

Responsible NLP Checklist

Paper title: *GeLoRA: Geometric Adaptive Ranks For Efficient LoRA Fine-tuning*

Authors: *Abdessalam Ed-dib, Zhanibek Datbayev, Amine M. Aboussalah*

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 improving the efficiency of parameter-efficient fine-tuning through a theoretically grounded method for adaptive rank selection. As it does not introduce new data, tasks, or deployment practices, we believe it does not pose novel risks beyond those already inherent to existing LLM fine-tuning methods.

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

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

We used publicly available pre-trained models and benchmarks (e.g., DeBERTa, LLaMA, GLUE, SQuAD, MT-Bench), and cited their original creators in Section 5, "Experiments."

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

We used publicly available datasets (GLUE, SQuAD, MT-Bench) and models (DeBERTa, LLaMA) under their respective licenses, which permit academic research use. As we did not redistribute any artifacts or modify their licensing terms, we did not explicitly discuss licenses 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)?

We used all artifacts (datasets and models) strictly for academic research purposes, consistent with their intended use. Since we did not repurpose, redistribute, or deploy them outside standard research settings, we did not explicitly include this discussion in the paper.

☒ 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 used only publicly available benchmark datasets (GLUE, SQuAD, MT-Bench), which are widely adopted in the community and presumed to have undergone appropriate filtering and anonymization. As we did not collect or modify any data ourselves, we did not explicitly discuss these aspects in the paper.

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.

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

We relied on established benchmark datasets (GLUE, SQuAD, MT-Bench) whose documentation is publicly available. Since we did not create or modify these datasets, we did not repeat their documentation in the paper.

- ☒ 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?

Relevant dataset statistics, including the number of examples and details about train/dev/test splits, are provided in Appendix.

☒ **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 report model sizes (e.g., number of parameters in DeBERTa and LLaMA) and describe the computing infrastructure used (e.g., GPU type and number) in Section 5, "Experiments." Total compute time is also discussed to contextualize efficiency comparisons.

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

We detail the experimental setup, including training procedures, LoRA rank selection strategies, and key hyperparameters (learning rate, batch size, etc.) in Section 5, "Experiments." Best-found hyperparameter values are reported for reproducibility in Appendix.

- ☒ 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 report mean performance across multiple runs and include standard deviations where applicable. It is clearly stated in Section 5, "Experiments," whether results are from single runs or averaged across multiple seeds.

- ☒ 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 specify the packages used (e.g., Hugging Face Transformers, Scikit-Dimension, PyTorch, ...) along with key implementation details and parameter settings in Section 5, "Experiments."

☒ **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 annotators or participants were 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)?

No human annotators or participants were involved.

- ☐ 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 used publicly available datasets with existing consent protocols; we did not collect new data.

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

No new data collection or human subjects research was conducted.

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

No human annotators or participants were involved.

☒ **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?

AI assistants were used only to correct grammatical mistakes in the paper.