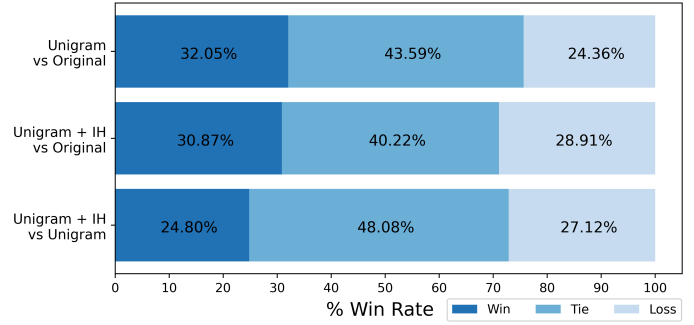


Fig. 4. Human evaluation results of Saiga models

models were compared pair-wise in independent sessions and each question was evaluated by 15 random annotators that were selected after preliminary testing. The annotators were instructed to select the answer variants with highest factual consistency (or logical accuracy) and personal relevance (i.e. the answer with most helpful information). For the cases of equal preference the annotators had "Tie" option.

The results of human evaluation are presented in Figure 4. As can be seen from the results, both unigram model variants outperform the original model by a significant margin. However, the annotators often found the difference in generation quality marginal, labeling the question pairs as "Tie" situation. Interestingly, exclusive LM head initialization (denoted as "+IH") has a negative effect on instruction-tuning efficiency and in side-by-side comparison this Unigram variant loses to the simpler embedding-tuning approach, which corroborates zero-shot Russian Super Glue results from section 5.B.

D. Performance

The main benefit of tokenization optimization is guaranteed speed up of model inference and training. We evaluated the Russian text generation resource efficiency for the original LLaMa 7b and our Unigram modification. To measure the pure sampling efficiency, the generation process was guided using ForceTokensLogitsProcessor which limited the token choice at each generation timestep to predefined token sequence. The sequence was broken down by sentence count boundaries at 1,5,10 and 15 sentences. The inference was done on RTX 4090 in max performance setting with native FP16 model weights and Flash attention 2 enabled.

Figure 5 shows the generation time and additional memory consumed for text length breakpoints. As can be seen the improvement in average token count substantially boosts the resource efficiency. The difference in time and memory consumption for short sequences (1 sentence) is marginal, however it becomes more evident with increasing text length. To produce a story of 15 sentences the generation process for our Unigram model would take approximately 17 seconds, while the original backbone would spend 27 seconds, thus achieving a ~60% speed up. Similarly, fine-tuning the model with the new tokenization on the instruct dataset is 35% faster,

⁹<https://huggingface.co/ai-forever/ruGPT-3.5-13B>

¹⁰<https://huggingface.co/ai-forever/ruT5-large>

TABLE I
LoRA FINE-TUNE RESULTS ON RUSSIAN SUPER GLUE

	LiDiRus	RCB	PARus	MuSeRC	TERRa	RUSSE	RWSD	DaNetQA	RuCoS	mean
llama7b	0,361	0,462	0,672	0,799	0,860	0,624	0,682	0,866	0,802	0,681
llama7b_rulm_raw	0,392	0,494	0,688	0,805	0,859	0,631	0,669	0,871	0,791	0,689
llama7b_rulm_bpe	0,365	0,509	0,684	0,782	0,844	0,626	0,747	0,824	0,737	0,680
llama7b_rulm_unigram	0,412	0,561	0,732	0,800	0,875	0,660	0,675	0,865	0,756	0,704
llama7b_rulm_unigram_hm	0,387	0,546	0,750	0,815	0,866	0,660	0,740	0,812	0,758	0,704

TABLE II
ZERO SHOT RESULTS OF SAIGA INSTRUCTION-TUNE ON RUSSIAN SUPER GLUE

	LiDiRus	RCB	PARus	MuSeRC	TERRa	RUSSE	RWSD	DaNetQA	RuCoS	mean
saiga7b	0,084	0,412	0,528	0,311	0,514	0,484	0,675	0,676	0,319	0,445
saiga7b_rulm_raw	0,025	0,373	0,610	0,310	0,523	0,587	0,584	0,783	0,474	0,474
saiga7b_rulm_bpe	0,149	0,429	0,596	0,344	0,647	0,478	0,636	0,757	0,397	0,493
saiga7b_rulm_unigram	0,194	0,432	0,568	0,313	0,591	0,587	0,630	0,789	0,477	0,509
saiga7b_rulm_unigram_hm	0,198	0,413	0,584	0,349	0,533	0,587	0,578	0,789	0,475	0,501

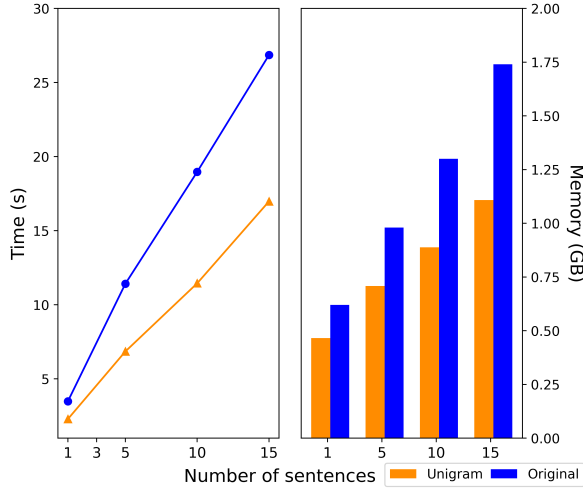


Fig. 5. Resource consumption for different text lengths

taking 20 hours instead of the original 27 hours at the same training configuration.

VI. CONCLUSIONS

We explored the possibility of adapting large language models (LLaMA) to Russian language with vocabulary substitution method. Our tokenization experiments with Unigram and BPE algorithms showed that Unigram algorithm is better suited for Russian language, having better morphological accuracy and comparable token count. This result was also confirmed on Russian Super Glue benchmark which demonstrated considerable quality improvements for all variants of vocabulary substitution with Unigram tokenizer bringing the highest improvements. Additional Human Evaluation of Saiga instruction-tuned models indicated better instruction utilization after Unigram vocabulary substitution which has higher user preference over the original Saiga7b model. We also showed that vocabulary substitution is an efficient method for resource optimization, achieving significant speed up both

for fine-tuning (35%) and inference (up to 60%) with memory consumption reduction.

ACKNOWLEDGEMENTS

The work of Mikhail Tikhomirov (general concept, method implementation, experiments) was supported by Non-commercial Foundation for Support of Science and Education "INTELLECT". The work of Daniil Chernyshev (dataset collection, result interpretation, survey) was supported by Non-commercial Foundation for Support of Science and Education "INTELLECT". The research is carried out using the equipment of the shared research facilities of HPC computing resources at Lomonosov Moscow State University.

REFERENCES

- [1] Q. Zeng, et al, "GreenPLM: Cross-Lingual Transfer of Monolingual Pre-Trained Language Models at Almost No Cost" //arXiv preprint arXiv:2211.06993. – 2022.
- [2] Y. Cui, et al, "Efficient and Effective Text Encoding for Chinese LLaMA and Alpaca" //arXiv preprint arXiv:2304.08177. – 2023.
- [3] T. Brown, et al, "Language Models are Few-Shot Learners" //arXiv preprint arXiv:2005.14165. – 2020.
- [4] L. Ouyang, et al, "Training language models to follow instructions with human feedback" //arXiv preprint arXiv:2203.02155. – 2022.
- [5] H. Touvron, et al, "LLaMA: Open and Efficient Foundation Language Models" //arXiv preprint arXiv:2302.13971. – 2023.
- [6] J. Zhang, et al, "PEGASUS: Pre-Training with Extracted Gap-Sentences for Abstractive Summarization," in Proceedings of the 37th International Conference on Machine Learning, 2020.
- [7] M. Tikhomirov, et al, "Using bert and augmentation in named entity recognition for cybersecurity domain," in Natural Language Processing and Information Systems: 25th International Conference on Applications of Natural Language to Information Systems, NLDB 2020, Saarbrücken, Germany, June 24–26, 2020, Proceedings 25, 2020, pp. 16–24.
- [8] W. Yoon, et al, "Pre-trained language model for biomedical question answering," in Joint European Conference on Machine Learning and Knowledge Discovery in Databases, 2019, pp. 727–740.
- [9] J. Devlin, et al, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers), 2019, pp. 4171–4186.
- [10] K. Bostrom, G. Durrett, "Byte Pair Encoding is Suboptimal for Language Model Pretraining," in Findings of the Association for Computational Linguistics: EMNLP 2020, 2020, pp. 4617–4624.

- [11] S. Lakew, et al, "Transfer Learning in Multilingual Neural Machine Translation with Dynamic Vocabulary," in Proceedings of the 15th International Conference on Spoken Language Translation, 2018, pp. 54–61.
- [12] Kuratov Y., Arkhipov M. Adaptation of deep bidirectional multilingual transformers for Russian language //arXiv preprint arXiv:1905.07213. – 2019.
- [13] P. Rust, et al, "How Good is Your Tokenizer? On the Monolingual Performance of Multilingual Language Models," in Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers), 2021, pp. 3118–3135.
- [14] A. Nayak, et al, "Domain adaptation challenges of BERT in tokenization and sub-word representations of Out-of-Vocabulary words," in Proceedings of the First Workshop on Insights from Negative Results in NLP, 2020, pp. 1–5.
- [15] V. Hofmann, J. Pierrehumbert, H. Schutze, "Superbizarre Is Not Superb: Derivational Morphology Improves BERT's Interpretation of Complex Words," in Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers), 2021, pp. 3594–3608.
- [16] W. Vries, M. Nissim, "As Good as New. How to Successfully Recycle English GPT-2 to Make Models for Other Languages," in Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021, 2021, pp. 836–846.
- [17] T. Shavrina, et al, "RussianSuperGLUE: A Russian language understanding evaluation benchmark" //arXiv preprint arXiv:2010.15925. – 2020.
- [18] Hu E. J. et al. Lora: Low-rank adaptation of large language models //arXiv preprint arXiv:2106.09685. – 2021.