

Responsible NLP Checklist

Paper title: *From Implicit Exploration to Structured Reasoning: Guideline and Refinement for LLMs*

Authors: *Jiaxiang Chen, Zhuo Wang, mingxi Zou, Zhucong Li, Zhijian Zhou, Song Wang, Zenglin Xu*

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?

(left blank)

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

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

Section 3 (Method) cites model sources such as GPT-4o.

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

Section 4 (Experimental Setup) discusses licensing and public usage terms of BBH and GPT APIs.

- ☒ 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)?

Section 4 clarifies that our use of models and datasets (e.g., BBH) follows their intended research use.

- ☐ 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?

(left blank)

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

No new artifacts were created, so documentation is not applicable.

- ☒ 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?

Section 4 includes dataset statistics and benchmark coverage.

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.

☒ **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?

Section 4.2. We report the models used (e.g., GPT-4o, GPT-4o-mini, LLaMA-3Qwen3), and include parameter scales for open-source models. All models are accessed via APIs (e.g., OpenAI), so GPU hours and hardware configurations are not applicable.

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

Section 4.2. All models are used via inference APIs with fixed temperature and decoding settings. No hyperparameter tuning was conducted.

- ☒ 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 5. We report mean accuracy across three runs. Since model outputs are stable under fixed decoding settings, we omit variance reporting.

- ☒ 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?

Section 4.1. We use accuracy as the sole evaluation metric and compute it with simple scripts; no external evaluation packages were used.

☒ **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?

We used ChatGPT to assist in drafting the paper.