

LLMs for Detection and Classification of Persuasion Techniques in Slavic Parliamentary Debates and Social Media Texts

Julia Jose and Rachel Greenstadt

Department of Computer Science and Engineering
New York University
New York, NY, USA
{jj3545, rg195}@nyu.edu

Abstract

We present an LLM-based method for the Slavic NLP 2025 shared task on detection and classification of persuasion techniques in parliamentary debates and social media. Our system uses OpenAI’s GPT models (gpt-4o-mini) and reasoning models (o4-mini) with chain-of-thought prompting, enforcing a ≥ 0.99 confidence threshold for verbatim span extraction. For subtask 1, each paragraph in the text is labeled "true" if any of the 25 persuasion techniques is present. For subtask 2, the model returns the full set of techniques used per paragraph. Across Bulgarian, Croatian, Polish, Russian, and Slovenian, we achieve Subtask 1 micro-F1 of 81.7%, 83.3%, 81.6%, 73.5%, 62.0%, respectively, and Subtask 2 F1 of 41.0%, 44.4%, 41.9%, 29.3%, 29.9%, respectively. Our system ranked in the top 2 for Subtask 2 and top 7 for Subtask 1.

1 Introduction

Persuasion techniques consist of rhetorical and psychological tactics (logical fallacies, emotional appeals, personal attacks, etc) that work to influence public opinion and behavior. In today’s information ecosystem, automated detection of these techniques helps facilitate fact-checking and content moderation. Da San Martino et al. (2019) defined 18 propaganda techniques widely used in recent news articles and consequently Da San Martino et al. (2020) invited submissions to detect instances of these techniques in English news articles at SemEval 2020 Task 11. Expanding this taxonomy to more broadly study news framing and persuasion in multiple languages, Piskorski et al. (2023) introduced 23 persuasion techniques and invited submissions to detect instances of these techniques in multi-lingual news articles at the paragraph level at SemEval 2023 Task 3. Similarly, Piskorski et al. (2024) expanded the detection task to span-level persuasion detection at CLEF 2024 (CheckThat! Lab Task 3) in French, German, Italian, and so on.

The Slavic NLP 2025 Workshop’s shared task on *Detection and Classification of Persuasion Techniques in Slavic Languages* focuses on texts in five Slavic languages: Bulgarian, Croatian, Polish, Russian, and Slovenian (Piskorski et al., 2025) in parliamentary debates and social media posts and participants are invited to submit solutions to two subtasks— (a) a binary detection problem where each paragraph is analyzed for the presence of any of the 25 persuasion techniques in Piskorski et al. (2025), and (b) a multi-class multi-label classification problem where the specific techniques within each paragraph must be identified.

Our Team’s (PSAL_NLP) submission to the task is a system that uses one of OpenAI’s GPT models, *gpt-4o-mini*, and one of their reasoning models, *o4-mini*. Our model includes a chain-of-thought prompt that checks each paragraph against each of the 25 techniques, instructs the model to return yes/no label per technique, extracting verbatim spans for any “yes” decisions, only accepting spans with confidence ≥ 0.99 , and returning those techniques as the final list of detected techniques per paragraph. We participated in both subtasks and experimented with both zero- and few-shot settings: in the few-shot setting, we added example phrases for each technique (obtained from the train dataset).

2 Related Work

Prior work in propaganda detection has produced several large datasets and detection mechanisms. For example, QProp (Barrón-Cedeno et al., 2019) contains 51,000 news articles (5,700 propaganda, 45,600 non-propaganda) labeled via Media Bias/Fact Check (MBFC) (Check, 2022), though distant supervision learns source signals rather than true propagandistic features. To address this, Da San Martino et al. (2019) developed the PTC dataset with phrase-level annotation of

18 propaganda techniques in English news articles. They also proposed a multi-granular neural network model designed to detect these techniques. Subsequently, [Dimitrov et al. \(2021\)](#) extended this work to multi-modal content in memes.

On the multi-lingual persuasion detection front, [Alam et al. \(2022\)](#) developed Arabic tweets dataset for propaganda detection. [Piskorski et al. \(2023\)](#) developed a dataset for 9 languages such as English, French, German, Italian, Polish, Russian, Georgian, Greek, and Spanish, with paragraph-level annotations of 23 persuasion techniques. As a follow-up, [Piskorski et al. \(2024\)](#) developed phrase-level annotations of the 23 techniques across this dataset and released a new dataset covering Arabic, Bulgarian, English, Portuguese, and Slovene.

Detection systems (for persuasion techniques) rely to a huge extent on transformer-based architectures. For example, [Jurkiewicz et al. \(2020\)](#) developed a RoBERTa-CRF model, achieving 62.07% micro-averaged F1 for techniques classification on the SemEval 2020 Task 11 dataset ([Da San Martino et al., 2020](#)). Likewise, [Purificato and Navigli \(2023\)](#) and [Hromadka et al. \(2023\)](#) used multilingual transformer models to achieve the top ranking in 7 of 9 languages in SemEval 2023 Task 3.

Researchers have also explored using LLMs for the detection of propaganda techniques in English news articles, only to find that they significantly underperform compared to the transformer-based counterparts ([Jose and Greenstadt, 2024](#); [Szwoch et al., 2024](#); [Hasanain et al., 2024](#)). However, their ability to detect such techniques at the paragraph level in Slavic languages remains unexplored and is the focus of our paper.

3 System Overview

Our system comprises of OpenAI’s LLMs, *gpt-4o-mini* and *o4-mini*, in a thoroughly prompt-engineered, zero- and few-shot setting. The temperature values were set to 0.1 for reproducibility, and every prompt begins with a system message instructing the model that it is an expert in Slavic persuasion techniques detection. We use chain-of-thought prompting to elicit step-by-step reasoning and maintain strict confidence thresholds. See Appendix for exact prompt.

3.1 Subtask 1: Binary Detection

We use *o4-mini* exclusively. Each paragraph is prefixed with definitions of all 25 persuasion tech-

Algorithm 1: Two-Pass CoT Prompt-Based Persuasion Technique Detection

Input: Paragraph p , technique groups T_1 and T_2 , confidence thres. $\tau = 0.99$

Output: Detected techniques list L

$L \leftarrow []$;

foreach $T_{\text{half}} \in \{T_1, T_2\}$ **do**

 Build a prompt that;

 (1) includes definitions of all techniques in T_{half} ;

 (2) asks to assign yes/no per technique;

 (3) asks to extract verbatim spans for “yes” at confidence $\geq \tau$;

 (4) asks to return a list $\hat{L}$ of techniques with confirmed spans;

$\hat{L} \leftarrow \text{ModelInference}(p, \text{prompt})$;

$L.\text{extend}(\hat{L})$;

return L ;

niques and instructed to output 1 if any are present (0 otherwise). We enforce a confidence score of 0.99, that is, the model should only return 1 if confidence score ≥ 0.99 . (Note: this is a semantic prompt cue, not a true calibrated cutoff). See Table 7 for exact prompt.

3.2 Subtask 2: Multi-Class Multi-Label Classification

For this subtask, we compare both *gpt-4o-mini* and *o4-mini* in zero- and few-shot settings. The zero-shot prompt lists all techniques (and definitions) plus output instruction, while few-shot prompt has example phrases for each technique (taken from the provided training dataset). We then apply a two-pass CoT prompt (Algorithm 1), feeding half the techniques per pass since feeding all 25 at once degraded performance (see Sections 4–5)—which boosts recall on less frequent techniques.

4 Experiments

Since we don’t train models, we used the additional training data (parliamentary debates) ([Piskorski et al., 2025](#)) as dev set, to evaluate model outputs. Section 5 contains official test-set results. In this section, we present additional dev-set experiments such as prompt ablations, zero/few-shot, and other settings that informed our final submission strategies for subtask 2. We used *gpt-4o-mini* for these

experiments. Croatian (HR) didn't have train/dev data and hence is only evaluated on test.

4.1 Prompt Engineering Ablations

For all languages except Croatian, we compare:

- **Simple Prompt:** list all 25 techniques + definitions (prompt in Table 9).
- **CoT+Th:** add chain-of-thought and confidence threshold ≥ 0.99 to the simple prompt. We add chain-of-thought reasoning by asking it to compare the given paragraph against each technique, assign "yes/no" per technique, and return verbatim spans for "yes" only if its confidence ≥ 0.99 . Note that this confidence score is not a true calibrated cut-off, but is intended to encourage the model to think about accuracy and only return ones it is highly confident about.

The CoT+Th prompt boosted precision, recall, and F1 compared to Simple prompt for all 4 languages. For Bulgarian (BG), there was a +6.2 pp increase in micro F1, +9.5 pp for Polish (PL), +6.4 pp for Russian (RU), and +11 pp for Slovenian (SI). See Table 4.

4.2 Zero- and Few-shot Ablations

Using *CoT+Th* as our base prompt, we compare:

- **Zero-shot:** same as CoT+Th prompt.
- **Few-shot:** zero-shot prompt + two positive examples per technique (from the additional training data on parliamentary debates). Since HR did not have additional train data to choose examples from, we only test HR in zero-shot settings. Furthermore, SI train data did not have examples of Slogans and Conversation_Killer so these techniques did not contain examples. We also omitted few-shot examples for Repetition across languages because our model processes only one paragraph at a time, and cross-paragraph repetition examples wouldn't apply unless the repetition was contained within that single paragraph. Tables 13, 14, 15, and 16 in the Appendix show the techniques and corresponding examples across languages.

Few-shot boosted all metrics compared to Zero-shot for all 4 languages. For BG, there was a +2.4 pp increase in micro F1 (with a small trade-off in micro-recall of -0.4), +3.1 pp for PL, +4.0 pp for RU, and +2.6 pp for SI. See Table 5.

4.3 Two-Pass vs. Single-Pass

We evaluate single-pass (all 25 techniques in one go) against two-pass split (techniques 1–13 and 14–25; see Algorithm 1). Using *CoT+Th* + *Few-Shot* as base prompts, we compare:

- **Single-pass:** same as CoT + Th + Few-shot prompt.
- **Two-pass:** techniques are split into 2 groups (techniques 1–13 and 14–25) to reduce the cognitive overload on the LLM. We selected the two-pass configuration heuristically to balance better reasoning capabilities with API cost: while a single pass overwhelms the model with too many techniques, finer-grained splits (e.g., three- or four-pass, or one-technique-per-call) would have substantially increased inference expense and were not feasible under our resource constraints.

In all languages, two-pass consistently improved recall and F1, but at the expense of precision (more false positives). See Table 6.

4.4 Implementation Details

We use OpenAI API to access gpt-4o-mini (gpt-4o-mini-2024-07-18), and o4-mini (o4-mini-2025-04-16). For gpt-4o-mini, we set *temperature* = 0.1, and *top_p* = 0.1. Each team could submit up to five runs per language per subtask.

For Subtask 1, we submitted one run per language using o4-mini.

For Subtask 2, we made four submissions per language: (a) CoT+Th+Few-shot+Two-pass using gpt-4o-mini, (b) CoT+Th+Zero-shot+Two-pass using gpt-4o-mini, (c) CoT+Th+Few-shot+Two-pass using o4-mini, and (d) CoT+Th+Few-shot+Single-pass using o4-mini.

We ran both models on our best dev-set prompt (CoT+Th+Few-shot+Two-pass). Because Two-pass produced more false positives on the dev set than Single-pass, we also ran both these prompts with o4-mini to see if the reasoning model exhibited the same pattern. Lastly, because we observed a slight decrease in micro-recall for Few-shot compared to Zero-shot gpt-4o-mini in Bulgarian, we also submitted a CoT+Th+Zero-shot+Two-pass run using gpt-4o-mini.

Language	Team	Run	Rank	F ₁
BG	FactUE	2	1	0.878
BG	PSAL_NLP	1	6	0.817
HR	FactUE	1	1	0.955
HR	PSAL_NLP	1	6	0.833
PL	oplot	1	1	0.903
PL	PSAL_NLP	1	7	0.816
RU	INSAntive	3	1	0.871
RU	PSAL_NLP	1	7	0.734
SI	UFAL4DEM	3	1	0.856
SI	PSAL_NLP	1	7	0.619

Table 1: Comparison of PSAL_NLP (ours) and the top-ranked systems on Subtask 1 (official test set).

Lang	Team	Run	Rank	Micro-F ₁	Macro-F ₁
BG	PSAL_NLP	1	1	0.410	0.319
BG	INSAntive		2	0.344	0.208
HR	Gradient-Flush		1	0.491	0.359
HR	PSAL_NLP	3	2	0.443	0.320
PL	PSAL_NLP	3	1	0.419	0.296
PL	Gradient-Flush		2	0.409	0.276
RU	INSAntive		1	0.295	0.158
RU	PSAL_NLP	2	2	0.292	0.207
SI	Gradient-Flush		1	0.323	0.190
SI	PSAL_NLP	2	2	0.298	0.263

Table 2: Comparison of PSAL_NLP (ours) and top systems on Subtask 2 (test data). Run ID mapping: 1=CoT+Th+Few-shot+Two-pass gpt-4o-mini, 2=CoT+Th+Few-shot+Two-pass o4-mini, 3=CoT+Th+Zero-shot+Two-pass gpt-4o-mini, 4=CoT+Th+Few-shot+Single-pass o4-mini. HR uses zero-shot only.

5 Results

5.1 Subtask 1: Binary Detection

For Subtask 1, we used a straightforward, definition-only prompt (see Table 7) to establish a consistent baseline across all five languages. Although we required the model to return “1” only if its confidence score ≥ 0.99 , the absence of chain-of-thought reasoning, few-shot examples, and the use of a single-pass prompt containing all 25 techniques likely constrained its reasoning capacity and recall, resulting in substantially poorer performance compared to Subtask 2.

We made one submission per language. Table 1 shows our results, compared to the top-performing team. Our team, PSAL_NLP, ranked sixth for Bulgarian and Croatian and seventh for all other languages. Our system did not outperform the XLM-RoBERTa baseline (Piskorski et al., 2025) nor the top-performing systems that relied on fine-tuned

Lang	Run	Micro-F ₁	Macro-F ₁
BG	1	0.410	0.319
BG	2	0.373	0.305
BG	3	0.403	0.318
BG	4	0.358	0.300
HR	2	0.422	0.309
HR	3	0.44	0.320
HR	4	0.422	0.309
PL	1	0.397	0.299
PL	2	0.390	0.315
PL	3	0.419	0.296
PL	4	0.379	0.297
RU	1	0.240	0.187
RU	2	0.292	0.207
RU	3	0.222	0.201
RU	4	0.281	0.199
SI	1	0.282	0.165
SI	2	0.298	0.263
SI	3	0.240	0.190
SI	4	0.214	0.153

Table 3: PSAL_NLP’s official runs for Subtask 2. Run ID mapping: 1=CoT+Th+Few-shot+Two-pass gpt-4o-mini, 2=CoT+Th+Few-shot+Two-pass o4-mini, 3=CoT+Th+Zero-shot+Two-pass gpt-4o-mini, 4=CoT+Th+Few-shot+Single-pass o4-mini. HR uses zero-shot only.

transformer models (from BERT family). We noticed significant drops in recall compared to all other systems. Incorporating a two-pass method using chain-of-thought reasoning and few-shot examples, as in Subtask 2, could help recover some of this lost recall in future work.

5.2 Subtask 2: Multi-Class Multi-Label Classification

Table 2 compares our results with the top-ranked system for this subtask. Our team, PSAL_NLP, ranked first for Bulgarian and Polish, and second for Croatian, Russian, and Slovenian. Out of three of our second-place finishes, two (RU and SI) used CoT+Th+Few-shot+Two-pass with o4-mini, and HR used CoT+Th+Zero-shot+Two-pass with gpt-4o-mini (since HR did not have train data to select few-shot examples from). Notably, our gpt-4o-mini CoT+Th+Few-shot+Two-pass model was the top-performing system for Bulgarian, and our gpt-4o-mini CoT+Th+Zero-shot+Two-pass model was the top-performing system for Polish.

Table 3 summarizes all four of our official submissions per language (precision/recall in Table 12 in the Appendix). We observe the following:

- For BG, RU, and SI, few-shot outperformed zero-shot. For PL, zero-shot was better by

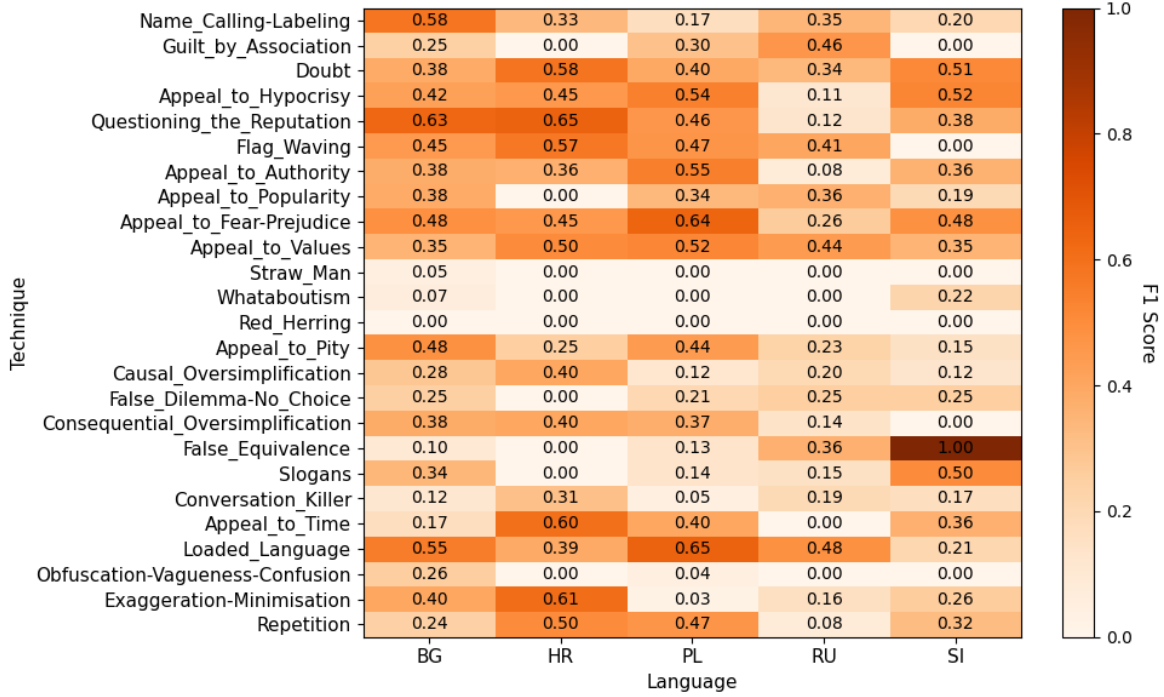


Figure 1: Per-Class (Persuasion Technique) F₁ Scores for PSAL_NLP (Best Run per Language)

+2.2 pp micro-F1. (HR was only tested zero-shot due to lack of training data)

- For all 5 languages, two-pass outperformed single-pass, also replicating our dev-set finding where two-pass generates more false positives (lower micro-precision) than single-pass.
- For HR, RU, and SI, o4-mini outperformed gpt-4o-mini, whereas for BG and PL, gpt-4o-mini outperformed o4-mini.

These results show that chain-of-thought prompts with high confidence thresholds and few-shot examples enable LLMs to outperform fine-tuned transformer baselines on persuasion technique classification. For example, against Team *INSANTive* (Wang et al., 2025) that used XLM-RoBERTa with LLM-generated explanations of techniques, we obtain improved micro and macro-F1 in BG, HR, PL, and SI. Likewise, compared to Team *GradientFlush* (Senichev et al., 2025) that fine-tuned multilingual transformer models on CLEF 2024 CheckThat! Lab data (Piskorski et al., 2024) alongside LLM-generated translations of instances of techniques, we achieved better performance for BG, PL, and RU. See Table 12 for full evaluation metrics.

Figure 1 shows per-class F1 for best model per language (see Table 2). Techniques like Loaded

Language, Questioning the Reputation, Appeal to Fear/Prejudice, Flag-Waving, and Doubt have higher F1 than techniques such as Strawman, Red-Herring, and Whataboutism, likely due to their higher prevalence (Piskorski et al., 2023).

We observe that for techniques such as Straw Man, Whataboutism, and Red Herring, there is a near-zero F1 across all languages. This potentially stems from their need for broader context for analysis (multiple paragraphs) since misrepresentations and topic diversions cannot be judged from a single paragraph alone. This implies that truly context-dependent techniques such as these would require broader contexts for accurate judgment.

5.3 Conclusion

We presented an LLM-based method using gpt-4o-mini and o4-mini with chain-of-thought prompting and ≥ 0.99 confidence thresholding to detect 25 persuasion techniques in five Slavic languages, as part of Slavic NLP 2025 shared task. Our system was ranked in the top 2 for the technique classification task and ranked 7th for technique detection. Ablation studies confirm that chain-of-thought, few-shot examples, and a two-pass strategy are key to improving performance.

Limitations

Our method uses OpenAI’s “mini” models; larger models that are not compressed might outperform these models. But they could incur higher costs and latency. Furthermore, future work could look into fine-tuning these models to improve performance.

The “confidence ≥ 0.99 ” instruction is a prompt-level nudge rather than a well-calibrated probability cutoff, so it does not guarantee statistically meaningful uncertainty estimation.

For subtask 1, the definition-only, single-pass prompt resulted in low recall and substantially poorer binary-detection performance, highlighting the importance of extensive prompt engineering efforts as seen in subtask 2.

Finally, for subtask 2, we heuristically used a two-pass split to balance performance and API cost; finer-grained or per-class passes might be more beneficial, however, these were computationally infeasible under our constraints.

References

- Firoj Alam, Hamdy Mubarak, Wajdi Zaghouni, Giovanni Da San Martino, and Preslav Nakov. 2022. [Overview of the WANLP 2022 shared task on propaganda detection in Arabic](#). In *Proceedings of the Seventh Arabic Natural Language Processing Workshop (WANLP)*, pages 108–118, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.
- Alberto Barrón-Cedeno, Israa Jaradat, Giovanni Da San Martino, and Preslav Nakov. 2019. Propopy: Organizing the news based on their propagandistic content. *Information Processing & Management*, 56(5):1849–1864.
- Media Bias/Fact Check. 2022. Questionable sources. <https://mediabiasfactcheck.com/fake-news/>, as of February 15, 2023.
- Giovanni Da San Martino, Alberto Barrón-Cedeño, Henning Wachsmuth, Rostislav Petrov, and Preslav Nakov. 2020. [SemEval-2020 task 11: Detection of propaganda techniques in news articles](#). In *Proceedings of the Fourteenth Workshop on Semantic Evaluation*, pages 1377–1414, Barcelona (online). International Committee for Computational Linguistics.
- Giovanni Da San Martino, Yu Seunghak, Alberto Barrón-Cedeno, Rostislav Petrov, Preslav Nakov, and 1 others. 2019. Fine-grained analysis of propaganda in news article. In *Proceedings of the 2019 conference on empirical methods in natural language processing and the 9th international joint conference on natural language processing (EMNLP-IJCNLP)*, pages 5636–5646. Association for Computational Linguistics.
- Dimitar Dimitrov, Bishr Bin Ali, Shaden Shaar, Firoj Alam, Fabrizio Silvestri, Hamed Firooz, Preslav Nakov, and Giovanni Da San Martino. 2021. [Detecting propaganda techniques in memes](#). 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)*, pages 6603–6617, Online. Association for Computational Linguistics.
- Maram Hasanain, Fatema Ahmed, and Firoj Alam. 2024. Can gpt-4 identify propaganda? annotation and detection of propaganda spans in news articles. *arXiv preprint arXiv:2402.17478*.
- Timo Hromadka, Timotej Smolen, Tomas Remis, Branislav Pecher, and Ivan Srba. 2023. [KInITVer-aAI at SemEval-2023 task 3: Simple yet powerful multilingual fine-tuning for persuasion techniques detection](#). In *Proceedings of the 17th International Workshop on Semantic Evaluation (SemEval-2023)*, pages 629–637, Toronto, Canada. Association for Computational Linguistics.
- Julia Jose and Rachel Greenstadt. 2024. [Are large language models good at detecting propaganda?](#) In *Workshop Proceedings of the 18th International AAAI Conference on Web and Social Media*, 5th International Workshop on Cyber Social Threats (CySoc 2024). AAAI Press.
- Dawid Jurkiewicz, Łukasz Borchmann, Izabela Kosmala, and Filip Graliński. 2020. Applcaai at semeval-2020 task 11: On roberta-crf, span cls and whether self-training helps them. *arXiv preprint arXiv:2005.07934*.
- Jakub Piskorski, Dimitar Dimitrov, Filip Dobrani’c, Marina Ernst, Jacek Haneczok, Ivan Koychev, Ivo Moravski, Nikola Ljubešić, Michał Marci’nczuk, Arkadiusz Modzelewski, and Roman Yangarber. 2025. SlavicNLP 2025 Shared Task: Detection and Classification of Persuasion Techniques in Parliamentary Debates and Social Media. In *Proceedings of the 10th Workshop on Slavic Natural Language Processing 2025 (SlavicNLP 2025)*, Vienna, Austria. Association for Computational Linguistics.
- Jakub Piskorski, Alípio Jorge, Maria da Purificação Silvano, Nuno Guimarães, Ana Filipa Pacheco, and Nana Yu. 2024. Overview of the clef-2024 checkthat! lab task 3 on persuasion techniques. In *Proceedings of the 15th Conference and Labs of the Evaluation Forum (CLEF 2024)*, pages 299–310, Grenoble, France.
- Jakub Piskorski, Nicolas Stefanovitch, Giovanni Da San Martino, and Preslav Nakov. 2023. [SemEval-2023 task 3: Detecting the category, the framing, and the persuasion techniques in online news in a multilingual setup](#). In *Proceedings of the 17th International Workshop on Semantic Evaluation (SemEval-2023)*, pages 2343–2361, Toronto, Canada. Association for Computational Linguistics.
- Antonio Purificato and Roberto Navigli. 2023. [APatt at SemEval-2023 task 3: The sapienza NLP system for](#)

ensemble-based multilingual propaganda detection. In *Proceedings of the 17th International Workshop on Semantic Evaluation (SemEval-2023)*, pages 382–388, Toronto, Canada. Association for Computational Linguistics.

Sergey Senichev, Aleksandr Boriskin, Nikita Krayko, and Daria Galimzianova. 2025. Gradient flush at slavic nlp 2025 task: Leveraging slavic bert and translation for persuasion techniques classification. In *Proceedings of the 10th Workshop on Slavic Natural Language Processing 2025 (SlavicNLP 2025)*, Vienna, Austria. Association for Computational Linguistics.

Joanna Szwoch, Mateusz Staszko, Rafal Rzepka, and Kenji Araki. 2024. Limitations of large language models in propaganda detection task. *Applied Sciences*, 14(10):4330.

Yutong Wang, Diana Nurbakova, and Sylvie Calabretto. 2025. Team insantive at slavichlp-2025 shared task: Data augmentation and enhancement via explanations for persuasion technique classification. In *Proceedings of the 10th Workshop on Slavic Natural Language Processing 2025 (SlavicNLP 2025)*, Vienna, Austria. Association for Computational Linguistics.

Lang	Prompt	Micro P	Micro R	Macro P	Macro R	Micro F ₁	Macro F ₁
BG	Simple	0.4223	0.3544	0.2517	0.2260	0.3854	0.1944
	CoT + Th	0.4521	0.4420	0.3094	0.3759	0.4470	0.2954
PL	Simple	0.3603	0.2562	0.2971	0.2079	0.2994	0.1890
	CoT + Th	0.4736	0.3383	0.3662	0.2677	0.3946	0.2730
RU	Simple	0.1556	0.2278	0.1100	0.1697	0.1849	0.0984
	CoT + Th	0.2067	0.3122	0.1883	0.2498	0.2487	0.1652
SI	Simple	0.5000	0.2905	0.2713	0.2602	0.3675	0.1981
	CoT + Th	0.5928	0.4006	0.3922	0.3391	0.4781	0.3037

Table 4: Ablation results on the development set: Simple prompt vs. CoT + Th prompt across four languages.

Lang	Prompt	Micro P	Micro R	Macro P	Macro R	Micro F ₁	Macro F ₁
BG	Zero-shot	0.4521	0.4420	0.3094	0.3759	0.4470	0.2954
	Few-shot	0.5095	0.4379	0.5269	0.4472	0.4710	0.3805
PL	Zero-shot	0.4736	0.3383	0.3662	0.2677	0.3946	0.2730
	Few-shot	0.5362	0.3530	0.4622	0.2931	0.4257	0.3223
RU	Zero-shot	0.2067	0.3122	0.1883	0.2498	0.2487	0.1652
	Few-shot	0.2492	0.3418	0.2776	0.2970	0.2883	0.2257
SI	Zero-shot	0.5928	0.4006	0.3922	0.3391	0.4781	0.3037
	Few-shot	0.6313	0.4190	0.4056	0.3433	0.5037	0.3127

Table 5: Zero-Shot vs. Few-Shot ablation on the development set for Subtask 2.

Lang	Prompt	Micro P	Micro R	Macro P	Macro R	Micro F ₁	Macro F ₁
BG	Single-pass	0.5095	0.4379	0.5269	0.4472	0.4710	0.3805
	Two-pass	0.4476	0.5132	0.3840	0.5117	0.4782	0.3810
PL	Single-pass	0.5362	0.3530	0.4622	0.2931	0.4257	0.3223
	Two-pass	0.4652	0.4171	0.4213	0.3437	0.4398	0.3363
RU	Single-pass	0.2492	0.3418	0.2776	0.2970	0.2883	0.2257
	Two-pass	0.2242	0.4219	0.2207	0.3576	0.2928	0.2398
SI	Single-pass	0.6313	0.4190	0.4056	0.3433	0.5037	0.3127
	Two-pass	0.5921	0.5015	0.5849	0.5444	0.5430	0.4932

Table 6: Comparison of single-pass vs. two-pass prompting on the development set.

Table 7: Prompt used for all languages in Subtask 1

System: You are an expert at determining if a given text fragment contains one or more persuasion techniques in a given taxonomy of persuasion techniques.

User: You are given a text fragment and the following list of persuasion techniques. Your task is to determine if the text fragment contains one or more of these persuasion techniques.

List of Persuasion Techniques:

1. Name_Calling-Labeling: a form of argument in which loaded labels are directed at an individual or a group, typically ..
2. ...

Output Instructions:

- Return 1 if the text fragment contains one or more of the persuasion techniques from the list above.
- Return 0 if it does not.
- Only return 1 if confidence ≥ 0.99 .

Text Fragment to Analyze:

Language	Run	Accuracy (%)	Precision (%)	Recall (%)	F ₁ (%)
BG	1	82.5	93.7	72.5	81.7
HR	1	86.5	96.2	73.5	83.3
PL	1	79.8	96.5	70.8	81.6
RU	1	66.9	92.8	60.8	73.5
SI	1	81.1	88.2	47.8	62.0

Table 8: PSAL_NLP Subtask 1 performance by language.

Table 9: Simple Prompt used for Subtask 2

System: You are an expert at determining if a given text fragment contains one or more persuasion techniques in a given taxonomy of persuasion techniques.

User: You are given a text fragment and the following list of persuasion techniques. Your task is to identify the persuasion techniques that this fragment uses.

List of Persuasion Techniques:

1. Name_Calling-Labeling: a form of argument in which loaded labels are directed at an individual or a group, typically ..
2. ...

Output Instructions:

- Return a Python list containing all the persuasion technique(s) that the following text fragment uses.

Table 10: CoT+Th Prompt used for Subtask 2 (This is also the Zero-Shot Prompt)

System: You are an expert at determining if a given text fragment contains one or more persuasion techniques in a given taxonomy of persuasion techniques.

User: You are given a text fragment and the following list of persuasion techniques. Your task is to identify the persuasion techniques that this fragment uses.

List of Persuasion Techniques:

1. Name_Calling-Labeling: a form of argument in which loaded labels are directed at an individual or a group, typically ..
2. ...

Output Instructions:

- For each technique listed above, check if the text fragment uses the technique, and return yes or no beside the technique name, along with the detected span(s) (verbatim) that correspond to the technique.
- You should only return yes if you are extremely confident about your judgment (confidence $\geq$ 0.99).
- At the end of your output, return a list of all the persuasion technique(s) that you said yes to, as a python list.

Table 11: Few-Shot Prompt used for Subtask 2 (Add examples beside each technique)

System: You are an expert at determining if a given text fragment contains one or more persuasion techniques in a given taxonomy of persuasion techniques.

User: You are given a text fragment and the following list of persuasion techniques. Your task is to identify the persuasion techniques that this fragment uses.

List of Persuasion Techniques:

1. Name_Calling-Labeling: a form of argument in which loaded labels are directed at an individual .. Examples:
2. ...

Output Instructions:

- For each technique listed above, check if the text fragment uses the technique, and return yes or no beside the technique name, along with the detected span(s) (verbatim) that correspond to the technique.
- You should only return yes if you are extremely confident about your judgment (confidence $\geq$ 0.99).
- At the end of your output, return a list of all the persuasion technique(s) that you said yes to, as a python list.

Lang	Run	Accuracy	Micro P	Micro R	Micro F ₁	Macro P	Macro R	Macro F ₁
BG	1	44.4%	39.7%	42.4%	41.0%	35.8%	38.4%	32.0%
BG	2	45.2%	40.8%	34.3%	37.3%	37.5%	32.2%	30.5%
BG	3	43.9%	39.2%	41.4%	40.3%	32.5%	38.9%	31.8%
BG	4	46.1%	49.1%	28.2%	35.8%	43.6%	27.7%	30.0%
HR	2	54.1%	56.4%	33.8%	42.2%	43.8%	27.7%	30.9%
HR	3	54.1%	46.2%	42.7%	44.4%	39.6%	34.6%	32.0%
HR	4	54.1%	56.4%	33.8%	42.2%	43.8%	27.7%	30.9%
PL	1	36.6%	48.1%	33.7%	39.7%	42.6%	28.3%	29.9%
PL	2	36.9%	49.4%	32.3%	39.0%	39.4%	28.0%	31.5%
PL	3	35.5%	45.2%	39.1%	41.9%	33.6%	31.4%	29.7%
PL	4	38.1%	56.0%	28.7%	37.9%	43.0%	24.3%	29.7%
RU	1	22.4%	23.3%	24.9%	24.0%	17.9%	25.4%	18.7%
RU	2	22.9%	29.1%	29.4%	29.3%	20.6%	24.3%	20.8%
RU	3	21.7%	20.4%	24.4%	22.2%	21.3%	27.5%	20.1%
RU	4	24.7%	32.1%	25.0%	28.1%	24.6%	22.0%	19.9%
SI	1	66.5%	33.9%	24.1%	28.2%	24.7%	21.2%	16.5%
SI	2	66.5%	40.0%	23.9%	29.9%	34.3%	28.0%	26.3%
SI	3	66.3%	30.4%	19.8%	24.0%	25.9%	23.6%	19.0%
SI	4	66.3%	34.7%	15.5%	21.4%	27.8%	13.1%	15.3%

Table 12: PSAL_NLP Subtask 2 runs on test sets, showing accuracy, precision, recall, and F₁ for each run ID. Run ID mapping: 1=CoT+Th+Few-shot+Two-pass gpt-4o-mini, 2=CoT+Th+Few-shot+Two-pass o4-mini, 3=CoT+Th+Zero-shot+Two-pass gpt-4o-mini, 4=CoT+Th+Few-shot+Single-pass o4-mini. HR uses zero-shot only.

Table 13: Few-shot Examples Used for BG Subtask 2

Technique	Example(s)
Name _ Calling-Labeling	"абсолютно безсмислено", "по-опасно и по-срамно"
Guilt _ by _ Association	"Това не е просто проруска партия, това са директно думите на Кремъл, изречени от тази трибуна", "Едно от нещата, които казвате, това е класическа руска опорка"
Doubt	"Учудвам се, че сте председател на Комисията по отбрана, господин Гаджев!", "Виждате ли какъв аргумент?"
Appeal _ to _ Hypocrisy	"Искате да влезете в Шенген, а дори не можете да опазите границата на България", "А докато беше в БСП до миналата година, дали беше проруска партия?!"
Questioning _ the _ Reputation	"Малко фактология защо стигаме до това безумие от страна на управляващите", "той не носи отговорност за глупостта на Тагарев или на Денков"
Flag _ Waving	"но не можете да нарежете паметта на българския народ", "Бъдете наистина българи!"
Appeal _ to _ Authority	"Цитирам само официални източници, за да няма двусмислие, да няма обвинения", "Съединените щати са първи в това отношение"
Appeal _ to _ Popularity	"Всъщност огромна, огромна е подкрепата на целия демократичен цивилизован свят за Украйна", "това е проблемът на проблемите, който в момента вълнува и света, и Европа, респективно и България"
Appeal _ to _ Values	"Всяка нова година идва с нова надежда, с нови очаквания, с искане за перспектива, сигурност и стабилност", "Това е най-добрият начин да демонстрираме нашето единство и солидарност по отношение на възпирането и отбраната и споделянето на тежестите"
Appeal _ to _ Fear-Prejudice	"Ей, хора, с тази тема не си играйте, ще взривите държавата!", "Воюването означава агресия и атака"
Straw _ Man	"той дава пресконференция, на която с половин уста каза – между другото, само където не", "Означава ли това, че всъщност тук, както каза някой, ще се възстановява Османската империя? От ПП-ДБ го казва"
Red _ Herring	"ще Ви направя един цитат: „С лека ръка фашистите евроатлантици“ – и така нататък продължава цитатът", "Ще се върна назад по време на предизборната кампания, когато ние от БСП предупреждавахме, че Политическа партия ГЕРБ и „Продължаваме Промяната – Демократична България“ след изборите ще се съберат и ще управляват заедно"
Whataboutism	"Събираме капачки за децата, облагаме бизнеса с безумни данъци и в същото време харчим милиарди в посока към Украйна в една братоубийствена война", "Има един доклад от 2022 г., не е лошо да го прочете, така, както сте отишли в Секретното деловодство да четете какво точно е изпращано в Украйна"
Appeal _ to _ Pity	"Според Организацията на обединените нации са избити 14 хиляди етнически руснаци и граждани на Украйна", "Така че не е непредизвикана агресията, убити са хиляди хора"
Causal _ Oversimplification	"Спират тая подкрепа, значи и България трябва да спре своята подкрепа за Украйна", "Госпожо Назарян, очевидно отивате на избори, защото вчера не можахте да сформирате кабинет"
False _ Dilemma-No _ Choice	"трябва да влезе единствено и само с оставката си, или да бъде арестуван тук в залата на Народното събрание за национално предателство", "Няма нормален човек в Европейския съюз, който да вярва в това"
Consequential _ Oversimplification	"Тоест, ако ние постоянно говорим за конфликта, няма как да не обвържем военната професия с него", "Повтарям Ви – след това решение следващата стъпка е тази, за която Ви казах в самото начало"
False _ Equivalence	"Нека министър-председателят академик Денков да не бъде в ролята на Богдан Филев – 1944 г.", "Вие като техни комунистически отрочета правите абсолютно същото"
Slogans	"Искаме обяснение за това!", "Замислете се за това!"
Conversation _ Killer	"Лъжеш!", "Това искате Вие"
Appeal _ to _ Time	"Главният прокурор да се вземе в ръце и незабавно да вземе мерки", "И аз съм сигурен, че и това ще стане, но въпросът е кога ще стане, защото има голяма разлика"
Loaded _ Language	"тежки словесни", "уволнявала, махала, премахвала и наказвала"
Obfuscation-Vagueness-Confusion	"Има неща, които на тази среща България ще трябва да потвърди, или да не се съгласи с неща, неща, за които те нямат санкции от Народното събрание"

continued on next page

continued from previous page

Technique	Example(s)
Exaggeration-Minimisation	"да превърнат България в най-големия бежански лагер в Европа", "Самите европейски държави си подхвърлят един на друг нелегалните емигранти като топки за пинг-понг"
Repetition	—

Table 14: Few-shot Examples Used for PL Subtask 2

Technique	Example(s)
Name_Calling-Labeling	"upolityczniony trybunał", "bulwersujący wyrok"
Guilt_by_Association	"Jesteście forpoczta cywilizacji śmierci", "płk Dusza, ten, który negocjował umowę z FSB, z rosyjską służbą specjalną"
Doubt	"Niszczycie relacje polsko-amerykańskie, angażując się po jednej stronie sporu politycznego w Ameryce", "starsi panowie w garniturach nie będą mówić kobietom, co mają robić ze sobą"
Appeal_to_Hypocrisy	"został uchwalony program wieloletni, który przewidywał także konkretne środki na realizację stopnia wodnego w Siarzewie, a pani minister teraz mówi, że wszystko jest winą PiS-u", "Jakie to wygodne"
Questioning_the_Reputation	"podejmowaliśmy stosowne działania, państwo nas krytykowaliście", "Ich fałszywy heroizm i wygodnictwo nie byłyby jednak możliwe, gdyby nie wsparcie ruchu aborcyjnego, który rośnie w siłę od dekad"
Flag_Waving	"Mówmy polskim głosem w Unii Europejskiej wspólnie, razem", "My potrzebujemy w Polsce dobrego prawa"
Appeal_to_Authority	"W poniedziałek w holenderskim parlamencie przy udziale holenderskiej minister zdrowia odbyła się debata na temat dostępu do tabletek aborcyjnych dla osób z Polski", "Mówię i występuję tutaj jako ojciec siedmiorga dzieci"
Appeal_to_Popularity	"W czasie wojen przemysłowych współczynnik ten w wielu krajach wynosił ok. 7 %, a my marzymy o 3 %", "W ostatnim miesiącu w Polsce 9 tys. kobiet przerwało ciążę"
Appeal_to_Values	"Rozmawiamy o życiu", "One powinny mieć prawo do decydowania"
Appeal_to_Fear-Prejudice	"bronią polskiej granicy wschodniej przed zalewem nielegalnej imigracji", "kiedy dochodzi do śmierci, jak to było w przypadku pani Izabeli czy pani Doroty"
Straw_Man	"Ale prawda jest taka, że mam nieodparte wrażenie, że lewicy chodzi tylko o dyskusję, tak samo jak prawicy, a problemy kobiet do tej pory są nierozwiązane", "Najważniejsze to urodzić"
Red_Herring	"My niewiele mniej wydaliśmy na laptopy+ dla czwartoklasistów", "mimo że tylko jedna z nich meldowała, że Rosjanie wejdą na Ukrainę"
Whataboutism	"Chociażby dlaczego pana nie było na Monte Cassino kilka dni temu?", "bo w polskich szpitalach wciąż dzisiaj łatwo o relikwie, ale trudno o aborcję"
Appeal_to_Pity	"Podobno powodem są jakieś limity, limity w otwartości i w uśmiechu", "Koniec ze zmuszaniem kobiet do heroizmu"
Causal_Oversimplification	"I to my, osoby na tej sali, możemy sprawić, żeby takie tragedie jak Izy z Pszczyny więcej się po prostu nie powtarzały", "Na razie dzięki ustawie aborcyjnej muszą leżeć i nic nie mogą zrobić"
False_Dilemma-No_Choice	"Tylko kobieta i lekarz powinni decydować o przebiegu ciąży", "Lewica składa ustawę dotyczącą dekryminalizacji pomocy w aborcji. Bo aborcje były, są i będą"
Consequential_Oversimplification	"Dzięki dostępności aborcji farmakologicznej można wcześniej, a tym samym bezpieczniej przerwać ciążę", "Zakaz aborcji zabija i nie likwiduje aborcji"
False_Equivalence	"bo skoro 460 posłów powinno zagłosować w tej sprawie, to czy 30 mln Polaków to nie jest więcej niż 460 posłów?", "Katoliczki mogą nie chcieć przerywać ciąży, to jest ich wybór. A te kobiety, które chcą – również powinny go mieć"
Slogans	"ze zmuszaniem kobiet do heroizmu", "nie bój się, nie jesteś sama, pomogę ci"
Conversation_Killer	"są całkowicie nieakceptowalne i powinny być odrzucone już w pierwszym czytaniu", "Bo aborcje w Polsce były, są, będą"
Appeal_to_Time	"Dzisiaj jest ten moment, kiedy jeszcze możecie zmienić zdanie", "Nie pierwsza, ale mam nadzieję, że jedna z ostatnich"
Loaded_Language	"wymazywania kobiet", "Kto by tam się przejmował"
Obfuscation-Vagueness-Confusion	"Wyrokiem, który tak naprawdę nie jest wyrokiem, Trybunału Konstytucyjnego, który tak naprawdę nie jest Trybunałem Konstytucyjnym", "Bo, szanowni państwo, życie to nie jest Instagram"
Exaggeration-Minimisation	"Bo przecież kobiety w Polsce wciąż są w bardzo niebezpiecznym momencie", "a nie trwać w dalszym dręczeniu, w dalszym straszeniu i w dalszym upokarzaniu"
Repetition	—

Table 15: Few-shot Examples Used for RU Subtask 2

Technique	Example(s)
Name_Calling-Labeling	"Зеленский не политик и не государственный деятель", "США и страны Запада всегда были одержимы идеей мирового господства"
Guilt_by_Association	"Через «Пласт» в своё время прошло практически всё командование УПА — Бандера, Шухевич, Кук и др", "Терроризм радикального ислама, казалось, был уже забыт в России, но в этот раз без помощи Украины не обошлось"
Doubt	"Ни нормальной медкомиссии, ни запроса в ПНД по месту жительства не было", "Но наша власть, похоже, обманывается не только в сфере внешней политики и миграции - есть еще много интересных направлений"
Appeal_to_Hypocrisy	"А то один сидит, а другие, кто миллиардами воровал пошли на повышение или пересаживают на СВО бурю, а потом еще УВБД получают и все льготы, будут говорить как они героически защищали Родину", "Если администрация Байдена против кого-то и хочет на данном этапе ввести санкции в связи с «Северными потоками», так это надо делать против Байдена и Нуланд — к уничтожению данного проекта призывали именно они"
Questioning_the_Reputation	"Эти "тарые новые люди" это не команда Трампа, а команда, использующая Трампа (Маск, Тиль и проч)", "В том числе потому, что пресечение нелегальной миграции или даже ограничения стихийной миграции выгодно и властям Узбекистана"
Flag_Waving	"Эта история говорит о том, что я никогда не буду молиться за души врага", "Территория Курской области в ближайшее время будет полностью освобождена от противника"
Appeal_to_Authority	"Сталин в войну всех нужных вытаскивал, чтобы трудились", "Но как только Россия ответит, то молчание закончится"
Appeal_to_Popularity	"Саммит показал, что Западу не получилось сделать Россию изгоем на международной арене. Наоборот, к нам тянутся многие влиятельные региональные игроки как к одному из главных акторов в большой геополитической игре", "Сегодня помимо США в той или иной форме законодательство об иноагентах, но более либеральное, чем американское, действует в Британии, Израиле, Австралии, других странах. В России, кстати, закон принят в 2012 году"
Appeal_to_Values	"В исторической России такого не было. Напротив, Россия во все периоды своего существования защищала свои символы", "Набиулиной решило отказаться «от изображения объектов религиозного назначения» читай, от крестов"
Appeal_to_Fear-Prejudice	"А если закон не работает, так будет линчевание", "Контракт - добровольно или принудительно В воинских частях заставляют срочников и мобилизованных подписывать контракт"
Straw_Man	"Если это не игра, то фактически, Трамп дает карт-бланш Путину — доводи свое дело до конца, а мы будем наблюдать", "Хотя левые поверят и в это"
Red_Herring	"Внес законопроект, который наделяет английский статусом языка «международного общения» на Украине", "Одних суверенных резервов, оказавшихся теперь под арестом на Западе, было 300 миллиардов"
Whataboutism	"История закона об иноагентах началась в 1938 году в США. Надо сказать, его нормы там до сих пор остаются самыми жесткими в мире", "Резонанс от этого интервью был серьезный, в США его эффект пытались перебить информацией о выводе Россией в космос спутника с ядерным оружием, что, естественно, оказалось фейком"
Appeal_to_Pity	"Недавний случай: инвалид 3 группы (по умственной отсталости), состоит на учёте в Рязанском психдиспансере"
Causal_Oversimplification	"Сверху рисуют план по набору на контракт, регионы берут под козырек и привлекают людей выплатами, мошенники зарабатывают", "Все, кроме правительства, для которого важны суммы потраченных денег и красивые отчеты, а не количество детей"
False_Dilemma-No_Choice	"А вот коррупционеры, которые отсиживаются в добровольческих формированиях, чтобы их не закрыли, нужны", "С одной стороны, продолжающееся СВО, которое никак нельзя заканчивать до того, как его реальные задачи будут реализованы"
Consequential_Oversimplification	"Любой кризис, если его однажды не переломить, не остановить и не повернуть вспять, заканчивается катастрофой", "На то место, где была Россия, придут другие народы, жизненное пространство будет кем-то заполнено"

continued on next page

continued from previous page

Technique	Example(s)
False_Equivalence	"Если это так, хотелось бы чтобы до нее довели картинки из Сирии, которая прямо сейчас, на глазах превращается в террористический анклав"
Slogans	"Не нужно ужесточать закон, закон должен работать", "... Глас народа — глас Божий"
Conversation_Killer	"Что это за надругательство над здравым смыслом уходящей администрации?", "Гнетущее, жуткое ожидание"
Appeal_to_Time	"Однако нынешняя ситуация, когда аборт приравнивают к дежурной операции вроде удаления аппендицита, продолжаться не может", "Вымирание ускоряется"
Loaded_Language	"Рецепты известны и в целом все с ними согласны", "в штурм в один конец"
Obfuscation-Vagueness-Confusion	"Что за заявления? Неясно", "которые отличаются особой дерзостью"
Exaggeration-Minimisation	"МИГРАЦИОННОЕ ЦУНАМИ УГРОЖАЕТ УТОПИТЬ АНГЛОГОВОРЯЩИЕ СТРАНЫ ЗАПАДА", "Так что штормить будет всех. А взрывы и стрельба — это только начало"
Repetition	—

Table 16: Few-shot Examples Used for SI Subtask 2

Technique	Example(s)
Name_Calling-Labeling	"Golobisti", "birokrati"
Guilt_by_Association	"Zamislite si, NSi problema z uporabo nacističnih in fašističnih simbolov nima, kajne, danes pa bi rušila ministra", "sicer ne posvetuje z ustreznimi strokovnjaki, ampak pogovor opravi z Jašo Jenulom, torej osebo, ki je bila v času vodenja janšistov večkrat kaznovana, ker je pozival k neprijavljenim protestom"
Doubt	"seveda s poslušnim delom", "Navsezadnje pa gre za odgovornost ministra tudi zato, ker je na vodilno mesto v policiji imenovan neprimeren kader"
Appeal_to_Hypocrisy	"Kljub navedenemu pa je pod ministrom in nekdanjim generalnim direktorjem v okviru CVZ celo napredoval", "On je namreč obljubljal eno, delal je popolnoma drugače"
Questioning_the_Reputation	"Gospod minister, nekdo laže, nekdo laže in vas spravlja v neroden položaj", "kadrovske načrte pa lahko vsakoletno praktično prilagodimo po lastnih preferencah"
Flag_Waving	"vlado, ki se s civilno družbo, recimo, ne pogovarja preko vodnih topov, pendrekov, solzivca in nasilja, ampak za mizo, civilizirano in strpno"
Appeal_to_Authority	"ni utemeljena na dejstvih", "številke so vas vzele, izdale"
Appeal_to_Popularity	"Tako čutijo ljudje, tako govorijo ankete", "v javnosti seveda odmeva"
Appeal_to_Values	"skupnim ukrepanjem", "transparentno in zakonito"
Appeal_to_Fear-Prejudice	"Potem so ti podatki odtekali morda tudi mafijskim kriminalnim združbam", "problematike kot so Romi in pa migracije"
Straw_Man	"Denimo NSi ustvarja vtis, kot da je migracijska situacija maltene katastrofična in da so migranti ogrožajoč element", "be stranki kot prednostno nalogo EU vidita v tem, da je treba čim bolj zabarikadirati zunanje meje Evropske unije in zavrniti čim več tistih, ki jim uspe priti na ozemlje trdnjave Evropa, ter jih čim prej vrniti tja, od koder so prišli"
Red_Herring	"kot Novomeščanka", "Za nekatere se je to leto res začelo srečno, veselo in zdravo, za rudarje v rudniku Velenje pač ne"
Whataboutism	"Istočasno pa nihče ne poskrbi za varovanje tožilke Gončin", "Tudi nakazilo Svetlane Makarovič ni bilo nezakonito, samo brez pravne podlage je bilo"
Appeal_to_Pity	"Hvaležna sem jim kot Novomeščanka", "Seveda bomo imeli minuto molka"
Causal_Oversimplification	"Vsak dan smo priča eni novi aferi", "za rešitev te problematike, torej povečanega števila, torej, problematike kot so Romi in pa migracije, poveča število policistov na terenu"
False_Dilemma-No_Choice	"se upokojijo ali pa si poiščejo boljše zaposlitev in podajo odpoved", "Niste učinkoviti in nikoli ne boste vedeli kako učinkoviti bi bili, če bi pred enim letom in pol sprejeli zakone, ki so jih napisali župani"
Consequential_Oversimplification	"Bo moral kdo umreti, da boste priznali resnost razmer kot minister in predložili sprejem akcijskega načrta, kjer bi začeli ta ozka grla odpravljati?", "Namesto ustreznega ukrepanja policija na podhodu železniške postaje v Ljubljani namesti nalepke z napisom: Če ste sami žrtev spolnega"
False_Equivalence	"Če primerjamo torej ceno mobilne hiške, ki jo lahko kupimo na tržišču v velikosti 32 m2, z vso opremo, torej kuhinjo, torej hladilniklerjem, sedežno"
Slogans	—
Conversation_Killer	—
Appeal_to_Time	"Situacijo je treba nemudoma začeti reševati", "Nič ne gre čez noč, tudi reševanje romske problematike ne"
Loaded_Language	"očitno namenoma", "medijski cirkus"
Obfuscation-Vagueness-Confusion	"prisilnih sredstev", "po mnenju poznavalcev"
Exaggeration-Minimisation	"enoto policije, od katere so odvisna življenja in varnost oseb, ki jih ogrožajo mafijske združbe", "spomin zlate ribice, da je nastopila vsesplošna in množična amnezija"
Repetition	—