
B SUPPLEMENTARY MATERIALS

ETHICS STATEMENT

This work adheres to the ICLR Code of Ethics.¹⁰ Our research focuses on developing methods for token embedding alignment in generative recommenders. The experiments rely solely on publicly available datasets (Amazon Product Reviews and Yelp Open Dataset), which contain no personally identifiable information beyond what is publicly released. We do not foresee direct risks of harm to individuals or groups arising from this research. Nevertheless, as with all recommender systems, potential societal impacts include bias amplification and unintended reinforcement of popularity effects. We note these risks and emphasize that our contributions are methodological rather than application-specific; the proposed techniques can be combined with fairness-aware or debiasing mechanisms. No human subjects were involved, and no IRB approval was required.

REPRODUCIBILITY STATEMENT

We are committed to facilitating reproducibility and transparency of our work. To this end, we intend to fully open source **all** of our code, data, and trained models after publication.

- **Code and Implementation:** We will provide an open-sourced codebase link to the full implementation and training scripts after publication.
- **Datasets:** All datasets used (Amazon Product Reviews and Yelp Open Dataset) are publicly available. Detailed preprocessing steps, including the 5-core filtering strategy and sequence truncation to length 20, are described in Appendix A.1. For semi-synthetic search query datasets we constructed, we will also publish it to accelerate research in search recommendation systems.
- **Model and Training Details:** Hyperparameters (learning rates, batch sizes, epochs, optimizer choices) and architectural specifications (Qwen3-0.6B configuration) are included in Section 3.1 and Appendix A.4.
- **Evaluation:** Metrics, evaluation protocols, and baselines are fully documented in Section 3.2.1 and Appendix A.3.

Together, these materials should enable independent researchers to reproduce our findings.

THE USE OF LARGE LANGUAGE MODELS

Large language models were used in two ways during the preparation of this manuscript. First, we employed commercial LLMs solely to edit for clarity, grammar, and academic style, without altering the authors’ intended meaning or contributions. Second, we used open-source LLMs—with a clearly specified prompting strategy A.2.2—to generate synthetic search recommendation datasets. In both cases, the authors exercised full oversight and accept responsibility for all claims, analyses, and conclusions.

¹⁰<https://iclr.cc/public/CodeOfEthics>