

ABILITIES OF CONTRASTIVE SOFT PROMPTING FOR OPEN DOMAIN RHETORICAL QUESTION DETECTION

Josef BALOUN, Jiří MARTÍNEK

*Department of Computer Science and Engineering
University of West Bohemia
Plzeň, Czech Republic
e-mail: {balounj, jimar}@kiv.zcu.cz*

Cristophe CERISARA

*CNRS LORIA
Université de Lorraine
Nancy, France
e-mail: cerisara@loria.fr*

Pavel KRÁL

*Department of Computer Science and Engineering
University of West Bohemia
Plzeň, Czech Republic
e-mail: pkral@kiv.zcu.cz*

Abstract. In this work, we start by demonstrating experimentally that rhetorical question detection is still a challenging task, even for the state-of-the-art Large Language Models (LLMs). We propose an approach that boosts the performances of such LLMs by training a soft prompt in a way that enables building a joint embedding space from multiple loosely related corpora. The advantages of using a soft-prompt compared to fine-tuning is to limit the training costs and combat overfitting and forgetting. Soft prompting is often viewed as a way to guide the model towards a specific known task, or to introduce new knowledge into the model

through in-context learning. We show that soft prompting may also be used to modify the geometry of the embedding space, so that the distance between embeddings becomes semantically relevant for a target task, similarly to what is commonly achieved with contrastive fine-tuning. We exploit this property to combat data scarcity for the task of rhetorical question detection by merging several datasets into a joint semantic embedding space. On the standard Switchboard dataset we demonstrate that the resulting BERT-based model nearly divides by 2 the number of errors as compared to Flan-T5-XXL with only 5 few-shot labeled samples, thanks to this joint embedding space. We have chosen in our experiments a BERT model because it has already been shown with S-BERT that contrastive fine-tuning of BERT leads to semantically meaningful representations. Therefore, we also show that this property of BERT nicely transfers to the soft-prompting paradigm. Finally, we qualitatively analyze the resulting embedding space and propose a few heuristic criteria to select appropriate related tasks for inclusion into the pool of training datasets.

Keywords: Soft prompts, prompt-tuning, rhetorical questions, contrastive learning, triplet loss, pre-trained language models

1 INTRODUCTION

Nowadays, large pre-trained language models (PLM) are essential components of many natural language processing (NLP) applications. It has been demonstrated in many works that they encode valuable linguistic information, especially at the level of morphology, syntax, and semantics, as well as factual/world knowledge, such as cities, famous people, and historical events. It is also possible to extract procedural knowledge and multi-step reasoning from the largest PLMs (chain-of-thoughts, tree-of-thoughts, . . .), for instance, to solve maths problems.

However, there are still gaps in some types of information that these PLMs hardly capture, or maybe we just do not know yet how to extract it from the PLM. Examples of such challenges are detecting implicatures and rhetorical questions (RQs). For such tasks, the majority of popular PLMs with prompt-based approaches fail, and PLM performance is close to random guessing, when considering zero-shot or few-shot scenarios.

This can be seen for instance in [1], where the authors evaluate several PLMs with a prompt-based strategy for the task of conversational implicatures. The task is to automatically determine the implicature (i.e., the meaning yes or meaning no) for the question and its response (e.g., question: “Can you make a cake?”, response: “Can birds fly?”¹). Most models in their experiments obtained around 60 % accuracy on the test set (random guessing performance is 50 %) despite multiple prompt templates. Furthermore, the rhetorical question implicature (similar to the

¹ The response is a rhetorical question that semantically means the answer “yes”.

example above) seems to be especially difficult to detect compared to the other implicature types (e.g., “world knowledge”).

A natural option to try and solve these challenges would be to find datasets annotated with rhetorical questions and fine-tune a PLM on them. However, such datasets are relatively rare, small and heterogeneous, and fine-tuning large PLM incurs a prohibitive computational cost. We address the latter issue by exploiting soft-prompting [2, 3, 4] instead of fine-tuning, and show that contrastive training of the soft prompt enables to build a semantic embedding space adapted to the task, which to the best of our knowledge, has never been demonstrated before. We further show that this embeddings space may be obtained from several related datasets, which are thus projected into a joint embedding space. Finally, we solve the former issue by adapting this embedding space to the target dataset in a few-shot way, without exploiting any development corpus at this stage.

1.1 Research Question: Building a Joint and Semantic Embeddings Space

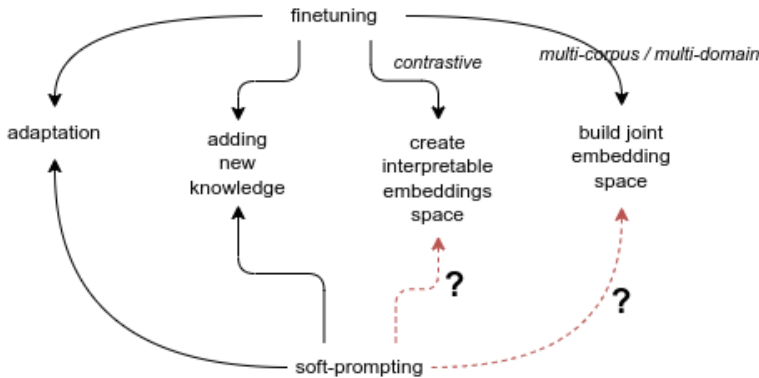


Figure 1. Illustration of the research questions addressed in this work: does contrastive training of soft prompts enable to build a semantic embedding space for rhetorical questions? Does multi-corpus training of soft-prompting enable to build a joint embedding space across domains?

Fine-tuning the parameters of a LLM may be used to achieve two main purposes: either adapting the model to a new domain or style, or adding new knowledge. It has been shown previously that both objectives may also be achieved with parameter-efficient training, and soft-prompting in particular. However, while fine-tuning an LLM contrastively may be used to produce a semantic embeddings space, it is still unclear whether the same can be achieved with soft prompting. The authors in [5] provide some hints that it might be possible, but with adapters and text-vision models. Nevertheless, such a conclusion is far from obvious given that soft prompting only affects the input sequence, which is a priori not designed to modify

the global properties of the embeddings space; if this is true, as our next experiments suggest it is, then this might mean that it may be possible to alter the embeddings space just by manually crafting carefully designed prompts.

Similarly, it is well-known that joint embedding spaces across diverse domains or corpora may be obtained through multi-corpus fine-tuning, but whether parameter efficient training enables to do the same is still an open question.

We propose to experimentally demonstrate, on the challenging task of rhetorical question detection, that both objectives may also be reached with properly training the soft-prompt of a LLM.

1.2 Contributions

1. **Progress in rhetorical question detection:** We gather several datasets related to rhetorical question detection and propose a parameter-efficient approach that outperforms by a large margin the state-of-the-art (SOTA) open-source and reproducible Flan-T5-XXL.
2. **Claim and support two novel properties of soft-prompting:** We show that contrastive soft-prompting may also be used to create semantic embeddings space; and to build a joint multi-corpus embedding space.
3. **New insights about desirable structure of a joint embedding space for rhetorical questions:** We qualitatively analyse the resulting semantic embedding space and exhibit some interesting structure, e.g., the fact that the sarcastic feature is not well aligned with the binary rhetorical feature.
4. **Few-shot without development corpus:** when fine-tuning a system with few-shots, hyper-parameters must still be adjusted, which is often done on a development corpus. The fact to use a development corpus questions whether the model is really few-shot. We describe a methodology that does not require any development corpus and thus guarantees that the resulting model is indeed a truly few-shot model.

2 CORPORA

Rhetorical questions (RQs) typically occur in two data sources: dialogues and social networks – 42 % of all questions on English Twitter are rhetorical [6]. In dialogues, though, rhetorical questions are among the least frequent categories of dialogue acts (see MRDA [7] and/or Switchboard [8] corpora). Therefore, methods working with little training data are required. Although there are some previous contributions for RQ detection in Twitter, e.g., [9], but the corpora change over time as tweets are deleted.

In this work, we focus on the automatic detection of RQs in dialogues. Indeed, detecting rhetorical questions is important for advanced conversational agents that target a more natural and smooth conversation. However, as summarized in the related works section, a variety of subtly different tasks exist around this notion of

rhetorical questions. Therefore, in this work we will study corpora which contain rhetorical question labels (e.g., Switchboard, Meeting Recorder Dialog Act Corpus – MRDA) as well as corpora dealing with “non-information seeking questions” (RQuet) and containing “sarcastic questions” (Sarcasm V2).

We included the following corpora in this work. All contain dialogues with RQs or related labels:

- The Switchboard Dialogue Act (**SWDA**) corpus is a set of manually transcribed conversations resulting in more than 220 000 utterances with 42 labels. We use the pre-processed balanced version of the SWDA corpus described in [10].
- The **MRDA** corpus contains over 100 000 hand-annotated dialogue act labels from dozens of naturally-occurring meetings. The corpus contains three levels of annotations; we use the RQ tags from the set of general labels;
- The **RQuet** corpus² [11] contains information seeking questions (ISQs) and non-information seeking questions (NISQs);
- The Sarcasm v2 (**SarcV2**) corpus³ [12] contains more than 6 500 sarcastic and non-sarcastic internet posts. We consider this additional corpus because the task might be related to RQs and because we want to study the possibility to transfer information from a related but distant task.

The statistics are provided in Appendix A.

3 RELATED WORK

3.1 Prompt-Based Zero-/Few-Shot Approaches

Fine-tuning pre-trained language models is still a method of choice to adapt PLMs to a new task or domain, but the cost required to update every single parameter of the model might be cumbersome, especially with large PLM with hundreds of billions of parameters. A possible solution is proposed by [13] with text prompts strategy. The authors show that text prompts are very effective with a frozen GPT-3 model, and obtain zero-shot predictions just by using a natural language sentence before the actual text input, such as: “*Translate the following English sentence to German: <sentence>*” → <german sentence>. In-context few-shot learning may be achieved by showing a few task demonstrations (e.g., several examples of translations). Either way, there are no gradient updates and no fine-tuning. We rely on the model’s ability to “understand” a prompt template.

However, such discrete (hard) prompts have several drawbacks. The main problem is that prompts which humans consider logical and suitable are not necessarily effective for the language models. Besides, we do not know whether our prompt is optimal in the context of the domain and the model we use. Last but not least, the

² <https://github.com/kkalouli/RQuet>

³ <https://nlds.soe.ucsc.edu/sarcasm2>

pre-trained models are sensitive to the choice of prompts [14]. The aforementioned problems have resulted in seeking “optimal” prompts in an automatic way relying on gradient-based searching [15] or automatic prompt generation [16].

Another set of methods is the usage of continuous (or soft) prompts [2, 17, 18]. The main idea behind this is adding additional learnable parameters (tunable “virtual text tokens”) that are injected into a PLM. Usually, these tunable tokens are somehow appended to the input text and optimized while keeping the whole model frozen. This training paradigm allows us to share large-scale models for various tasks more efficiently (only a small portion of optimized parameters are stored). Approaches based on prompt tuning have shown interesting and promising results compared to the manually created prompts [2]. The soft-prompting strategy belongs to the more general family of parameter-efficient training methods, which also include for instance low-rank adaptation, adapters and prefix tuning.

3.2 Rhetorical Question Detection

In dialogues, question types may be clustered into two groups: factoid (information-seeking) and non-factoid [11]. The NLP community often focuses on factoid questions (e.g., in the Question Answering task), as answering questions such as “Who was the president of the USA before Clinton?” is common in NLP applications, while the second group of questions plays a more subtle but still important role in the course of dialogues, e.g., to emphasize a statement or implicitly criticize a position. The most significant type of such questions is rhetorical questions, even though several sub-types exist [19].

Primary attention in the past has been paid to rhetorical questions [20, 21, 22, 23, 24], mainly focusing on linguistic aspects. Automatic detection of RQs in Switchboard corpus has been presented in [10]. The authors used a set of features combined with a SVM. They obtained good results on a balanced version of the SWDA corpus, with 81 % of accuracy without context, and 84 % when a context is used.

Authors in [11] created the corpus RQuet – Resource of Question Types, a collection of information-seeking questions (ISQs) and non-information-seeking-questions (NISQs). The authors evaluated several features, including speaker change, prior and subsequent contexts, and part-of-speech tags. They obtained a recognition accuracy of around 77 % when using questions without context.

Martínek et al. [25] focused on question type detection (including RQs) on the MRDA corpus. The authors used a linguistic rule-based approach together with a pre-trained BART [26] model to perform a zero-shot approach and found out that RQs are problematic for zero-shot.

Authors of [12, 27] have targeted the problem of rhetorical questions and sarcasm detection. As a source, they used debate forums, and labeling RQs was done automatically using heuristic rules.

In recent years, some efforts have been made to identify rhetorical questions in Twitter, notably [28, 29]. Zhuang et al. [9] created a corpus containing tweets

with questions (rhetorical and information-seeking) and carried out experiments with SVM and Bidirectional LSTM. The authors also integrated the prior tweet and a topic feature into the model and concluded that they seem to be helpful for this task.

Overall, the usage of Twitter corpora is problematic for reproducibility since a fraction of the tweets are regularly deleted over time, e.g., when a user deletes an account. We thus preferred to focus our work on dialogue corpora, such as MRDA and SWDA. Although RQs are less frequent in these corpora, long-term availability of data is guaranteed.

4 PRELIMINARY EXPERIMENTS

Before we delve into presenting our approach in detail, we describe the set of preliminary experiments demonstrating PLMs performance in the field of RQ detection. We focus on zero- or few-shot experiments with several open-source state-of-the-art PLMs. These experiments clearly show that all models dramatically fail at recognizing rhetorical questions, both in zero-shot (ZSL) and in-context few-shot (IC-FSL) setups. We explain this failure by the fact that recognizing rhetorical questions involves a reflexive, or meta-reasoning process, where the reader analyzes the semantic content of the question but also the intent of the speaker who is asking this question, intent that is different from the semantic content of the question.

Hence, the question “Do bicycles allow one to move from place to place?” might be interpreted by the PLM as a simple question with a clear semantic content, because the PLM’s training corpus contains many direct questions that look similar at first sight. But recognizing this question as a rhetorical one involves another level of reasoning out the context and the hidden intent of the speaker, who is likely not really interested in the factual answer to this question.

We show that such reflexive analyses are still out of reach of current open PLMs in Zero-Shot and IC-FSL modes, when 0 to 3 training samples per class are considered. We thus propose in Section 5 a solution that addresses the major flaws of the mainstream approaches commonly used to solve similar research questions, in particular fine-tuning and soft-prompt tuning. In these experiments, we test the following open state-of-the-art models:

- **Flan-T5-XXL** is a 11B-parameters model released by Google. It is considered as one of the best open PLM that outperforms GPT3 [30]. Because this model is an encoder-decoder transformer from the family of T5, we tested it with the following prompt templates and simply checked whether the answer is yes or no; note that it was always either “yes” or “no” in our experiments:
 - P1: “[input] Is the previous question a rhetorical question, yes or no?”
 - P2: “[input] Does the previous question require an answer, yes or no?”
 - P3: “Is the question a rhetorical question, yes or no? Question: [input]”

- **OPT-6.7B** is a 6.7B-parameters model released by Meta-AI [31]. It is one of the best open PLM from the GPT family. Since this model is a decoder-only transformer from the family of GPT, the question-answering strategy used for Flan-T5-XXL does not work, as other words than yes or no were often generated. We then used the P1 and P2 prompts and chose the minimum perplexity when either “yes” or “no” was concatenated to the prompt.
- **TK-11B** is a 11B-parameters model released by Allen-AI [32]. It is an open PLM trained with the “instruction-tuning” strategy that has led, along with reinforcement learning from human feedback (RLHF), to the success of ChatGPT. We followed the recipe given with TK-11B. We used the following prompt, and selected the lowest perplexity among “yes” and “no” continuation because all answers are not yes/no:
 - P4: “Definition: is the following question a rhetorical question, yes or no?
Input: [input] Output:”

In Table 1, we present the results of the zero-shot evaluation (see column 0), and the few-shot evaluation with 1 and 3 randomly chosen samples (columns 1 and 3). We have not tested more few-shot samples because we exceeded the maximum input length in some cases. Few-shot experiments are run five times, and the average accuracy is reported. **X** corresponds to experiments that have failed because of limited GPU VRAM (32 GB). The results show that the vast majority of the experiments resulted in an accuracy of around 50 % which is similar to the majority classifier or random guessing (note that the SWDA rhetorical question test corpus is balanced).

Model	Prompt	0	1	3
Flan-T5-XXL	P1	52.7	–	–
Flan-T5-XXL	P2	56.3	52.2	53.4
Flan-T5-XXL	P3	49.2	–	–
OPT-6.7B	P1	52.3	–	–
OPT-6.7B	P2	52.0	51.2	54.2
TK-11B	P4	55.9	57.7	X

Table 1. RQ accuracy on the balanced test set of SWDA, for an increasing number of training samples per class: 0 (ZSL), 1 and 3 (FSL)

5 PROPOSED APPROACH

Several parameter-efficient fine-tuning approaches, such as LOW-Rank-Adaptation and adapters, have been proposed recently to fine-tune large PLMs in a cost-efficient way. They often obtain comparable performances to full fine-tuning, even when the main model’s parameters are quantized to 4 bits as with qLoRA. We rather investigate the properties of the soft-prompting method because, conversely to the

above-mentioned mainstream parameter-efficient approaches, soft-prompting does not modify the internal parameters of the model, in particular the MLPs in the self-attention layers. The fact that they do not modify the MLP parameters, where information is mainly stored, leaves open the question of the real impact that the soft prompts may have on the embedding space that is build at the output of the transformer: is it possible to build a semantic embedding space with contrastive soft prompting, in a similar way than with S-BERT? Is it possible to build a joint embedding space that projects multiple domains into the same embedding space?

In this work, we investigate both properties of soft-prompt tuning in the context of the challenging task of rhetorical question detection. We thus propose a contrastive approach for soft-prompt tuning with the objective of building a semantic embedding space for rhetorical questions, and then exploit it for few-shot instance-based detection of rhetorical questions. We further investigate whether training this soft prompt contrastively on multiple domains/corpora enable building a joint embedding space across domains/corpora. For reproducibility, we provide a GIT repository⁴. We ran our experiments with relatively small uni- and bidirectional representatives of PLMs (namely smallest *gpt2* with 124M parameters and *bert-base-cased*), similarly as in the experiments of [17]. We expect that, by demonstrating both these properties of soft-prompting on small models, these properties shall remain valid on larger models, while the other way is more problematic because similar properties of PLMs are known to only emerge when the PLM is large enough, such as in-context learning, chain-of-thoughts...

5.1 Prompt-Tuning – Model Details

In all experiments, we use the BERT [33] pre-trained model (*bert-base-cased*). It was the most robust model in terms of hyper-parameters settings while obtaining good results with relatively low computational costs, compared to the other models we tested, such as GPT-2, Bloom, and Flan-T5 (see Appendix B).

The model parameters are frozen and we train only a final classification layer and 10 new embedding vectors (the soft prompts P_1 to P_{10}) based on preliminary experiments (see Appendix C). These trainable parameters are shown in green in Figure 2.

To preserve the CLS token pre-trained position, the soft prompts P_1 to P_{10} are systematically appended at the end of every input sequence due to the positional encoding. We use padding tokens and an adequate attention mask in order to guarantee the same position of the soft prompts for each input. The input text $T_1 \dots T_n$ is composed of the current question without context, to facilitate comparison across corpora. The last linear layer has one output neuron for binary classification.

⁴ <https://github.com/balounjosef/SoftPromptRQ>

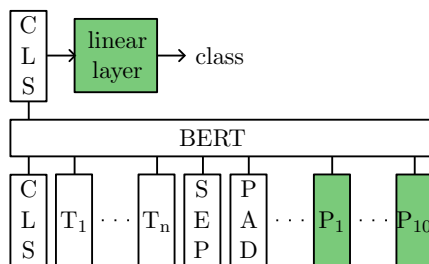


Figure 2. Soft-prompt model, trainable parameters are filled in green

5.2 Leveraging More Corpora

An important aspect of our approach is that we are using more than one corpus to take advantage of similar domains/tasks and also evaluate the generalization of the model.

5.2.1 Zero-Shot Cross-Corpus Transfer

The simplest way to test generalization across corpora consists in tuning the soft prompts on the training part of one corpus (called the *source* corpus) and evaluating the performance on the test part of another corpus (the *target* corpus). The experiments are presented in Section 6.1. This follows the intuition of transferring knowledge of similar labels in the same domain (e.g., RQs in dialogues). However, this simple approach does not solve the bias problem mentioned before. This is called Zero-shot Cross-corpus Transfer, because no label of the target corpus is used.

5.2.2 Pseudo-Few-Shot Cross-Corpus Transfer

After training the soft prompts on the source corpus, we may further continue training them on the target few-shot samples: we show in Section 6.2 that good results can be obtained. However, this approach is still limited by the fact that continued prompt-tuning involves setting some hyper-parameters, e.g. the number of epochs, typically using a development corpus of the target task. That is why we call this approach “pseudo-few-shot”.

5.3 Contrastive Alignment of Embeddings

The above-mentioned solutions are problematic because the hyper-parameters used to fine-tune the prompts on the target few-shot data strongly depends on the target corpus; assuming the best hyperparameters are found on the source corpus, or another distant dataset, they are likely to induce a large variance of performances when applied on the target few shots. This is why it is not rare to find published works

that use a development corpus of the target domain to find these hyper-parameters, which is not compatible with the few-shot setting [34].

We propose adopting a contrastive learning strategy to align both the source and target corpora into the same embedding space. Hence, few-shot classification on the target task may be realized within the framework of “instance-based learning” or “prototype-learning”, just by comparing the distances in this joint embedding space between the unknown test sample and the few-shot samples (see Section 6.3 and Figure 4).

Let s denote a sequence of tokens and $E : s \mapsto \mathbb{R}^d$ an embedding space. Prompt-tuning on one (source) corpus creates an embedding space E_{source} that is different from the target embedding space E_{target} . Then, continued prompt-tuning progressively aligns E_{source} towards E_{target} .

In our experiment (see Section 6.3), we try to align E_{MRDA} and E_{SWDA} into a single common embedding space. Contrastive learning is a method of choice to train such joint embeddings [35]. Furthermore, once such a joint embedding space is available, a parameter-free instance-based learning strategy can be used to perform few-shot learning by computing the distance between an unknown test sample and the known few-shots in this embedding space, and deciding on the output class with K-nearest neighbors. Figure 3 shows the idea behind of contrastive alignment of embeddings. Let us assume that a source and target corpora contains positive and negative classes (depicted as circles and crosses). The goal is to bring close together the positive samples in both target and source corpora and enable the use of classification based on measuring distances between vectors. Another option would be to train a discriminator model that decides, based on two input vectors across datasets, whether or not they share a class. Such an option, however, is more complicated and requires further training and more samples.

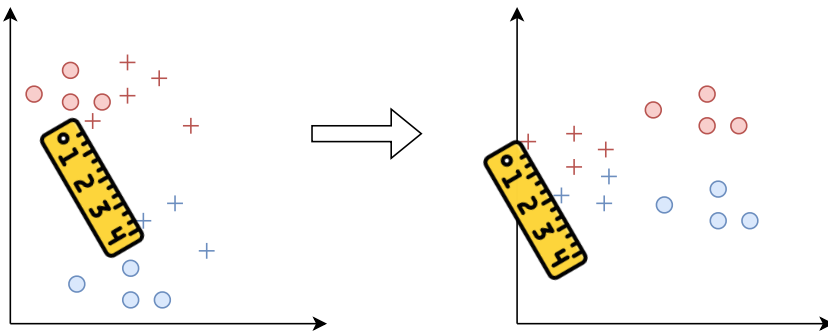


Figure 3. An illustration of simplified building the join embedding space for two corpora in 2D. Note that Euclidean distance is visualized, but in higher dimensions, the more reasonable option is to choose a cosine distance.

6 EXPERIMENTS

First of all, we provide a comparison of the number of samples used for a particular approach we have tested to solve our research question (see Table 2). An important difference is that the Pseudo-few-shot Cross-corpus Transfer, compared to Contrastive Prompt Tuning, requires additional validation data from the target domain.

Method	MRDA	RQuet	SWDA
Zero-Shot Learning (Section 4)	0	0	0
In-context Few-Shot Learning (Section 4)	0	0	3
Zero-shot Cross-corpus Transfer (Section 6.1)	542	1 588	0
Pseudo-few-shot Cross-corpus Transfer (Section 6.2)	542	1 588	shots + 142 (dev)
Contrastive Prompt Tuning (Section 6.3)	542	1 588	shots

Table 2. Number of samples used during tuning of each method. The *shots* relates to the number of samples used for few-shot (1, 5, 10 and 25-shots).

6.1 Zero-Shot Cross-Corpus Transfer

This experiment does not assume any annotated target labels, and directly transfer information from one corpus to another: concretely, prompt-tuning of the pre-trained language model is realized independently on each of the four corpora, SWDA, MRDA, RQuet and SarcasmV2, leading to four task-dependent models. Then, every model is directly evaluated on the test part of all corpora. The mapping between the respective four sets of binary labels is shown in Table 3. Cross-corpus accuracies of the BERT model are reported in Table 4.

class	SWDA	MRDA	RQuet	SarcV2
0	non-RQ	non-RQ	ISQ	not sarcastic
1	RQ	RQ	NISQ	sarcastic

Table 3. Mapping between labels in the four corpora

Obviously, the best results are obtained when the train and test sets belong to the same corpus (see the diagonal in Table 4).

We can further conclude from these results that:

1. The soft-prompt BERT model obtained competitive results with (SOTA);
2. The transfer from MRDA and RQuet corpora to SWDA seems possible (see values in italics in Table 4);
3. Conversely, the transfer from SWDA to MRDA fails as the results are close to random guessing, which may be due to a bias in a corpus;
4. sarcastic and rhetorical questions do not seem to be well aligned.

Train	Test			
	SWDA	MRDA	RQuet	SarcV2
SWDA	80.9	53.9	61.7	48.6
MRDA	72.7	68.6	55.0	47.5
RQuet	67.2	42.2	75.0	51.2
SarcV2	42.2	43.1	51.7	74.4
SOTA	81.3	—	77.7	—

Table 4. Single-corpus (diagonal) and cross-corpus accuracy [%] of soft-prompt model (*bert-base-cased*). SOTA results: SWDA [10], RQuet [11]. For SarcV2, the authors reports the F1 score 76.0 % [27].

We focus next on transferring to the SWDA corpus, thereafter called the *target* corpus, while MRDA, RQuet and SarcV2 are the *source* corpora. Since the SWDA is the largest rhetorical question corpus in the collection, the results are more reliable.

6.2 Pseudo-Few-Shot Cross-Corpus Transfer

The cost of producing a few annotated examples for the target task is often moderate, it is reasonable to assume that a small number of labeled samples per target class is available. Exploiting them in our model may be done with *continued prompt-tuning* on the few-shot samples, i.e., simply continuing standard prompt-tuning iterations but on a new set of training samples⁵. Hence, in this experiment, we randomly sample 1, 5, 10 and 25 samples per class from the SWDA corpus. As this sampling may induce variable results, we run this experiment ten times and average the results. We report in Table 5 the results of the following models:

- BERT prompt-tuned from scratch (directly prompt-tuned on the target SWDA few-shots). In this case, there is no transfer and the green trainable parameters in Figure 2 are randomly initialized (the first column with the header None in Table 5);
- BERT prompt-tuned on one of three source corpora (MRDA, RQuet, SarcV2) or their combination (MRDA+RQuet and All) and further continued-prompt-tuned on the target SWDA few-shots. The best results were obtained by first prompt-tuning on the MRDA train corpus and then continued-prompt-tuned on the target SWDA 25 shots.

The learning rate for prompt-tuning is $2e-4$; it has been chosen with preliminary experiments on the MRDA corpus only. However, note that although the model parameters are only trained on the few-shot samples, we also perform early stopping on the development corpus of SWDA; so the experiment should not be really considered as few shot-learning, which explains the term “pseudo-few-shot” in the

⁵ To the best of our knowledge, we introduce here the term “continued prompt-tuning”, which is inspired by the more common term “continued pretraining”.

Shots	Pseudo-Few-Shot Cross-Corpus Transfer					
	None	MRDA	RQuet	SarcV2	MRDA + RQuet	All
1	58.8	73.0	68.2	43.2	72.7	65.6
5	63.3	73.9	71.2	48.6	73.8	67.6
10	64.8	75.1	72.5	50.5	73.8	67.3
25	66.6	76.4	74.1	51.8	73.9	68.9

Table 5. 10 runs average accuracies [%] on SWDA corpus: we gradually performed continued-prompt-tuning of three models from the previous experiment (Section 6.1 and Table 4)

section’s title. This is a generic issue with few-shot learning experiments, which is usually “solved” with one of the two following approaches: either use a development corpus in addition to the few-shots (as we did), or further split the few-shots into a train and development corpus. As discussed in [34], none of these options is satisfying: the first one because more than a few data points are used to train the model; and the second one because setting hyper-parameters with a standard tuning procedure on so little data inevitably induces a significant variance of the error, and in most concrete use cases, the few-shots are given, and there is no way to tell whether performances will be far from the average or not. We propose a real few-shot learning approach in Section 6.3.

Continued prompt-tuning from both the MRDA and RQuet models on the target few-shots always improves results. However, initial prompt-tuning of BERT on the combined MRDA and RQuet corpora does not help, as compared to simply prompt-tuning BERT on MRDA. Similarly, initial prompt-tuning of BERT on SarcV2 is useless, as the accuracy stays close to a random guess.

This experiment indicates that, just like fine-tuning, prompt-tuning may also be successfully used to transfer information from related tasks when applied in the few-shot setting, and that this approach may be useful for the rhetorical question detection task, which is difficult to handle for large pre-trained models as discussed in Section 4.

6.3 Contrastive Prompt Tuning – Alignment of Embeddings

In this scenario, we replace the final classification layer in Figure 2 by a projection layer that reduces the 768-dimensional BERT output to a 32-dim embedding. This dimensionality reduction is useful due to the “curse of dimensionality” challenge, which makes similarity search in high dimensions difficult. The final embedding size of 32 is based on a few trials made on the source corpora. The trainable parameters (10 prompts and final projection layer) are now trained with the triplet loss [36]:

$$\mathcal{L}(A, P, N) = \max[d(f_A, f_P) - d(f_A, f_N) + \alpha, 0], \quad (1)$$

where f is an embedding obtained by our model, d is a distance function in the embedding space and $\alpha = 0.5$ is the margin. The goal of the margin is to reduce

overfitting, stabilize a training process and, last but not least, helps the model to generalize better. We experimented on MRDA corpus with different distance functions (Euclidean, Manhattan, Cosine). Similar results were obtained, probably because of layer-normalization in BERT, as the norms of the projected embedding vectors are more or less similar. Finally, we chose the cosine distance.

We create triplets that consist of three distinct samples: *anchor* (A), *positive* (P), and *negative* (N). An anchor and positive sample share the same class (label), while the negative sample is representative of a different class. Optimizing the parameters of the network brings A and P closer in the feature space and pushes A and N farther apart [37].

Figure 4 summarizes the proposed approach. As before, BERT is first prompt-tuned on one or several source corpora, but now with the triplet loss. Then, in the second transfer step, BERT is further prompt-tuned with triplets composed of anchors that come from either a source corpus or the target few-shots, while positive and negative samples are sampled from the target few-shots. Hence, many of the second-step triplets mix samples from the source and target corpora, which consequently creates a joint embedding space.

Another advantage of the proposed contrastive approach is that there is no hyper-parameter (learning rate and early stopping) to tune on a target development corpus: contrastive training is run until convergence⁶ and then few-shot classification is achieved with the nearest neighbors method. Thus, this approach achieves “real few-shot classification”, as described in [34].

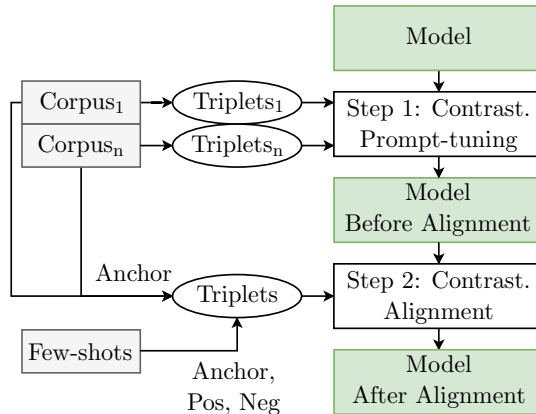


Figure 4. Outline of the contrastive training strategy

The main objective is to build a joint embedding space for rhetorical questions that is common to all related corpora. First, we discard SarcV2, which is not related

⁶ We use randomly generated triplets from training data and measure the ratio of correctly predicted triplets so that positive distance is lower than negative distance: $d(f_A, f_P) < d(f_A, f_N)$.

enough to the task, as discussed in Section 7. We thus consider both MRDA and RQuet as source corpora and SWDA as the target corpus.

Table 6 compares the accuracy obtained with few-shot K-NN classification, before and after the second step of alignment (the last two columns). The remaining columns on the left are results from Table 5. In all configurations, we set the K to the closest odd number $\leq \sqrt{N}$ where N is the number of shots available, since it is standard practice to deal with outliers. So for the 1, 5, 10, and 25 shots, the K is set to 1, 1, 3, and 5 respectively, but we have also observed that $K = 1$ gives similar results. We can draw several conclusions from Table 6:

1. Building a better embedding space with interpretable distances comes at the cost of lower accuracy, as compared to the best pseudo-few-shot learning method that directly optimizes the accuracy. This is particularly visible in the first-step models. Note that for these two contrastive models, no SWDA development data is used;
2. Aligning the embedding space in step 2 clearly improves classification over the unaligned embedding space in step 1;
3. Aligning the embeddings in step 2 requires enough few-shots. The results are poor with only 1-shot, but they tend to get closer to the best pseudo-few-shot model with more shots.

Shots	Pseudo-Few-Shot Cross-Corpus Transfer		Before Alignment	After Alignment
	MRDA + RQuet	MRDA + RQuet	MRDA + RQuet	MRDA + RQuet
	Dev set ✓	Dev set ✗	Dev set ✗	Dev set ✗
1	72.7	71.1	60.9	68.7
5	73.8	73.3	67.9	73.6
10	73.8	73.6	69.1	74.8
25	73.9	73.6	68.7	75.1

Table 6. Overall summary; 10 runs average accuracies [%] on SWDA. The first column is the continued-prompt-tuning of pre-trained BERT prompt-tuned on the specified data (Section 6.2); the next column is the contrastive model with independent triplets per corpus; the last column is the contrastive model after the alignment step with mixed triplets across corpora (both covered in Section 6.3). Note that Dev set ✓ and Dev set ✗ indicates whether or not the method requires target development set (as already indicated in Table 2).

7 DISCUSSION

Although the best absolute performance is obtained with the pseudo-few-shot model, the comparison with the joint embeddings model is not totally fair, as the latter

does not use any development corpus from the target task. Furthermore, building a joint embedding space in a contrastive way is more interesting as the distances between samples are interpretable, and information coming from the three corpora is encoded into this unique common space, which might be in the future enriched with new related tasks and corpora.

Despite our initial intuition, that sarcastic questions may often also consist of rhetorical questions, our attempts at merging both tasks failed experimentally. We analyze next qualitatively the relations between the embeddings space of each corpus independently obtained after the initial prompt-tuning phase.

7.1 How to Interpret Visualizations?

Figures 5 and 6 show the visualizations of embedding spaces. In this way, we try to plot the relations between corpora and the quality of our “transformation” provided by the contrastive training strategy. Each of the figures contains two or three circles. Each circle represents a particular corpus. The points within represents either a positive or negative class sample. The samples are classified based on cosine distance therefore only the angle matters. The origin is in the center of the circle.

The most important is the outer circle since it is related to our target corpus (SWDA). If all corpora (or their embedding spaces more precisely) are well-aligned it can be seen in the image that all the red marks will occur in one area of the circle (within spherical coordinates) and all the blue ones will be in the opposite area across all the circles.

For the visualization of corpus alignments given a specific model, we take the same number of class 0 (red) and class 1 (blue) samples for each corpus. Since we use cosine distance for the training and evaluation phase, we do not care about the norms of the embedding vectors and normalize them. Then, we reduce the dimensionality to 2 utilizing PCA. Since we are interested in the angles, we further re-normalize the vectors so that each concentric circle represents a particular corpus. The classes are distinguished in the same way and also by a different color. For visualization purposes, we further add a small portion of noise to visualize the density. In this way, the alignment between corpora should be visible since the red and blue colored symbols creates “clusters” across circles.

7.2 Analysis

In Figure 5, the contrast between left and right images is striking: in the left image, the red (non-RQ and ISQ) samples of both RQuet and SWDA are mainly located on the left, while the blue (RQ and NISQ) samples of both corpora are mainly located on the right. So, both corpora look pretty well aligned.

Conversely, the right part of Figure 5, the red (non-RQ and not sarcastic) samples of SWDA and SarcV2 seem opposed, while the blue (RQ and sarcastic) samples seem more aligned on the left; even worse, the two main PCA-dimensions do not

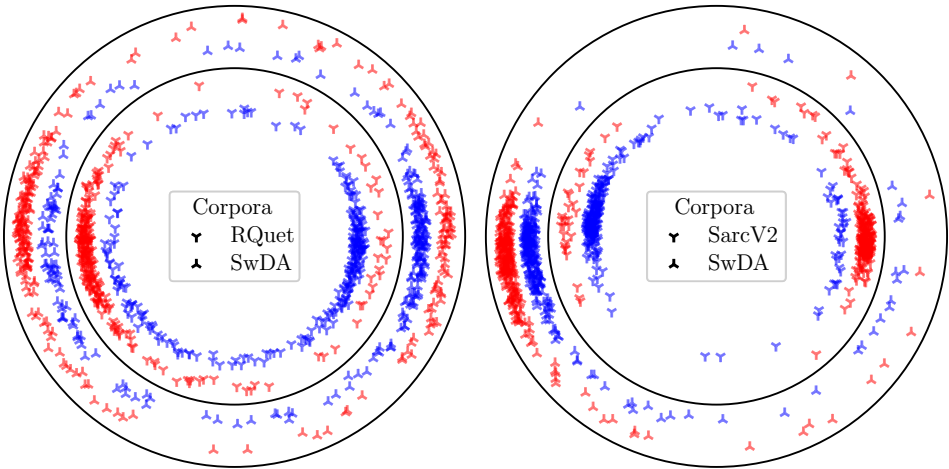


Figure 5. Vizualization of the model trained on RQuet (left) and SarcV2 (right) before alignment

enable to distinguish between RQ and non-RQ from SWDA. So the respective embeddings spaces of SWDA and SarcV2 are not aligned at all, which might explain the poor performances obtained when trying to transfer information from SarcV2 to SWDA.

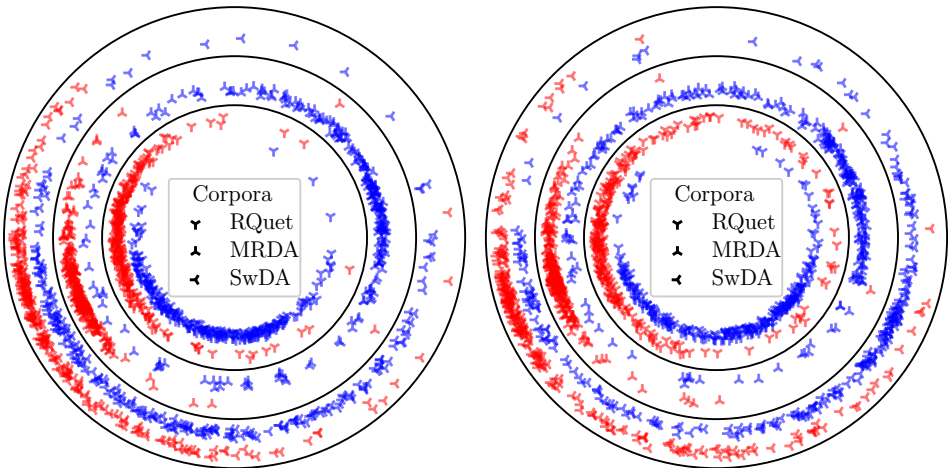


Figure 6. Vizualization of the model trained on MRDA + RQuet before (left) and after (right) alignment with 25-shots from SWDA

Note that these visualizations might be used to answer the difficult question of deciding whether a new task is close enough to the target rhetorical question detection task to be considered and merged within the common embedding space built by the three corpora: MRDA, RQuet, and SWDA.

We can further see in Figure 6 the impact of aligning the embeddings space contrastively: if we compare the right image to the left one, the blue dots seem better aligned across the three corpora, which confirms our intuition about the positive effect of aligning the models with cross-corpus triplets.

8 CONCLUSION

In this work, we have presented two new properties of soft-prompt tuning that appear even for relatively small-sized pretrained language models: first, the fact that, just like with full fine-tuning, using a contrastive training loss builds a semantic embedding space, which is not intuitive given that no internal parameter of the model is modified; mere soft prompts thus seem to be able to more strongly alter the resulting embedding space than expected. The second property is that, also just like with full fine-tuning, soft-prompts may build a joint embedding space where multiple corpora and labels can be aligned semantically.

We addressed a task, rhetorical question detection, which is difficult to achieve in a zero-shot setting even with large pre-trained language models, as we observed both in the literature on this topic and with our own preliminary experiments (Section 4). An additional challenge comes from the fact that when fine-tuning or prompt-tuning such models on one or several of the few existing dialogue corpora annotated with rhetorical questions, the relatively good performances obtained in the training conditions hardly transfer to another corpus with a different application context.

The proposed contrastive soft-prompting strategy (see Figure 3) enables to partly solve both challenges by transferring information from several corpora into the same joint semantic embedding space, using only few shots from the target task. When evaluated in a *real few-shot* setup, our approach obtained better results than 5 out of 6 prompt-tuning models that exploit a much larger development corpus for hyperparameter tuning and proved competitive against the best prompt-tuning model.

Finally, we visualize and analyze the resulting embedding spaces, and suggest a qualitative criterion to select appropriate related corpora that may be further included in this common embedding space.

Acknowledgments

The work of Josef Baloun has been supported by the Grant No. SGS-2025-022 – New Data Processing Methods in Current Areas of Computer Science. The work of Jiří Martínek and Pavel Král has been supported by the project R & D of Technologies for Advanced Digitalization in the Pilsen Metropolitan Area (DigiTech)

No. CZ.02.01.01/00/23_021/0008436. The work of Christophe Cerisara has been supported by the project ANR LLM4ALL. Computational resources were partly provided by GENCI-IDRIS (Grant 2023-AD011011668).

REFERENCES

- [1] RUIS, L.—KHAN, A.—BIDERMANN, S.—HOOKER, S.—ROCKTÄSCHEL, T.—GREFENSTETTE, E.: Large Language Models Are Not Zero-Shot Communicators. CoRR, 2022, doi: 10.48550/arXiv.2210.14986.
- [2] LESTER, B.—AL-RFOU, R.—CONSTANT, N.: The Power of Scale for Parameter-Efficient Prompt Tuning. CoRR, 2021, doi: 10.48550/arXiv.2104.08691.
- [3] HE, Y.—ZHENG, S.—TAY, Y.—GUPTA, J.—DU, Y.—ARIBANDI, V.—ZHAO, Z.—LI, Y.—CHEN, Z.—METZLER, D.—CHENG, H. T.—CHI, E. H.: HyperPrompt: Prompt-Based Task-Conditioning of Transformers. In: Chaudhuri, K., Jegelka, S., Song, L., Szepesvari, C., Niu, G., Sabato, S. (Eds.): Proceedings of the 39th International Conference on Machine Learning. PMLR, Vol. 162, 2022, pp. 8678–8690, <https://proceedings.mlr.press/v162/he22f.html>.
- [4] VU, T.—LESTER, B.—CONSTANT, N.—AL-RFOU, R.—CER, D.: SPoT: Better Frozen Model Adaptation Through Soft Prompt Transfer. CoRR, 2021, doi: 10.48550/arXiv.2110.07904.
- [5] KHAN, Z.—FU, Y.: Contrastive Alignment of Vision to Language Through Parameter-Efficient Transfer Learning. The Eleventh International Conference on Learning Representations (ICLR 2023), 2023, <https://openreview.net/forum?id=x0BPR9iXc1>.
- [6] PAUL, S. A.—HONG, L.—CHI, E. H.: What Is a Question? Crowdsourcing Tweet Categorization. Workshop on Crowdsourcing and Human Computation at the Conference on Human Factors in Computing Systems (CHI 2011), 2011, <https://www.humancomputation.com/crowdcamp/chi2011/papers/paul.pdf>.
- [7] SHRIBERG, E.—DHILLON, R.—BHAGAT, S.—ANG, J.—CARVEY, H.: The ICSI Meeting Recorder Dialog Act (MRDA) Corpus. Technical Report. 2004, <https://aclanthology.org/W04-2319/>.
- [8] GODFREY, J. J.—HOLLIMAN, E. C.—MCDANIEL, J.: SWITCHBOARD: Telephone Speech Corpus for Research and Development. ICASSP-92: 1992 IEEE International Conference on Acoustics, Speech and Signal Processing, Vol. 1, 1992, pp. 517–520, doi: 10.1109/ICASSP.1992.225858.
- [9] ZHUANG, Y.—RILOFF, E.: Exploring the Role of Context to Distinguish Rhetorical and Information-Seeking Questions. In: Rijhwani, S., Liu, J., Wang, Y., Dror, R. (Eds.): Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop (ACL 2020). 2020, pp. 306–312, doi: 10.18653/v1/2020.acl-srw.41.
- [10] BHATTASALI, S.—CYTRYN, J.—FELDMAN, E.—PARK, J.: Automatic Identification of Rhetorical Questions. Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference

- on Natural Language Processing (Volume 2: Short Papers), 2015, pp. 743–749, <https://aclanthology.org/P15-2122.pdf>.
- [11] KALOULI, A.L.—KEHLBECK, R.—SEVASTJANOVA, R.—DEUSSEN, O.—KEIM, D.A.—BUTT, M.: Is That Really a Question?: Going Beyond Factoid Questions in NLP. In: Zarrieß, S., Bos, J., van Noord, R., Abzianidze, L. (Eds.): Proceedings of the 14th International Conference on Computational Semantics (IWCS 2021). Association for Computational Linguistics, 2021, pp. 132–143, <https://aclanthology.org/2021.iwcs-1.13/>.
 - [12] ORABY, S.—HARRISON, V.—REED, L.—HERNANDEZ, E.—RILOFF, E.—WALKER, M.: Creating and Characterizing a Diverse Corpus of Sarcasm in Dialogue. 2016, pp. 31–41, doi: 10.18653/v1/W16-3604.
 - [13] BROWN, T.—MANN, B.—RYDER, N.—SUBBIAH, M.—KAPLAN, J.D. et al.: Language Models Are Few-Shot Learners. In: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M.F., Lin, H. (Eds.): Advances in Neural Information Processing Systems 33 (NeurIPS 2020). Curran Associates, Inc., 2020, pp. 1877–1901, https://proceedings.neurips.cc/paper_files/paper/2020/file/1457c0d6bfc4967418bfb8ac142f64a-Paper.pdf.
 - [14] ZHAO, Z.—WALLACE, E.—FENG, S.—KLEIN, D.—SINGH, S.: Calibrate Before Use: Improving Few-Shot Performance of Language Models. In: Meila, M., Zhang, T. (Eds.): Proceedings of the 38th International Conference on Machine Learning. PMLR, Vol. 139, 2021, pp. 12697–12706, <https://proceedings.mlr.press/v139/zhao21c.html>.
 - [15] SHIN, T.—RAZEGHI, Y.—LOGAN IV, R.L.—WALLACE, E.—SINGH, S.: Auto-Prompt: Eliciting Knowledge from Language Models with Automatically Generated Prompts. In: Webber, B., Cohn, T., He, Y., Liu, Y. (Eds.): Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP). Association for Computational Linguistics, 2020, pp. 4222–4235, doi: 10.18653/v1/2020.emnlp-main.346.
 - [16] GAO, T.—FISCH, A.—CHEN, D.: Making Pre-Trained Language Models Better Few-Shot Learners. 2021, pp. 3816–3830, doi: 10.18653/v1/2021.acl-long.295.
 - [17] LIU, X.—ZHENG, Y.—DU, Z.—DING, M.—QIAN, Y.—YANG, Z.—TANG, J.: GPT Understands, Too. CoRR, 2021, doi: 10.48550/arXiv.2103.10385.
 - [18] LI, X.L.—LIANG, P.: Prefix-Tuning: Optimizing Continuous Prompts for Generation. 2021, pp. 4582–4597, doi: 10.18653/v1/2021.acl-long.355.
 - [19] DAYAL, V.: Questions. Oxford University Press, 2016.
 - [20] KIMBALL, J.P.—ADAMS, D.—CAMPBELL, M.A.—COHEN, V.—LOVINS, J.—MAXWELL, E.—NYGREN, C.—REIGHARD, J.: Papers from the 7th Regional Meeting of the Chicago Linguistic Society. University of Chicago, Chicago Linguistic Society, 1971.
 - [21] SCHMIDT-RADEFELDT, J.: On So-Called ‘Rhetorical’ Questions. Journal of Pragmatics, Vol. 1, 1977, No. 4, pp. 375–392, doi: 10.1016/0378-2166(77)90029-7.
 - [22] FRANK, J.: You Call That a Rhetorical Question?: Forms and Functions of Rhetorical Questions in Conversation. Journal of Pragmatics, Vol. 14, 1990, No. 5, pp. 723–738, doi: 10.1016/0378-2166(90)90003-V.

- [23] GUTIÉRREZ REXACH, J.: Rhetorical Questions, Relevance and Scales. *Alicante Journal of English Studies*, 1998, No. 11, pp. 139–155, doi: 10.14198/raei.1998.11.11.
- [24] SCHAFFER, D.: Can Rhetorical Questions Function as Retorts?: Is the Pope Catholic? *Journal of Pragmatics*, Vol. 37, 2005, No. 4, pp. 433–460, doi: 10.1016/j.pragma.2003.12.007.
- [25] MARTÍNEK, J.—CERISARA, C.—KRÁL, P.—LENC, L.—BALOUN, J.: Weak Supervision for Question Type Detection with Large Language Models. *Proceedings of Interspeech 2022*, 2022, pp. 3283–3287, doi: 10.21437/Interspeech.2022-345.
- [26] LEWIS, M.—LIU, Y.—GOYAL, N.—GHAZVININEJAD, M.—MOHAMED, A.—LEVY, O.—STOYANOV, V.—ZETTEMAYER, L.: BART: Denoising Sequence-to-Sequence Pre-Training for Natural Language Generation, Translation, and Comprehension. In: Jurafsky, D., Chai, J., Schluter, N., Tetreault, J. (Eds.): *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, 2020, pp. 7871–7880, doi: 10.18653/v1/2020.acl-main.703.
- [27] ORABY, S.—HARRISON, V.—MISRA, A.—RILOFF, E.—WALKER, M.: Are You Serious?: Rhetorical Questions and Sarcasm in Social Media Dialog. In: Jokinen, K., Stede, M., DeVault, D., Annie, L. (Eds.): *Proceedings of the 18th Annual SIGdial Meeting on Discourse and Dialogue*. Association for Computational Linguistics, 2017, pp. 310–319, doi: 10.18653/v1/W17-5537.
- [28] RANGANATH, S.—HU, X.—TANG, J.—WANG, S.—LIU, H.: Identifying Rhetorical Questions in Social Media. *Proceedings of the International AAAI Conference on Web and Social Media*, Vol. 10, 2016, pp. 667–670, doi: 10.1609/icwsm.v10i1.14771.
- [29] ZHAO, Z.—MEI, Q.: Questions About Questions: An Empirical Analysis of Information Needs on Twitter. *Proceedings of the 22nd International Conference on World Wide Web (WWW '13)*, 2013, pp. 1545–1556, doi: 10.1145/2488388.248852.
- [30] CHUNG, H. W.—HOU, L.—LONGPRE, S.—ZOPH, B.—TAY, Y.—FEDUS, W. et al.: Scaling Instruction-Finetuned Language Models. *CoRR*, 2022, doi: 10.48550/arXiv.2210.11416.
- [31] ZHANG, S.—ROLLER, S.—GOYAL, N.—ARTETXE, M.—CHEN, M.—CHEN, S. et al.: OPT: Open Pre-Trained Transformer Language Models. *CoRR*, 2022, doi: 10.48550/arXiv.2205.01068.
- [32] WANG, Y.—MISHRA, S.—ALIPOORMOLABASHI, P.—KORDI, Y.—MIRZAEI, A.—NAIK, A. et al.: Super-NaturalInstructions: Generalization via Declarative Instructions on 1600+ NLP Tasks. In: Goldberg, Y., Kozareva, Z., Zhang, Y. (Eds.): *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, 2022, pp. 5085–5109, doi: 10.18653/v1/2022.emnlp-main.340.
- [33] DEVLIN, J.—CHANG, M. W.—LEE, K.—TOUTANOVA, K.: BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. 2019, pp. 4171–4186, doi: 10.18653/v1/N19-1423.
- [34] SCHICK, T.—SCHÜTZE, H.: True Few-Shot Learning with Prompts – A Real-World Perspective. *Transactions of the Association for Computational Linguistics*, Vol. 10, 2022, pp. 716–731, doi: 10.1162/tacl.a.00485.

- [35] WANG, T.—ISOLA, P.: Understanding Contrastive Representation Learning Through Alignment and Uniformity on the Hypersphere. Proceedings of the 37th International Conference on Machine Learning (ICML '20), PMLR, Vol. 119, 2020, pp. 9929–9939, <https://proceedings.mlr.press/v119/wang20k.html>.
- [36] HOFFER, E.—AILON, N.: Deep Metric Learning Using Triplet Network. In: Fergan, A., Pelillo, M., Loog, M. (Eds.): Similarity-Based Pattern Recognition (SIMBAD 2015). Springer, Cham, Lecture Notes in Computer Science, Vol. 9370, 2015, pp. 84–92, doi: 10.1007/978-3-319-24261-3_7.
- [37] BALNTAS, V.—RIBA, E.—PONSA, D.—MIKOLAJCZYK, K.: Learning Local Feature Descriptors with Triplets and Shallow Convolutional Neural Networks. In: Wilson, R. C., Hancock, E. R., Smith, W. A. P. (Eds.): Proceedings of the British Machine Vision Conference (BMVC). BMVA Press, 2016, doi: 10.5244/C.30.119.

Appendix A CORPORA STATISTICS

The number of samples and the input size in terms of input tokens for *bert-base-cased* model are provided in Table A1 and Table A2, respectively.

Corpus	train	dev	test
SWDA	560	142	256
MRDA	476	66	102
RQuet	1 270	318	180
SarcV2	5 216	652	652

Table A1. Number of samples in train, development and test part of each balanced corpus

Corpus	Max	Mean	Median	90-perc.
SWDA	46	13.36	12	22
MRDA	65	10.69	9	20
RQuet	117	20.27	15	39
SarcV2	411	60.85	48	124

Table A2. The number of input tokens in samples from given corpus with *bert-base-cased* tokenizer (maximal, mean, median and 90 percentile value)

Appendix B OTHER MODEL RESULTS

The results of different prompt-tuning models on the full SWDA corpus are provided in Table B1.

The zero-shot cross-corpus results of gpt2 soft-prompt model are provided in Table B2. Compared to Table 4 from Section 6.1 where the *bert-base-cased* was employed, the results are generally worse except SarcV2. We also noticed better results when transferring from SWDA to MRDA and RQuet which indicates that

Model	Prompts	Learnable	SWDA
		Parameters	Acc
bert-base-cased	10	8 449	80.9
bert-base-cased	50	39 169	81.2
bert-large-cased	10	11 265	81.2
gpt2	10	8 449	78.9
google/flan-t5-large	10	11 265	75.4
google/flan-t5-large	50	52 225	80.1
bigscience/bloom-7.1b	10	45 057	80.5

Table B1. Results of different prompt-tuning models on full SWDA corpus

the transferability across corpora may also depend on the model. In that case, the proposed qualitative criterion is beneficial. It may help to decide whether to include the corpus given the concrete model or not.

Train	Test			
	SWDA	MRDA	RQuet	SarcV2
SWDA	78.9	60.8	64.4	46.6
MRDA	67.6	66.7	48.9	41.7
RQuet	67.2	48.0	71.1	50.6
SarcV2	45.3	45.1	53.3	78.1

Table B2. Single-corpus (diagonal) and cross-corpus accuracy [%] of gpt2 soft-prompt model

Appendix C NUMBER OF PROMPTS

The preliminary experiments to decide the number of prompts are provided in Table C1 and Table C2.

Prompts	RQuet	MRDA	SWDA	AVG
1	72.2	68.6	76.2	72.3
5	75.6	67.6	75.8	73.0
10	75	68.6	80.9	74.8
20	73.9	66.7	80.5	73.7
50	69.4	66.7	81.2	72.4

Table C1. Results of prompt-tuning *bert-base-cased* model with different number of prompts on full corpus

Prompts	SWDA + left context
1	77.8
10	81.1
50	73.9

Table C2. Results of prompt-tuning gpt2 model with different number of prompts on full SWDA corpus utilizing the left context



Josef BALOUN is a Ph.D. student at the Faculty of Applied Sciences, University of West Bohemia, in Pilsen, Czech Republic. His research focuses on deep learning, computer vision, and natural language processing, with an emphasis on document analysis. He is particularly interested in addressing challenges related to limited data, investigating the generalization capabilities of neural networks, and enhancing the quality of neural network representations. His work aims to optimize neural networks for effective performance in data-scarce environments while ensuring robust generalization in practical applications.



Jiří MARTÍNEK is a postdoctoral researcher at the Faculty of Applied Sciences, University of West Bohemia, in Pilsen, Czech Republic. He obtained his Ph.D. in 2022, with a research focus primarily on deep learning, particularly in the fields of natural language processing and image processing. He has contributed to more than a dozen publications, presented his work at prestigious international conferences, and co-authored articles in high-impact journals.



Cristophe CERISARA is a French researcher at CNRS (National Centre for Scientific Research), specialized in large language models and natural language processing. He created and has been leading the SYNALP research team composed of about 20 NLP researchers since 2012. He is also the leader of the AI-NLP axis of the LORIA laboratory since 2019, and he was referent for the French National Plan in AI in 2020. He has supervised 13 Ph.D. thesis, and has lead several projects in AI and about training Large Language Models in the past few years.



Pavel KRÁL is Associated Professor at the Department of Computer Science and Engineering of the University of West Bohemia (Czech Republic). He is also a head of the NLP research group (<http://nlp.kiv.zcu.cz>) at the NTIS research center of the University of West Bohemia. He received his Ph.D. degree in 2007 at the Department of Computer Science and Engineering at the University of West Bohemia and at Henri Poincaré University in Nancy (France). His research is focused on natural language processing and pattern recognition with a particular focus on semantics, dialogue act recognition, historical document

analysis. He published more than 100 research papers in prestigious international journals and conferences. He is also responsible for the specialization Natural Language Processing in the graduate (M.Sc.) study program – Computer Science and its Specialization.