

Responsible NLP Checklist

Paper title: *PEPE: Long-context Extension for Large Language Models via Periodic Extrapolation Positional Encodings*

Authors: *Jikun Hu, Dongsheng Guo, Yuli LIU, Qingyao Ai, Lixuan Wang, Xuebing Sun, Qilei Zhang, Quan Zhou, Cheng Luo*

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 4

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

The artifacts used in this work (e.g., models, datasets) are standard and publicly available resources that are widely adopted in the NLP community. Their licenses are already clearly documented in their original sources (e.g., Hugging Face, official repositories), and our usage fully complies with their terms. As our work does not modify or redistribute these artifacts, we did not include a separate license discussion in the paper.

- ☒ 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)?

The intended use of the existing artifacts (e.g., pre-trained models, datasets) is general-purpose NLP research, and our work falls within this scope. For the artifacts we created, such as code, the intended use is research replication and extension. While we did not explicitly discuss compatibility with access conditions in the paper, our usage strictly adheres to the original licenses and research-only restrictions where applicable.

- ☒ 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?

The data used in this work is publicly available and widely adopted in the NLP community. These

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.

datasets are generally considered free of personally identifiable information (PII) and offensive content.

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

We have released the code publicly (as noted in the abstract), but the paper does not include detailed documentation on domains, languages, or linguistic phenomena. A basic README is provided in the repository.

- ☒ 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

☒ **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 4.1

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

Section 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?

Sections 4.2 and 4.3

- ☒ 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 4.2

☒ **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 only used AI assistants to help check the paper writing and code, not for generating core content.