

Optimizing a Contextual Chatbot for the POLYCC: A Competition-Based LLM Fine-Tuning Experience

Siti Zaleha Ibrahim*, Norsuzila Shafie, Wan Siti Rodziah Mohd Nasir, Mohd Suhairi Md Suhamin

Politeknik Kota Bharu, Malaysia

szaleha@pkb.edu.my; norsuzila@pkb.edu.my; rodziah@pkb.edu.my; suhairisuhaimin@pkb.edu.my

Abstract— This paper reports a competition-based experience in fine-tuning a large language model (LLM) for the Malaysian Polytechnic and Community College (POLYCC) domain using the SynLoRA-SGS methodology. Participating in the student category of the POLYCC LLM League 2025, our team fine-tuned a Meta-Llama-3-8B-Instruct model through bilingual synthetic data generation and Low-Rank Adaptation (LoRA) hyperparameter optimization on Amazon SageMaker. Despite qualifying in last position (6th of 6 finalists) during the automated leaderboard stage, the team achieved second place overall in the final multi-source evaluation with a grand total of 149 points, including the highest spectator score among all finalists. This result demonstrates that focused dataset refinement between competition stages can produce substantial performance gains, and that a systematic fine-tuning approach can overcome initial ranking disadvantages when evaluated through holistic multi-source assessment.

Keywords— Large Language Model, Fine-Tuning, LoRA, POLYCC, Competition-Based Evaluation

I. INTRODUCTION

The POLYCC LLM League 2025, organized by Politeknik dan Kolej Komuniti (POLYCC) and Bahagian Instruksional dan Pembelajaran Digital (BIPD), is a competition-based evaluation framework in which participants fine-tune a provided large language model (LLM) to build domain-specific contextual chatbots. The competition featured two categories: (i) lecturer team (*Pasukan Pensyarah*), and (ii) student team (*Pasukan Pelajar*). Each category with separate leaderboards and finalists.

This paper documents our experience participating in the student category as the Politeknik Kota Bharu (PKB) team. The competition required participants to fine-tune a Meta-Llama-3-8B-Instruct model using Amazon SageMaker JumpStart, with performance determined solely by dataset construction and hyperparameter selection, no architectural modifications were permitted. For our detailed fine-tuning methodology, we refer readers to the SynLoRA-SGS Framework described in main article by lecturer team, which documents the four-stage pipeline applied in this work: multi-model synthetic data generation, quality-controlled curation, LoRA-based fine-tuning with selective grid search and model deployment.

The topic focused on an entrepreneurship and innovation initiative within the POLYCC ecosystem that promotes student-driven product development and marketplace activities [1]. Developing a chatbot for this domain required generating specialized bilingual (Malay–English) training data covering Technical and Vocational Education and Training (TVET) policies, operational guidelines and frequently asked questions which content not present in the base model's pre-training data. Bilingual NLP tasks in the Malaysian context present unique challenges due to code-mixing phenomena as demonstrated in prior work on bilingual area [2]. NLP applications targeting Malaysian bilingual contexts have expanded across domains including public security [3] and education.

Our team's trajectory presents an interesting case study: we qualified in last position (6th of 6 finalists) during the automated leaderboard stage yet achieved second place overall in the final evaluation. The following sections analyse the contributing factors to this progression and provide a detailed report of the competition results.

II. COMPETITION STRUCTURE

The POLYCC LLM League 2025 followed a two-stage evaluation framework. In the qualification stage, participants submitted fine-tuned models to an automated leaderboard powered by the AWS League of LLMs (LOL) platform. Models were evaluated through pairwise comparisons against other submissions, with performance measured by winning rate (WR), the percentage of pairwise matchups won. The top six teams in each category advanced to the final round.

The final round introduced a new, previously unseen topic to assess generalization capability. Teams rapidly constructed new domain-specific datasets and fine-tuned their models within a limited timeframe. Final

*Corresponding Author

evaluation employed a multi-source assessment framework comprising three weighted components: automated AI evaluation, expert judge scoring and live spectator (audience) voting. The grand total combined qualification-stage and final-stage scores to determine overall ranking.

III. APPROACH

Our fine-tuning approach followed the SynLoRA-SGS methodology described in main article. This section briefly summarises the key steps as applied to the TVET domain. Leaders are referred to the main article for full framework details and theoretical justification.

A. Dataset Construction

Bilingual instruction–response pairs were generated using three commercial LLM services (ChatGPT, Gemini and Claude) prompted with Maker Market source documents. This multi-model approach increased linguistic diversity and reduced single-model bias [4-6]. The raw outputs were aggregated and subjected to quality control: deduplication, fact-checking against source documents and formatting standardization [7-9]. Dataset sizes ranged from approximately 1,200 to 4,000 entries across iterations with bilingual (Malay–English) balance maintained throughout.

B. Fine-Tuning Configuration

LoRA fine-tuning [8, 10] was performed on Meta-Llama-3-8B-Instruct via Amazon SageMaker JumpStart using ml.g4dn.12xlarge instances. Hyperparameter selection followed the selective grid search strategy, exploring configurations across LoRA rank (r), scaling factor (α), dropout, learning rate and epoch count. Guided on the dominant configuration, the final-stage models converged on $r = 192$, $\alpha = 896$, with learning rates of $1.0\text{--}1.5 \times 10^{-4}$ and 3–10 epochs. Each team was allocated 30 hours of total training time, necessitating rigorous budget management across experimental iterations.

IV. RESULTS

A. Qualification Stage

Table I presents the qualification results for the student category (*Pasukan Pelajar POLYCC*). Six teams advanced to the final round based on their automated leaderboard performance.

TABLE I
QUALIFICATION STAGE RESULTS (STUDENT CATEGORY)

Rank	Institution	WR (%)	NOS
1	PSA	62	133
2	KK Gopeng	40	16
3	PMJ	38	16
4	POLIMAS	34	80
5	PUO	34	92
6	PKB (Ours)	32	1

WR = Winning Rate; NOS = Number of Submissions

PKB qualified in last position with a winning rate of 32% and a single submission. While this placed the team at the bottom of the finalist table, it secured qualification for the final round, where the evaluation framework shifted from automated pairwise comparison to a multi-source assessment that included human expert and audience judgement.

A. Final Stage Preparation and Fine-Tuning

The transition from qualification to the final round represented the most critical phase of our competition experience. Having qualified in last position, the team undertook an intensive preparation process targeting two areas: (i) constructing a higher-quality bilingual dataset for the new Maker Market topic and (ii) systematically refining the LoRA hyperparameter configuration based.

For dataset construction, the team applied the multi-model synthetic generation pipeline, processing TVET source documents through ChatGPT, Gemini and Claude in parallel. Unlike the qualification stage, where time constraints limited quality control, the final-stage dataset underwent more rigorous cross-lingual validation. Each generated Malay–English pair was manually reviewed for factual consistency against the source material.

The curated dataset emphasized concise, contextually grounded responses to maximize domain relevance within the training budget.

For LoRA fine-tuning, the team adopted the dominant hyperparameter region, centring on $r = 192$ and $\alpha = 896$, and systematically varied epoch count and dropout to optimize convergence. Table II summarizes the iterative fine-tuning process during the final stage.

TABLE II
FINAL STAGE FINE-TUNING ITERATIONS

Model	Ep	(r , α)	Loss↓	PPL↓	Observation
Ctx-1	3	(192, 896)	0.697	2.01	Correct facts, concise responses
Ctx-2	7	(192, 896)	0.438	1.55	Mixed response length, stable
Ctx-3	5	(192, 896)	0.409	1.51	Improved contextual alignment
Ctx-4	9	(192, 896)	0.410	1.51	Longer answers, minor BM flaws
Ctx-5	10	(192, 896)	0.300	1.35	Grounded, context-aware

Ep = Epochs; Loss = Eval Loss; PPL = Perplexity; BM = Bahasa Malaysia

Across five iterations, Evaluation Loss improved from 0.697 to 0.300 and Perplexity dropped from 2.01 to 1.35, indicating progressive gains in contextual language modeling. Notably, the epoch progression was non-monotonic (3, 7, 5, 9, 10); Ctx-3 at 5 epochs achieved lower loss than Ctx-2 at 7 epochs, indicating that the dataset refinements applied between iterations contributed more to performance than extended training alone. This finding is consistent with the data-quality-over-size principle.

The final submitted model (Ctx-5) demonstrated strong contextual grounding for Maker Market queries in both Malay and English with factually accurate and domain-relevant responses. Some sensitivity to prompt phrasing remained, particularly for Malay inputs requiring elaboration but overall the model exhibited reliable conversational behaviour within the target domain.

B. Competition Outcome

Under the multi-source final evaluation comprising automated AI scoring, expert judgment and live spectator voting, the PKB team achieved a grand total of 149.0 points, securing second place in the student category (*Pasukan Pelajar*) among six finalist teams. The team attained the joint-highest spectator score (6.0), tied with the first-place team, indicating that audience perception of the chatbot’s practical usefulness was comparable to the top-ranked finalist.

The progression from last-place qualifier (6th, WR 32%) to second-place finalist represents a substantial improvement attributable to the intensive final-stage preparation process. While the qualification stage relied on a single baseline submission, the final stage benefited from systematic dataset refinement and iterative LoRA optimization as detailed in Table II. This outcome validates the effectiveness of a quality-focused fine-tuning strategy over submission-volume-driven approaches.

All experiments were conducted using AWS SageMaker with Hugging Face training utilities, following the official workflow and constraints specified by the POLYCC LLM League 2025 organizers. The competition infrastructure enforced standardised training, evaluation and submission procedures to ensure fairness and reproducibility across participants.

V. DISCUSSION

The PKB team’s trajectory from last-place qualifier to second-place finalist highlights several insights relevant to applied LLM fine-tuning in constrained settings.

First, the iterative refinement process documented in Table II demonstrates that systematic preparation between competition stages can overcome initial ranking disadvantages. Each model iteration involved not only hyperparameter adjustments but also targeted dataset improvements by removing low-quality entries, rebalancing bilingual coverage and adding contextually specific Maker Market content. The non-monotonic epoch pattern across iterations confirms that these dataset-level changes, rather than training duration alone, drove performance gains.

Second, the TVET topics domain presented specific challenges for synthetic data generation. Unlike general-knowledge topics, TVET content is institutionally specific and not represented in the base model’s pre-training corpus. The challenge of maintaining model performance when transferring across linguistic contexts has been documented in related NLP domains [5]. This required careful curation of source documents and cross-lingual validation of generated instruction–response pairs to avoid hallucinated policy details. The multi-model

generation strategy proved particularly effective here, as different LLM teachers produced complementary perspectives on TVET operations, enriching the training data diversity.

Third, the competition experience revealed practical constraints rarely discussed in literature. The 30-hour training budget per team required disciplined experimental planning: each training run consumed 20–40 minutes, meaning the team could execute approximately 45–90 experiments in total. The selective grid search strategy proved essential in navigating the hyperparameter space efficiently within this budget, converging on the dominant configuration ($r = 192$, $\alpha = 896$) without exhaustive exploration.

VI. CONCLUSION

This paper reported the experience of a student team from Politeknik Kota Bharu in the POLYCC LLM League 2025, applying the SynLoRA-SGS methodology to fine-tune a contextual LLM chatbot for the Maker Market domain. Despite qualifying in last position, the team achieved second place in the overall student category through focused dataset refinement and systematic hyperparameter optimization. The results demonstrate that competition-based LLM evaluation frameworks provide valuable structured learning environments for students and that a quality-focused fine-tuning strategy can yield competitive performance under strict resource constraints.

ACKNOWLEDGEMENT

We would like to express our sincere appreciation to Politeknik dan Kolej Komuniti (POLYCC) and Bahagian Instruksional dan Pembelajaran Digital (BIPD) for organizing the POLYCC LLM League 2025. Special thanks are extended to AWS Malaysia for providing the cloud infrastructure and hosting support that enabled the experimentation and evaluation processes. We also gratefully acknowledge the support and encouragement from management and colleagues at Politeknik Kota Bharu (PKB), whose cooperation and commitment played an important role in the successful completion of this work.

DECLARATION OF GENERATIVE AI USAGE

During the preparation of this work, the authors used ChatGPT, Gemini and Claude to generate synthetic training data and assist in proofreading the manuscript. The authors declare that they reviewed and edited the final output and take full responsibility for the content.

REFERENCES

- [1] W. Zahari and M. S. Md Suhaimin, Impact of a Project-Based Learning Competition on Engineering Students' Interest and Skills. 2025.
- [2] M. S. Md Suhaimin, A. Wibowo, E. G. Mounq, P. Anthony, and M. H. Ahmad Hijazi, "Multitask deep learning for sentiment analysis with sarcasm detection in bilingual code-mixed social media content," 2026, Bilingual code-mixed; Deep learning; Hybrid feature engineering; Language model; Multitasking vol. 15, no. 2, p. 12, 2026-04-01 2026, doi: 10.11591/eei.v15i2.10935.
- [3] M. S. M. Suhaimin, M. H. A. Hijazi, and E. G. Mounq, "MSSThreD: Multitask Social Media Sentiment Analysis with Sarcasm for Public Security Threat Detection," in 2024 11th International Conference on Information Technology, Computer, and Electrical Engineering (ICITACEE), 2024: IEEE, pp. 25-30.
- [4] S. Gunasekar et al., "Textbooks are all you need," arXiv preprint arXiv:2306.11644, 2023.
- [5] M. S. M. Suhaimin, M. H. A. Hijazi, W. S. R. M. Nasir, A. Wibowo, and E. G. Mounq, "The Challenge of Generalization: Preserving Sarcasm Detection in a Multitask Model Across Different Linguistic Contexts," in 2025 International Conference on Advances in Machine Intelligence, and Cybersecurity Technologies (AMICT), 2025: IEEE, pp. 49-53.
- [6] Y. Wang, Y. Kordi, S. Mishra, A. Liu, N. A. Smith, D. Khashabi, and H. Hajishirzi, "Self-instruct: Aligning language models with self-generated instructions," in Proceedings of the 61st annual meeting of the association for computational linguistics (volume 1: long papers), 2023, pp. 13484-13508.
- [7] Z. Ji et al., "Survey of hallucination in natural language generation," ACM computing surveys, vol. 55, no. 12, pp. 1-38, 2023.
- [8] S. Raschka, "Practical tips for finetuning llms using lora (low-rank adaptation)," Ahead of AI (Nov. 2023). url: <https://sebastianraschka.substack.com/p/practical-tips-for-finetuning-llms>, p. 67, 2024.
- [9] M. S. M. Suhaimin, M. H. A. Hijazi, E. G. Mounq, and M. A. M. Hamza, "Data Augmentation Approach for Language Identification in Imbalanced Bilingual Code-Mixed Social Media Datasets," in 2023 IEEE International Conference on Artificial Intelligence in Engineering and Technology (IICAET), 2023: IEEE, pp. 257-261.
- [10] E. J. Hu et al., "Lora: Low-rank adaptation of large language models," Iclr, vol. 1, no. 2, p. 3, 2022.