

Responsible NLP Checklist

Paper title: *DivScore: Zero-Shot Detection of LLM-Generated Text in Specialized Domains*

Authors: *Zhihui Chen, Kai He, Yucheng Huang, Yunxiao Zhu, Mengling Feng*

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 focuses on detecting LLM-generated text in specialized domains such as medicine and law, where the potential risks are not included due to the scope of discussion and limited space.

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

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

Appendix B.1

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

Owing to space constraints, we do not discuss detailed license information in this paper, while such details are provided in the released code repository: <https://github.com/richardChenzhihui/DivScore>.

☒ 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)?

Owing to space constraints, we do not discuss detailed license information in this paper, while such details are provided in the released code repository: <https://github.com/richardChenzhihui/DivScore>.

☒ 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?

We consider maintaining the original text necessary for LLM-generated content detection, as modifications alter the detection task. However, we carefully reviewed all datasets to ensure that they do not contain any personally identifying information or offensive content. Where applicable, we used publicly available deidentified medical and legal datasets that have been previously vetted for privacy and appropriateness. No additional personal or sensitive data was collected or included in our work.

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

Appendix B

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.

- ☒ 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?

Appendix B, Appendix C

☒ **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?

Appendix B

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

Section 4.1

- ☒ 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?

Section 4.2, Section 4.3, Section 4.4

- ☒ 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?

Appendix B

☒ **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.?

No human participants, human subjects, or annotators were involved in this work.

- ☐ 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)?

No human participants, human subjects, or annotators were involved in this work.

- ☐ 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)?

This work does not involve human participants, human subjects, or annotators. All data used in this study are from publicly available datasets with no additional data was collected.

- ☐ D4. Was the data collection protocol approved (or determined exempt) by an ethics review board?
- This work does not involve any collection of data from human participants, human subjects, or annotators.*

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

This work does not involve any collection of data from human participants, human subjects, or annotators.

☒ **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?

We strictly adhere to the ACL Policy on AI Writing Assistance. Our use of AI assistants was limited to checking the grammar and polishing the language of the author-written text.