

Responsible NLP Checklist

Paper title: *This is not a Disimprovement: Improving Negation Reasoning in Large Language Models via Prompt Engineering*

Authors: *Joshua Jose Dias Barreto, Abhik Jana*

How to read the checklist symbols:

- ☒ the authors responded 'yes'
- ☒ the authors responded 'no'
- ☐ the authors indicated that the question does not apply to their work
- ☐ the authors did not respond to the checkbox question

For background on the checklist and guidance provided to the authors, see the [Responsible NLP Checklist](#) page at ACL Rolling Review.

☒ A. Questions mandatory for all submissions.

☒ A1. Did you describe the limitations of your work?

This paper has a Limitations section.

☒ A2. Did you discuss any potential risks of your work?

Our work evaluates large language models' reasoning capabilities concerning negation using existing datasets and prompting techniques. It does not introduce new models, generate sensitive content, or involve real-world deployment. Therefore, it poses no ethical, societal, or security risks

☒ B. Did you use or create scientific artifacts? (e.g. code, datasets, models)

☒ B1. Did you cite the creators of artifacts you used?

2

☒ B2. Did you discuss the license or terms for use and/or distribution of any artifacts?

Our study uses existing large language models and publicly available datasets without creating new artifacts. All resources used are already governed by their respective licenses, which are adhered to in our experiments. Therefore, no additional discussion on licensing or distribution terms was necessary.

☒ B3. Did you discuss if your use of existing artifact(s) was consistent with their intended use, provided that it was specified? For the artifacts you create, do you specify intended use and whether that is compatible with the original access conditions (in particular, derivatives of data accessed for research purposes should not be used outside of research contexts)?

Our work solely utilized publicly available models and datasets within the scope of their intended research purposes. Since no new artifacts were created and all resources were used as intended by their original creators, a discussion on usage consistency or intended use was not necessary.

☒ B4. Did you discuss the steps taken to check whether the data that was collected/used contains any information that names or uniquely identifies individual people or offensive content, and the steps taken to protect/anonymize it?

Our study used publicly available datasets that were already anonymized and widely accepted for research purposes. Since no additional data collection was conducted and the datasets are known to contain no personally identifiable information or offensive content, further discussion on data privacy and anonymization was deemed unnecessary.

The [Responsible NLP Checklist](#) used at ACL Rolling Review is adopted from [NAACL 2022](#), with the addition of ACL 2023 question on AI writing assistance and further refinements based on ARR practice.

- ☒ B5. Did you provide documentation of the artifacts, e.g., coverage of domains, languages, and linguistic phenomena, demographic groups represented, etc.?

2

- ☒ B6. Did you report relevant statistics like the number of examples, details of train/test/dev splits, etc. for the data that you used/created?

2

☒ **C. Did you run computational experiments?**

- ☒ C1. Did you report the number of parameters in the models used, the total computational budget (e.g., GPU hours), and computing infrastructure used?

4

- ☒ C2. Did you discuss the experimental setup, including hyperparameter search and best-found hyperparameter values?

4

- ☒ C3. Did you report descriptive statistics about your results (e.g., error bars around results, summary statistics from sets of experiments), and is it transparent whether you are reporting the max, mean, etc. or just a single run?

4

- ☒ C4. If you used existing packages (e.g., for preprocessing, for normalization, or for evaluation, such as NLTK, SpaCy, ROUGE, etc.), did you report the implementation, model, and parameter settings used?

3

☒ **D. Did you use human annotators (e.g., crowdworkers) or research with human subjects?**

- ☐ D1. Did you report the full text of instructions given to participants, including e.g., screenshots, disclaimers of any risks to participants or annotators, etc.?

(left blank)

- ☐ D2. Did you report information about how you recruited (e.g., crowdsourcing platform, students) and paid participants, and discuss if such payment is adequate given the participants' demographic (e.g., country of residence)?

(left blank)

- ☐ D3. Did you discuss whether and how consent was obtained from people whose data you're using/curating (e.g., did your instructions explain how the data would be used)?

(left blank)

- ☐ D4. Was the data collection protocol approved (or determined exempt) by an ethics review board?

(left blank)

- ☐ D5. Did you report the basic demographic and geographic characteristics of the annotator population that is the source of the data?

(left blank)

☒ **E. Did you use AI assistants (e.g., ChatGPT, Copilot) in your research, coding, or writing?**

- ☒ E1. If you used AI assistants, did you include information about their use?

Information about the use of AI assistants was not included because they were solely used for grammar checking and paraphrasing. No AI tools were involved in developing the research methodology, data analysis, or generating original content.