

Responsible NLP Checklist

Paper title: *Deconstruct, Diagnose, and Deliberate: A Protocol-Adaptive Role-Specific Multi-Agent Framework for Fake News Detection*

Authors: Zhigan Zhang, Jie Sui

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?

In the Limitation.

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

☒ 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?

In the appendix O. Our work does not involve any personal identifiable information or offensive content. The data used in this study consists of publicly available news articles and claims.

☒ 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?

In the Appendix A.

☒ C. Did you run computational experiments?

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

In this paper, we did not conduct hyperparameter search or tuning. The experimental setup focused primarily on the framework design and performance evaluation, without specifically optimizing the models hyperparameters. Future work may consider hyperparameter tuning to further enhance the models performance and investigate the impact of different settings.

☒ 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?

We did not provide descriptive statistics primarily due to the limitations of experimental resources and computational costs in this study. To focus on the framework design and performance evaluation, we conducted single-run experiments and compared the model’s performance with existing baselines.

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. ACL 2026 used a subset of ARR checklist form.

☒ **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.?

In the Appendix B

- ☒ 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, we did not compensate the annotators, all annotators participated voluntarily.

- ☒ 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)?

In the appendix O. Participants were informed of the research purpose and gave their consent before proceeding.

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

No, the data collection plan has not been approved by the ethics review board because the data used is publicly available and does not involve the collection of personal data.

☒ **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?

In the appendix O. AI assistants were employed solely for language refinement and LaTeX structuring to improve readability. The core methodology and experimental analysis remain the original work of the authors.