

Responsible NLP Checklist

Paper title: *Uncovering Scaling Laws for Large Language Models via Inverse Problems*

Authors: *Arun Verma, Zhaoxuan Wu, Zijian Zhou, Xiaoqiang Lin, Zhiliang Chen, Rachael Hwee Ling Sim, Rui Qiao, Jintan Wang, Nhung Bui, Xinyuan Niu, Wenyang Hu, Gregory Kang Ruey Lau, Zi-Yu Khoo, Zitong Zhao, Xinyi Xu, Apivich Hemachandra, See-Kiong Ng, Bryan Kian Hsiang Low*

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?

This position paper argues how we can use inverse problems to efficiently uncover underlying scaling laws for LLMs. We do not foresee any potential risks.

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

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

We did not use or create any artifact.

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

We did not use or create any artifact.

- ☐ 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)?

We did not use or create any artifact.

- ☐ 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 did not use or create any artifact.

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

We did not use or create any artifact.

- ☐ 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?

We did not use or create any artifact.

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?

We did not use any models, as this is a position paper arguing a research direction rather than proposing a new methodology.

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

We did not include any experiments, as this is a position paper arguing a research direction rather than proposing a new methodology.

- ☐ 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 include any experiments, as this is a position paper arguing a research direction rather than proposing a new methodology.

- ☐ 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?

We did not use or include any packages.

☒ **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.?

There were no human participants involved.

- ☐ 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)?

There was not recruitment or participants.

- ☐ 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)?

We did not use or curate data.

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

We did not perform any data collection.

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

We did not introduce or use new data, and thus did not require any 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 did not use any AI assistant.