

Large Language Models in Targeted Sentiment Analysis for Russian

N. Rusnachenko^{1*}, A. Golubev^{2**}, and N. Loukachevitch^{3***}

(Submitted by V. V. Voevodin)

¹Newcastle University, Newcastle Upon Tyne, England, United Kingdom

²Lomonosov Moscow State University, Moscow, 119234 Russia

³Research Computing Center, Lomonosov Moscow State University, Moscow, 119234 Russia

Received April 2, 2024; revised April 15, 2024; accepted May 5, 2024

Abstract—In this paper, we investigate the use of decoder-based generative transformers for extracting sentiment towards the named entities in Russian news articles. We study sentiment analysis capabilities of instruction-tuned large language models (LLMs). We consider the dataset of RuSentNE-2023 in our study. The first group of experiments was aimed at the evaluation of zero-shot capabilities of LLMs with closed and open transparencies. The second covers the fine-tuning of Flan-T5 using the “chain-of-thought” (CoT) three-hop reasoning framework (THoR). We found that the results of the zero-shot approaches are similar to the results achieved by baseline fine-tuned encoder-based transformers (BERT_{base}). Reasoning capabilities of the fine-tuned Flan-T5 models with THoR achieve at least 5% increment with the base-size model compared to the results of the zero-shot experiment. The best results of sentiment analysis on RuSentNE-2023 were achieved by fine-tuned Flan-T5_{xl}, which surpassed the results of previous state-of-the-art transformer-based classifiers. Our CoT application framework is publicly available: <https://github.com/nicolay-r/Reasoning-for-Sentiment-Analysis-Framework>

DOI: 10.1134/S1995080224603758

Keywords and phrases: *sentiment analysis, large language models*

1. INTRODUCTION

In recent years, large language models (LLMs) based on the Transformer architecture have significantly changed the landscape of natural language processing (NLP). Such models are pre-trained on large volumes of unlabeled texts. The so-called instruction-tuned language models are further trained on large sets of instructions (prompts and correct answers). This pre-training made it possible to solve tasks without training (fine-tuning) models on target datasets in the so-called zero-shot or few-shot formats [1, 2]. The zero-shot format is based solely on the formulation of a special prompt (question) for a model [1, 3]. The few-shot format comprises a prompt and several examples of correct answers [1]. In addition, there are special prompts that contain instructions for reasoning, referred to as Chain-of-Thought (CoT) prompts [4].

In sentiment analysis, the application of pre-trained language models has led to a significant improvement in the performance in various tasks. Sentiment analysis tasks can be divided into two main categories: general sentiment analysis [5], and targeted sentiment analysis (TSA) [6]. General sentiment analysis aims to determine the overall sentiment of a text, while targeted sentiment analysis focuses on identifying the sentiment towards a specific entity [7–9], its characteristics (aspects) [10], or controversial issues.

*E-mail: rusnicolay@gmail.com

**E-mail: antongolubev5@yandex.ru

***E-mail: louk_nat@mail.ru

2. RELATED WORK

Evaluating language models in sentiment analysis, the authors of [11] compare the performance of LLMs such as ChatGPT (GPT-3.5 and GPT-4), PaLM [12], Flan-UL2 [13], and LLaMA [14] across 13 sentiment analysis tasks on 26 datasets and compare the results against small language models (SLMs) trained on domain-specific datasets. The tasks under study include document- and sentence-level sentiment analysis, aspect-based sentiment analysis, and also such tasks as implicit sentiment, irony, hate speech detection, etc. With such a diverse range of tasks and models, the authors created a standardized template for prompts, containing the task name, definition, and desired output format. The authors conclude that the zero-shot application of LLMs is already effective for simpler sentiment classification tasks, such as binary and trinary classification. However, for tasks that require structured sentiment outputs, such as aspect-based analysis, the performance of LLMs lags behind that of small models trained on specific domains: LMs often have a lower accuracy than fine-tuned ones.

In [15], the authors evaluate large language models GPT-3.5, BLOOMZ, and XGLM in sentiment classification in 34 languages, including 6 high/medium-resource languages, 25 low-resource languages, and 3 code-switching datasets. When prompting LLMs, six variants of prompts are used, and the obtained results are averaged. For high-resource languages, GPT-3.5 achieves 77.5% by F-measure, for low-resource (African) languages shows 38.3% by F-measure in 3-way classification.

In [16], authors illustrate the application of the CoT concept to extract implicit sentiments from users' reviews [17, 18]. According to the provided paradigms, the proposed Three-Hop-Reasoning (THoR) system could be treated as an emerged paradigm: task-agnostic schema of reasoning steps, with one-by-one components referred to sentiment analysis. With these steps, authors aimed at extraction of "aspects", with further "opinion" as atomic components for devising the final answer [16]. The authors conclude that the application of the fine-tuning process for instruction-tuned Flan-T5 [19] results in models that surpass few-shot systems across publicly open systems [20] and significantly outperform encoder-based classifiers [21]. When it comes to limitations, authors conclude their beliefs on unleashing the full LLMs reasoning capabilities by applying THoR towards large enough models [16].

For Russian, there are several directions of related investigations: aspect-based (SentiRuEval) [22] and entity-oriented sentiment analysis (RuSentNE-2023) [9, 23]. The most recent advances in both tasks show that the highest results are mainly achieved by fine-tuned encoder-based classification language models [21, 23–25]. The best results on RuSentNE-2023 evaluation were obtained by ensembles of BERT-like encoder models [9].

Several recent work explores the application of generative models in sentiment analysis tasks. The authors of [26] study generative models of the GPT family in the Aspect-Based Triplet Extraction Task (ASTE). The authors compare the few-shot strategies for the GPT-3 and ChatGPT (GPT-3.5) models with the fine-tuned Russian ruGPT-3_{small} and ruGPT-3_{large} models, based on the GPT2 architecture [27]. They found that a few-shot approach on ruGPT-3 family models did not produce adequate results: fine-tuned ruGPT-3 models showed a significant improvement. The authors of [28] adopt the T5 model [20] and compare it with encoder-based approaches in the RuSentNE-2023 evaluation. Two variants of the Russian-adapted models ruT5_{base} and ruT5_{large} were used. However, the best results obtained by ruT5_{large} are 7 percentage points lower than the top RuSentNE-2023 submission [9].

3. RUSENTNE-2023 EVALUATION AND DATASET

The RuSentNE-2023 dataset¹⁾ is annotated with sentiment towards named entities in news texts. News texts pose challenges for targeted sentiment analysis [9]. At first, such texts may contain opinions conveyed by different subjects, including the author(s)' attitudes, positions of cited sources, and relations of mentioned entities to each other. Secondly, some sentences contain several named entities with different sentiments, complicating the determination of sentiment towards each individual named entity. Thirdly, the majority of named entities in news texts are mentioned in neutral context, indicating a significant prevalence of the neutral class. Last but not least, the significant amount of sentiment in news texts is implicit in nature, primarily conveyed through entity actions.

¹⁾Resources are publicly available: <https://github.com/dialogue-evaluation/RuSentNE-evaluation>

Table 1. Distribution of entity types in training (**train**), development (**dev**) and test (**test**) sets of the RuSentNE-2023 dataset

	train		dev		test	
Entity type	#	%	#	%	#	%
Person	1934	29	857	30	480	25
Profession	1666	25	533	24	510	26
Organization	1487	23	653	23	484	25
Country	1274	19	686	19	363	19
Nationality	276	4	116	4	110	5
Total	6637	100	2845	100	1947	100

The source of the annotated sentiment in RuSentNE-2023 could be (i) an author, (ii) another cited source, and (iii) another entity mentioned in the text. Sentiment annotation is a three-scale and has the following labels: positive, negative, and neutral.

For example, in the following sentence there is a negative sentiment to *Hungary* and neutral sentiment to *German Economy Minister*. The source of the negative sentiment to *Hungary* is the *German Economy Minister*:

Министр экономики Германии критикует Венгрию за налог на иностранных инвесторов.
(German Economy Minister criticizes Hungary for tax on foreign investors.)

Named entities of the following types represent objects of sentiment [29]:

- Person—physical person regarded as an individual,
- Organization—an organized group of people or company,
- Country—a nation or a body of land with one government,
- Profession—jobs, positions in various organizations, and professional titles,
- Nationality—nouns denoting country citizens and adjectives corresponding to nations in contexts different from authority-related.

The distribution of entity types in the training (**train**), development (**dev**), and test (**test**) sets of the RuSentNE-2023 dataset is presented in Table 1.

4. EXPERIMENTAL SETUP

We experiment with the initial dataset RuSentNE-2023 and its English translation (RuSentNE-2023^{en}). Since most LLMs are trained on English data, we adopt GoogleTranslate²⁾ to automatically translate and compose RuSentNE-2023^{en}.

To evaluate the predicted results, we follow the competition rules [9] and adopt macro F1-measure ranged [0, 100] over: (1) positive and negative classes $F_1(PN)$, (2) all task classes $F_1(PN0)$.

We investigate LLMs reasoning capabilities with the following modes: (i) zero-shot and (ii) fine-tuning. To conduct the experiment, our computational resources were limited by access to a single NVIDIA A100 GPU (40GB). The model training and inference code are implemented in Python-3.8.

²⁾<https://pypi.org/project/googletrans/>

Table 2. List of LLMs utilized in Zero-shot experiments, separated on “closed models” with access via OpenAI API (GPT) and “open models” [19, 30–33] which size does not exceed 7 billion parameters

Models	Versions	Params	Reference
GPT-4	1106-preview	1.76T ^a	<code>gpt-4-1106-preview</code> (OpenAI API)
GPT-3.5	1106	175B ^b	<code>gpt-3.5-turbo-1106</code> (OpenAI API)
	0613	175B ^b	<code>gpt-3.5-turbo-0613</code> (OpenAI API)
Mistral [30]	v0.1	7B	https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.1
	v0.2	7B	https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2
DeciLM: [31]		7B	https://huggingface.co/Deci/DeciLM-7B-instruct
Microsoft-Phi-2 [33]		2.7B	https://huggingface.co/microsoft/phi-2
Gemma: [32]	Instructive	7B ^c	https://huggingface.co/google/gemma-7b-it
	Instructive	2B	https://huggingface.co/google/gemma-2b-it
Flan-T5: [19]	XL	3B	https://huggingface.co/google/flan-t5-xl
	Large	750M	https://huggingface.co/google/flan-t5-large
	Base	250M	https://huggingface.co/google/flan-t5-base

^a Non-disclosed by OpenAI; according to the other sources, GPT-4 yields of eight models 220B-sized parameters each connected by a Mixture of Experts (MoE) [48] ($\approx 1.76T$ params)

^b Non-disclosed by OpenAI, our assumption that the architecture is referred to GPT-3, size of 175B params [1]

^c The actual and non-official amount of hidden parameters for Gemma-7B-IT model is 9B

4.1. LLMs Zero-shot Experiments Setup

Our setup covers a range of recently popular models that fall into two main categories: (i) “closed models” and (ii) “open models” [19, 30–33].

In the case of closed models, the following chat-based dialogue assistants were used:

- GPT-4, version 1106 preview;
- GPT-3.5, versions 1106, 0613.

For the GPT models, we utilized a so-called system prompt, which is a special message used to assign a role to the assistant. System prompts prescribe the style and task for the chat-bot communication. We used the following system prompt: **System Message.** You are an AI assistant skilled in natural language processing and sentiment analysis. Your task is to analyze text inputs to determine the underlying sentiment, whether it’s positive, negative, or neutral. You should consider the nuances of language, including sarcasm, irony, and context. Your responses should include not only the sentiment classification but also a brief explanation of why a particular sentiment was assigned, highlighting key words or phrases that influenced the decision. In the case of open models, we experiment with instruction-tuned versions. The complete list of the assessed models is presented in Table 2 and includes:

- Mistral [30]—grouped-query attention (GQA) [34] for faster inference, coupled with sliding window attention [35] to effectively handle sequences of arbitrary length with a reduced inference cost. The authors adopt byte-fallback BPE tokenization [36]. There are several versions of publicly available instruction-tuned models: v0.1 and v0.2. The v0.1 has the input context window size of 8K tokens.³⁾ In v0.2, the related size has been increased up to 32K context.⁴⁾ For both versions, information on training data is not publicly available.

³⁾Technically is unlimited, with threshold originated by the 4K size of sliding window [35].

⁴⁾Our assumption that model is no longer adopts the sliding window attention [35].

- DeciLM [31]—is an auto-regressive language model using an optimized transformer decoder architecture that includes variable GQA [34]. Information on fine-tuning and utilized datasets is publicly available. Model has been fine-tuned for instruction with LoRA [37] on the SlimOrca [38] dataset. The context window size is 8K tokens.
- Microsoft-Phi-2 [33]—trained using the resource adopted in Phi-1.5 [33]; augmented with a new data source that consists of various NLP synthetic texts and filtered websites (for safety and educational value). Dataset represents combination of NLP synthetic data created by AOAI⁵⁾ GPT-3.5 and filtered web data from Falcon RefinedWeb [39] and SlimPajama [40], assessed by AOAI GPT-4. The context window size of the model is 2K tokens. For model training, authors utilize 96xA100-80G for 14 days.
- Gemma [32]— Built from the same research and technology used to create the Gemini models [41]. Represent a text-to-text, decoder-only large language models, with architecture similar to LLaMA [14]. Available in English, with open access to instruction-tuned variants. Represent well-suited for a variety of text generation tasks, including question answering, summarization, and reasoning. The context window size is 8K tokens.
- Flan-T5 [19]— Represent a specialized variant of the T5 [42], initially proposed as the unified text-to-text transformer. The concept of fine-tuning language models based on instructions (Flan) [19] highlights the significant improvement across series of tasks, including: MMLU [43], BBH [44], TyDiQA [45], MGSM [46], open-ended generation. Flan-T5 trained with the encoder context length of 1K tokens.

Given the sentence (X) with the target entity mentioned in it (t), we utilize prompting techniques of two types: original version (v1) [23] and precisely adapted the revised version, dubbed as (v2): **v1_{ru}**: Какая оценка тональности в предложении X , по отношению к t ? Выбери из трех вариантов: позитивная, негативная, нейтральная. **v1_{en}**: What's the attitude of the sentence X to the target t ? Select one from: positive, negative, neutral. **v2_{ru}**: Каково отношение автора или другого субъекта в предложении X к t ? Выбери из трех вариантов: позитивная, негативная, нейтральная. **v2_{en}**: What is the attitude of the author or another subject in the sentence X to the target t ? Choose from: positive, negative, neutral. In the case of proprietary models, we adopt OpenAI API service to access the OpenAI models. To reduce the charging cost for the tokens, we limit the result output for GPT-3.5 and GPT-4 models by using the response threshold of 75 tokens. We augment the initial prompt with the suffix that requires short answers as follows: “Create a very short summary that uses 50 `completion_tokens` or less”. For LLMs that can be downloaded for local use, in this paper, we experiment with models whose size does not exceed 7B parameters. We utilized `transformers` API [47] to conduct experiments on rented servers. In all models for zero-shot experiments, the value of the temperature parameter was chosen as 0.1.

To assess the inferred textual responses, in this study we adopt a universal annotation answer mapping strategy that involves the application of class-specific textual lower-cased templates for classes (positive, negative, and neutral), separately declared for each language. We use these templates during a sequential occurrence check in lower-cased output in the following order: (1) neutral, (2) positive, and (3) negative. In the case of absence of any class, the “UNK” placeholder (counted as “neutral”) was considered. We use “positive”, “negative”, and “neutral” templates for RuSentNE-2023_{en}, and “позитив” (positive), “негатив” (negative), and “нейтрал” (neutral) templates for experiments on RuSentNE-2023 dataset. We perform final checks involving regular expressions to guarantee that the final answers are not prefixed with the initial prompt.

4.2. LLMs Fine-tuning Setup

To experiment with the fine-tuning, we adopt encoder-decoder style instruction-tuned Flan-T5⁶⁾ as our backbone large language model for the proposed methodology for texts written in English. In this paper, we experiment with the following fine-tuning techniques:

⁵⁾<https://github.com/microsoft/sample-app-aoai-chatGPT>

⁶⁾https://huggingface.co/docs/transformers/en/model_doc/flan-t5

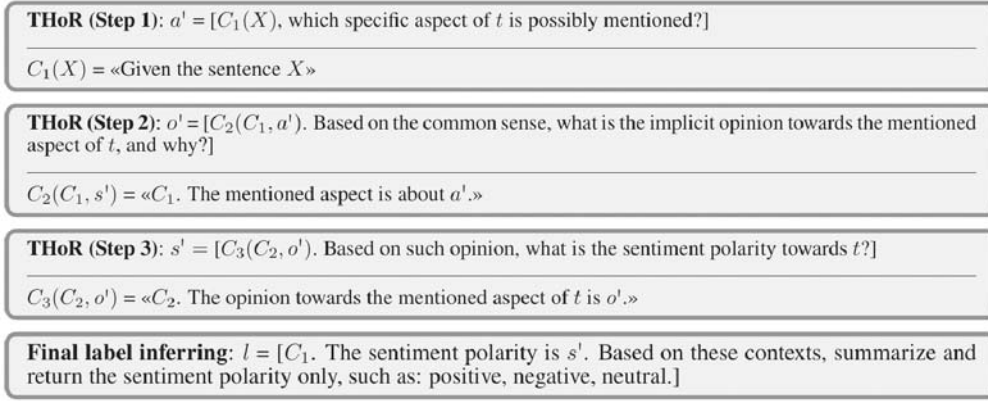


Fig. 1. Inferring sentiment s' using CoT three-hop reasoning framework (THoR), including “*final label inferring*” to answer one of the task classes [16].

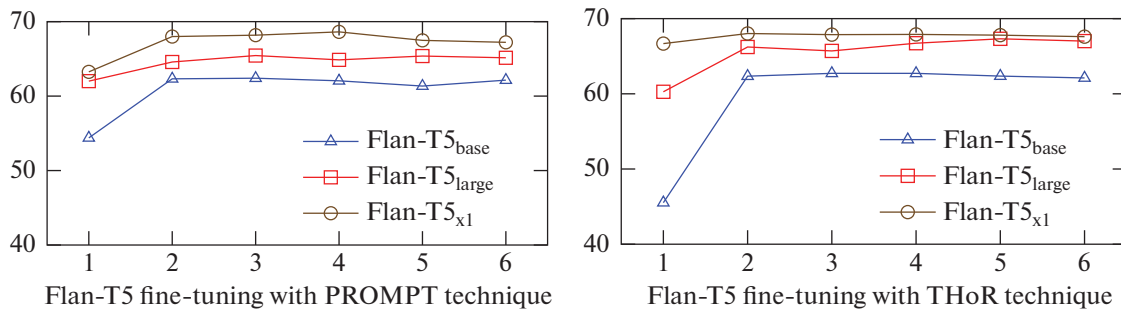


Fig. 2. Analysis of the Flan-T5 models results on RuSentNE-2023^{en} dev per each epoch (horizontal axis) by $F_1(PN)$ (vertical axis) during fine-tuning with PROMPT (left) and THoR technique (right) per different sizes.

- PROMPT—prompt-tuning with the $v1_{en}$ version of the prompt;
- THoR—Three-Hop-Reasoning [16] technique, which can be considered as an emerged paradigm: task-agnostic schema of reasoning steps, with one-by-one components referred to sentiment analysis. The experiments with three-hop reasoning are resource and time-intensive, so we currently study this technique only for one LLMs family.

Next, we cover THoR in greater details. With $C_i, i \in \overline{1..3}$ we denote the prompts that wrap the content in the input context. Given the sentence (X) with target entity mentioned in it (t), Fig. 1 illustrates the initial application of the three-step approach [16] for inferring the sentiment s' . According to Table 1, the result a' could be interpret as $a' = \operatorname{argmax} p(a|X, t)$, opinion o' as $o' = \operatorname{argmax} p(o|X, t, a'_1)$, and the final answer s' noted as: $s' = \operatorname{argmax} p(e|X, t, s', o')$. To guarantee the correctness of the final answer (s'), we use the following prompt message to inferring the sentiment label l (Fig. 1, bottom).

We experiment with: 250M (base), 750M (large), and 3B (XL) versions. We conduct experiments only for English-translated RuSentNE-2023^{en} dataset due to the pre-training specifics of Flan-T5 model. For mapping the Flan-T5 output towards task classes, we seek for the exact string from the set of textual task labels: “**positive**”, “**negative**”, and “**neutral**”. When it comes to fine-tuning parameters setup, we follow the settings proposed by authors of THoR framework [16]. In particular, we use AdamW [49] optimizer with learning rate $2 \cdot 10^{-4}$ and BATCH-SIZE of 32. To infer the answers, the maximal temperature of 1.0 was considered.

5. RESULTS AND DISCUSSION

Table 3 shows the zero-shot prompting results on the **test** set of RuSentNE-2023 across various LLMs, separately for the original and translated texts (RuSentNE-2023^{en}). The Table 3 contains two

Table 3. Results of the LLMs application in zero-shot mode for the **test** part of the RuSentNE-2023 dataset, separately for original texts and automatically translated in English (RuSentNE-2023^{en}); for “N/A%” column values (no-answer rate), “.” denotes cases when the amount of unknown answers does not exceed 1%; top results per each version of the dataset are bolded

Model	$F_1(PN)$	$F_1(PN0)$	N/A%	$F_1(PN)$	$F_1(PN0)$	N/A%
RuSentNE-2023 ^{en} [<i>Translated Texts into English</i>]						
Prompt Type	v1 _{en}			v2 _{en}		
GPT-4 _{1106-preview}	54.43	63.44	.	54.59	64.32	.
GPT-3.5 _{turbo-0613}	49.22	59.51	.	51.79	61.38	.
GPT-3.5 _{turbo-1106}	47.87	54.62	.	47.04	53.19	.
Mistral-Instruct-7B _{v0.1}	49.56	58.86	.	49.46	58.51	.
Mistral-Instruct-7B _{v0.2}	45.69	57.16	.	44.82	56.04	.
DeciLM	42.73	49.88	.	43.85	53.65	1.44
Microsoft-Phi-2	37.26	31.91	5.55	40.95	42.77	3.13
Gemma-7B-IT	40.58	45.94	.	40.96	44.63	.
Gemma-2B-IT	18.70	39.51	.	31.75	45.96	2.62
Flan-T5 _{xl}	35.35	31.51	.	48.14	57.33	.
Flan-T5 _{large}	34.86	23.34	.	36.05	24.27	.
Flan-T5 _{base}	32.64	21.81	.	31.05	20.84	.
RuSentNE-2023 [<i>Original texts written in Russian</i>]						
Prompt Type	v1 _{ru}			v2 _{ru}		
GPT-4 _{1106-preview}	49.44	58.74	.	48.04	60.55	.
GPT-3.5 _{turbo-0613}	45.97	56.10	.	45.85	57.36	.
GPT-3.5 _{turbo-1106}	38.95	45.93	.	35.07	48.53	.
Mistral-Instruct-7B _{v0.2}	48.71	57.10	.	42.60	48.05	.

main measures for evaluation ($F_1(PN)$, $F_1(PN0)$), accompanied by *no-answer rate* (N/A%). In general, all models were capable of handling instructions generated using v1 and v2 prompts for entity-oriented sentiment analysis of initial and translated versions of the RuSentNE-2023 dataset.

It can be seen from Table 3 that such models that support Russian language are tend to perform worse with texts from RuSentNE-2023 than for its translated variant RuSentNE-2023^{en}. There are a few models for which the proportion of N/A% answers (Table 3) is significantly higher than for other models. The Microsoft-Phi-2 model with both versions of prompts returns the prompt along with its response. Since the full response is limited by the **max_token** value, it is truncated, and we do not receive information about sentiment classification from the model. The average length of a response for examples for which the model did not provide an answer is 465 characters, for the rest it equals 288. DeciLM and Gemma-2B models with v2 prompt configuration generate garbage answers in the form of data scraps from the prompt or clearly indicate that they cannot determine the sentiment polarity (“...so i cannot answer this question from the provided context”).

According to $F_1(PN)$, in the case of proprietary OpenAI models we see the similar performance order of the models’ results irrespective of the prompts language source. In particular, the highest results on RuSentNE-2023^{en} were achieved by GPT-4 ($F_1(PN)$ = 54.53), followed by GPT-3.5_{turbo-0613} ($F_1(PN)$ = 49.22) and GPT-3.5_{turbo-1106} ($F_1(PN)$ = 47.87).⁷⁾ Switching to the original RuSentNE-2023 results in 10% and 7% decrease in result by $F_1(PN)$ for GPT-4 and GPT-3.5_{turbo-0613}, respec-

⁷⁾We believe that the reason of the worse behavior of the GPT-3.5_{turbo-1106} against the previous GPT-3.5_{turbo-0613} is caused by a higher tolerance level of the results responses.

tively. In turn, the only multilingual Mistral-Instruct-7B_{v0.2} illustrates a performance comparable to GPT-3.5_{turbo-0613} and outperforms GPT-3.5_{turbo-1106} by $\approx 19\text{--}25\%$. Also, we found Mistral-Instruct-7B_{v0.2} as more sensitive to prompts on original RuSentNE-2023 texts, rather OpenAI models (see last row, Table 3).

Most open LLMs are able to follow English instructions. For RuSentNE-2023^{en}, we see that the best performing models are Mistral [30] models ($45.69 \leq F_1(PN) \leq 49.56$), and DeciLM [31] as the closest competitor alternative ($F_1(PN) = 42.73$). The remaining families of Microsoft-Phi-2 and Flan-T5 series demonstrate the gap in results ($32.64 \leq F_1(PN) \leq 37.26$). Through the entire range of open models, Mistral [30] is the only one capable of handling RuSentNE-2023 in Russian. The official release Mistral-Instruct-7B_{v0.2} was better suited for non-English languages.

6. ERROR ANALYSIS

For the error analysis, we selected examples from the RuSentNE-2023 test set where most models' predictions did not align with human annotations. The following main types of discrepancies between models' predictions and manual annotations are found (denoted as “E1”, “E2”, and “E3”):

E1. A sentence mentions the positive author's attitude to the target person and some negative event with this person (trauma, death). The annotators treat such examples as positive for the person because in most cases traumas or death do no influence on existing positive attitude. Models in zero-shot mode answer on such examples by choosing the negative sentiment to the target person due to “negative effect” context. In turn, fine-tuned models annotate most of such examples correctly. For instance, in the following example we can see that the author has a positive attitude towards Chuck Berry despite the incident described.

Легендарный Чак Берри потерял сознание на концерте.
(Legendary musician Chuck Berry fainted during a concert in Chicago.)

E2. A sentence mentions several entities, while the negative sentiment is directed only at one of these entities. Models cannot distinguish the correct object of this negative attitude. For example,

Юлия же в свою очередь обвиняет бывшего супруга в том, что он не выполняет решение суда, и
уже объявлен в федеральный розыск.
(Yulia, in turn, accuses her ex-husband of not complying with the court decision and has already been
put on the federal wanted list.)

In this case, the models predict a negative attitude towards Yulia, but in fact Yulia has a negative opinion of her husband.

E3. Similar to the E2. A sentence with evident negative sentiment mentions a single entity, but sentiment is directed to some out-of-sentence entity. Almost all models in zero-shot mode predict a negative sentiment to the mentioned entity. In the following example, the models in zero-shot mode answer by choosing negative sentiment to America, but in fact the negative opinion is towards the “situation”.

Ситуация, однако, не может нас не беспокоить—объем экспорта в Америку упал на 8%.
(The situation, however, cannot but worry us—the volume of exports to America fell by 8%.)

7. CONCLUSIONS

In this paper, we investigated the application of large language models (LLM) in targeted sentiment analysis within news texts. We follow the RuSentNE-2023 competition both for experiments with LLM and evaluation. The RuSentNE-2023 evaluation was aimed at extracting sentiment towards the objects annotated in Russian-language sentences. In experiments, we used the RuSentNE-2023 dataset as well as its English automatically translated version (RuSentNE-2023^{en}).

LLM-based sentiment analysis was investigated in several directions, which cover the influence of aspects such as: (i) the language of the source texts (Russian or English), (ii) variants of prompts, and (iii) the effect of fine-tuning, including the application of the Chain-of-Thought technique. We have discovered the significance of the content translation, since most LLM models demonstrate reasonable comprehension capabilities primarily in English. Zero-shot approaches achieve results comparable to fine-tuned BERT-based RuSentNE-2023 competition baselines. Experiments with LLM fine-tuning (Flan-T5) have shown the improvement in $\approx 10\%$ performance by $F_1(PN)$ for all considered Flan-T5 models. The highest results were obtained by fine-tuned xl-sized Flan-T5_{xl} model (3B+ parameters), which are currently the best results achieved on the RuSentNE-2023 test set.

In further, we aim to continue experimenting with the Chain-of-Thought in the following directions: (i) reasoning revision techniques using extensive resources of auxiliary information, (ii) parameter-efficient tuning for larger models.

FUNDING

The work is supported by the Russian Science Foundation (grant no. 21-71-30003).

CONFLICT OF INTEREST

The authors of this work declare that they have no conflicts of interest.

REFERENCES

1. T. Brown et al., “Language models are few-shot learners,” *Adv. Neural Inform. Proces. Syst.* **33**, 1877–1901 (2020).
2. J. Wei et al., “Finetuned language models are zero-shot learners,” in *Proceedings of the International Conference on Learning Representations* (2021). <https://doi.org/10.48550/arXiv.2109.01652>
3. Z. Bowen et al., “How would stance detection techniques evolve after the launch of ChatGPT?,” *arXiv: 2212.14548* (2024). <https://doi.org/10.48550/arXiv.2212.14548>
4. J. Wei et al., “Chain-of-thought prompting elicits reasoning in large language models,” *Adv. Neural Inform. Proces. Syst.* **35**, 24824–24837 (2022).
5. P. Turney, “Thumbs up or thumbs down? Semantic orientation applied to unsupervised classification of reviews,” in *Proceedings of the 40th Annual Meeting on Association for Computational Linguistics* (2002), pp. 417–424. <https://dl.acm.org/doi/10.3115/1073083.1073153>
6. O. Toledo-Ronen et al., “Multi-domain targeted sentiment analysis,” in *Proceedings of Annual Conference of the North American Chapter of the Association for Computational Linguistics* (2022), pp. 2751–2762. <https://doi.org/10.18653/v1/2022.naacl-main.198>
7. E. Amigó et al., “Overview of replab 2013: Evaluating online reputation monitoring systems,” in *Proceedings of International Conference of the Cross-language Evaluation Forum for European Languages* (2013), pp. 333–352. https://doi.org/10.1007/978-3-642-40802-1_31
8. N. Loukachevitch and Y. Rubtsova, “Entity-oriented sentiment analysis of tweets: results and problems,” in *Proceedings of Text, Speech, and Dialogue: 18th International Conference* (2015), pp. 551–559. https://doi.org/10.1007/978-3-319-24033-6_62
9. A. Golubev et al., “RuSentNE-2023: Evaluating entity-oriented sentiment analysis on russian news texts,” in *Proceedings of the International Conference Dialogue* (2023), pp. 130–141. <https://doi.org/10.28995/2075-7182-2023-22-130-141>
10. M. Pontiki et al., “Semeval-2016 task 5: Aspect based sentiment analysis,” in *Proceedings of Workshop on Semantic Evaluation SemEval-2016* (2016), pp. 19–30. <https://doi.org/10.18653/v1/S16-1002>
11. Z. Wenxuan et al., “Sentiment analysis in the era of large language models: A reality check,” *arXiv: 2305.15005* (2023). <https://doi.org/10.48550/arXiv.2305.15005>

12. A. Chowdhery et al., “PaLM: Scaling language modeling with pathways,” *J. Machine Learn. Res.* **24** (240), 1–113 (2023).
13. T. Yi et al., “UL2: Unifying language learning paradigms,” in *Proceedings of the 11th International Conference on Learning Representations* (2022). <https://doi.org/10.48550/arXiv.2205.05131>
14. T. Hugo et al., “LLaMA: Open and efficient foundation language models,” arXiv: 2302.13971 (2023). <https://doi.org/10.48550/arXiv.2302.13971>
15. K. Fajri et al., “Zero-shot sentiment analysis in low-resource languages using a multilingual sentiment lexicon,” in *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics* (2024), Vol. 1, pp. 298–320.
16. F. Hao et al., “Reasoning implicit sentiment with chain-of-thought prompting,” in *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics* (2023), Vol. 2, pp. 1171–1182. <https://doi.org/10.18653/v1/2023.acl-short.101>
17. M. Pontiki et al., “SemEval-2014 task 4: Aspect based sentiment analysis,” in *Proceedings of the 8th International Workshop on Semantic Evaluation SemEval 2014* (2014), pp. 27–35. <https://doi.org/10.3115/v1/S14-2004>
18. Z. Li et al., “Learning implicit sentiment in aspect-based sentiment analysis with supervised contrastive pre-training,” in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing* (2021), pp. 246–256. <https://doi.org/10.18653/v1/2021.emnlp-main.22>
19. H. Chung et al., “Scaling instruction-finetuned language models,” *J. Machine Learn. Res.* **25** (70), 1–53 (2024).
20. C. Raffel et al., “Exploring the limits of transfer learning with a unified text-to-text transformer,” *J. Machine Learn. Res.* **21** (140), 1–6 (2020).
21. 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* (2019), Vol. 1, pp. 4171–4186. <https://doi.org/10.18653/v1/N19-1423>
22. N. Loukachevitch et al., “SentiRuEval: Testing object-oriented sentiment analysis systems in Russian,” in *Proceedings of International Conference Dialogue* (2015), Vol. 2, pp. 3–13.
23. A. Golubev and N. Loukachevitch, “Improving results on Russian sentiment datasets,” in *Proceedings of Conference on Artificial Intelligence and Natural Language* (2020), pp. 109–121. https://doi.org/10.1007/978-3-030-59082-6_8
24. L. Yinhan et al., “RoBERTa: A robustly optimized BERT pretraining approach,” arXiv: 1907.11692 (2021). <https://doi.org/10.48550/arXiv.1907.11692>
25. S. Smetanin and M. Komarov, “Deep transfer learning baselines for sentiment analysis in Russian,” *Inform. Proces. Manage.* **58** (3) (2021). <https://doi.org/10.1016/j.ipm.2020.102484>
26. S. Chumakov et al., “Generative approach to aspect based sentiment analysis with GPT language models,” *Proc. Comput. Sci.* **229**, 284–293 (2023). <https://doi.org/10.1016/j.procs.2023.12.030>
27. A. Radford et al., “Language models are unsupervised multitask learners,” OpenAI blog (2019). Accessed 2024.
28. I. Moloshnikov et al., “Named entity-oriented sentiment analysis with text2text generation approach,” in *Proceedings of the International Conference Dialogue* (2023), pp. 362–370. <https://doi.org/10.28995/2075-7182-2023-22-361-370>
29. N. Loukachevitch et al., “NEREL: A Russian information extraction dataset with rich annotation for nested entities, relations, and wikidata entity links,” *Language Resour. Eval.* **58** (2), 1–37 (2023). <https://aclanthology.org/2021.ranlp-1.100>
30. J. Albert, et al., “Mistral 7B,” arXiv: 2310.06825. <https://doi.org/10.48550/arXiv.2310.06825>
31. DeciAI Research Team, DeciLM-7B. <https://huggingface.co/Deci/DeciLM-7B>. Accessed 2024.
32. J. Banks and T. Warkentin, Gemma: Introducing new state-of-the-art open models. <https://blog.google/technology/developers/gemma-open-models/>. Accessed 2024.
33. L. Yuanzhi et al., “Textbooks are all you need II: Phi-1.5 technical report,” arXiv: 2309.05463 (2023). <https://doi.org/10.48550/arXiv.2309.05463>
34. J. Ainslie et al., “GQA: Training generalized multi-query transformer models from multi-head checkpoints,” in *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing* (2023), pp. 4895–4901. <https://doi.org/10.18653/v1/2023.emnlp-main.298>
35. Iz Beltagy et al., “Longformer: The long-document transformer,” arXiv: 2004.05150 (2020). <https://doi.org/10.48550/arXiv.2004.05150>

36. S. Rico et al., “Neural machine translation of rare words with subword units,” in *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics* (2016), Vol. 1, pp. 1715–1725. <https://doi.org/10.18653/v1/P16-1162>
37. E. Hu et al., “LoRA: Low-rank adaptation of large language models,” arXiv: 2106.09685 (2021). <https://doi.org/10.48550/arXiv.2106.09685>
38. L. Wing et al., “SlimOrca: An open dataset of GPT-4 augmented FLAN reasoning traces, with verification”. <https://huggingface.co/Open-Orca/SlimOrca>. Accessed 2024.
39. P. Guilherme et al., “The RefinedWeb dataset for Falcon LLM: outperforming curated corpora with web data, and web data only,” arXiv: 2306.01116 (2023). <https://doi.org/10.48550/arXiv.2306.01116>
40. D. Soboleva et al., “SlimPajama: A 627B token cleaned and deduplicated version of RedPajama”. <https://huggingface.co/datasets/cerebras/SlimPajama-627B>. Accessed 2024.
41. A. Rohan et al., “Gemini: A family of highly capable multimodal models,” arXiv: 2312.11805 (2023). <https://doi.org/10.48550/arXiv.2312.11805>
42. C. Raffel et al., “Exploring the limits of transfer learning with a unified text-to-text transformer,” *J. Mach. Learn. Res.* **21**, 5485–5551 (2020).
43. H. Dan et al., “Measuring massive multitask language understanding,” arXiv: 2009.03300 (2020). <https://doi.org/10.48550/arXiv.2009.03300>
44. M. Suzgun et al., “Challenging BIG-bench tasks and whether chain-of-thought can solve them,” in *Findings of the Association for Computational Linguistics: ACL 2023* (2023), pp. 13003–13051. <https://doi.org/10.18653/v1/2023.findings-acl.824>
45. J. Clark et al., “TyDi QA: A benchmark for information-seeking question answering in typologically diverse languages,” *Trans. Assoc. Comput. Linguist.* **8**, 454–470 (2020). <https://aclanthology.org/2020.tacl-1.30>
46. S. Freda et al., “Language models are multilingual chain-of-thought reasoners,” arXiv: 2210.03057 (2022). <https://doi.org/10.48550/arXiv.2210.03057>
47. W. Thomas et al., “Transformers: State-of-the-art natural language processing,” in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations, October 2020* (2020), pp. 38–45. <https://doi.org/10.18653/v1/2020.emnlp-demos.6>
48. M. Schreiner, “GPT-4 architecture, datasets, costs and more leaked”. <https://the-decoder.com/gpt-4-architecture-datasets-costs-and-more-leaked/>. Accessed 2024.
49. I. Loshchilov and F. Hutter, “Decoupled weight decay regularization,” in *Proceedings of the International Conference on Learning Representations* (2018); arXiv: 1711.05101. <https://doi.org/10.48550/arXiv.1711.05101>

Publisher’s Note. Pleiades Publishing remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.