

Introduction to the Special Issue: POLYCC LLM League 2025

Akashah bin Arshad*, Muhammad Syamil bin Rozmi, Zaidi bin Othman, Mohd Najib bin Hamdan, Mazlaine binti Husain, Nur Zatil Iman binti Abu Bakar, Nadirah binti Abd Aziz, Mohd Suhairi Md Suhamin

Bahagian Instruksional dan Pembelajaran Digital, Jabatan Pendidikan Politeknik dan Kolej Komuniti, Malaysia

akashah@mohe.gov.my; syamil.rozmi@mohe.gov.my, zaidi.othman@mohe.gov.my, najib.hamdan@mohe.gov.my, mazlaine@mohe.gov.my, zatil.bakar@mohe.gov.my, nadirah.aziz@mohe.gov.my; suhairisuhaimin@pkb.edu.my

I. THE EVOLVING LANDSCAPE OF APPLIED AI

Recent years have marked a transformative shift in how organizations interact with information, driven largely by the mainstreaming of Artificial Intelligence (AI). While global benchmarks often focus on high-resource models, there is an urgent need to address the scalability and adaptability of Large Language Models (LLMs) in specific, resource-constrained institutional environments.

The gamified learning Gamification is a popular pedagogical approach in computer science education, and advances in LLMs enable dynamic and engaging learning experiences, including competition driven by AI. For instance, a game-show-themed learning tool called study introduces Quack the Code, where students teach an LLM-powered "debug duck" while competing in an AI-hosted game [1]. AI was used to generate personalized avatars, usernames, questions, host responses, and interactions with game elements. In December 2024, Amazon Web Services (AWS) launched the AWS Large Language Model League dubbed as AWS LLM League during re:Invent 2024 [2]. Jabatan Pendidikan Politeknik & Kolej Komuniti (JPPKK) was given the honour to be the first public institution in Malaysia to hosts the AWS LLM league which later known as the POLYCC LLM league for two years [3]. POLYCC league was first held in 2024 after the great support and participation of POLYCC Deep Racer Competition in 2023. The Deep racer competition is a gamified competition was first introduced in 2028 during at re:Invent 2018 which transformed the complex, often intimidating world of reinforcement learning into a high-stakes, 1/18th-scale digital and physical playground. It invited developers of all skill levels to step away from dry theory and into the driver's seat, where they could train autonomous racing models in a cloud-based 3D simulator before deploying them onto physical tracks to navigate hair-pin turns and straightaways [4].

Ultimately, this special issue serves as more than a technical summary; it is a testament to the potential of the POLYCC community to lead the charge in Malaysia's digital transformation. By bridging the gap between rigorous academic inquiry and the gamified spirit of the POLYCC LLM League 2025, these contributions offer students and staff a unique avenue to master emerging technologies outside the traditional classroom. Furthermore, by fostering the development of digital agents that are as ethically grounded as they are technically proficient, these initiatives directly support the Ministry of Digital's vision of establishing Malaysia as a premier AI Nation by 2030 [5].

II. THE POLYCC LLM LEAGUE 2025: A COMPETITION FOR INNOVATION

JPPKK an organization under Malaysia's Ministry of Higher Education (MOHE), collaborated with AWS to launch new gamified training programs aimed at educators and students to improve Malaysia's technical and vocational education and training (TVET) community to close the country's AI skills gap [5]. The finale of POLYCC LLM League 2025 was held in Shah Alam Convention Centre on 17th August 2025 in conjunction with 13th CIDOS Inspiring Learning Awards. All POLYCC employees and students from 36 polytechnics and 106 community colleges are eligible to participate and compete in 2025. Students, lectures, and staff who form teams to play against other teams make up the 2501 registered participants. JPPKK achieved the Malaysia Book of Records (MBOR) for "Largest Participation in AI Large Language Model (LLM) League" with 2,501 participants on 19th August 19 2025, during the prize ceremony. Each team can use their allocated training hour credit to use Unified Sagemaker Studio to fine tune LLM Llama 3.2B model. Users fine-tune the LLM model based on selected theme using several hyperparameters to give the best output. they customized it with their curated datasets, adjusting hyperparameters. Some examples of hyperparameters such Epochs, Learning and LoRA parameter and few others advanced parameters can be utilized to fine tune LLM model based on participant strategy and creativity. The submitted model would be compared against a bigger 90B reference model with the quality of the responses decided using an LLM-as-a-Judge approach. Participants score a win for each question where the LLM judge deemed the fine-tuned model's response to be more accurate and

*Corresponding Author

comprehensive than that of the larger model. The competition enforced a rigorous two-stage evaluation protocol: Qualification Stage: An automated assessment focusing on technical robustness and leaderboard ranking. Participants need to train the model on selected topic Final Round: A multi-dimensional assessment that defied "common sense" benchmarks, requiring models to adapt to new, unseen topics.

III. CORE THEMES: CURATION, OPTIMIZATION, AND EVALUATION

The papers included in this special issue highlight three critical layers of contextual chatbot development:

A. *Synthetic Data Curation:*

Synthetic data generation (SDG) leveraging LLMs has emerged as a transformative solution to the persistent challenges of data scarcity, high labeling costs, and stringent privacy constraints. By utilizing techniques such as prompt-based generation in zero-shot or few-shot settings and fine-tuning on domain-specific corpora, LLMs can produce scalable, cost-effective alternatives to human-annotated datasets. These models effectively augment existing data to address class imbalances and enhance performance in specialized tasks, such as sentiment analysis and dependency parsing. Furthermore, the integration of differential privacy mechanisms enables the replication of complex real-world patterns while safeguarding sensitive information, a critical requirement in highly regulated sectors like healthcare and finance. In the context of the current competition, participants utilized the Amazon PartyRock playground to generate synthetic data. PartyRock provides access to various foundational models via Amazon Bedrock, facilitating the creation of no-code AI applications to generate synthetic training sets for fine-tuning. Additionally, participants were encouraged to employ diverse conversational agents to iteratively refine and improve the quality of the resulting question-and-answer datasets based on the designated themes [6-7].

Despite these advancements, significant challenges persist regarding the quality, realism, and ethical implications of LLM-generated data. A primary concern is the risk of bias amplification, wherein synthetic outputs reinforce or exacerbate existing prejudices inherent in the underlying training data. Additionally, achieving functional and distributional realism in high-dimensional or multimodal contexts remains computationally intensive and frequently fails to generalize across diverse institutional settings. These limitations are further compounded by the absence of standardized evaluation protocols, which complicates the consistent measurement of data utility [8]. Consequently, while SDG represents a promising frontier for machine learning, future research must prioritize the development of robust evaluation frameworks and iterative refinement techniques to ensure the reliability and ethical integrity of synthetic outputs.

B. *Resource-Optimized Fine-Tuning:*

Fine-tuning LLMs on structured question-response datasets necessitates the strategic optimization of hyperparameters, specifically through efficient Parameter-Efficient Fine-Tuning (PEFT) techniques like Low-Rank Adaptation (LoRA). The calibration of epochs and learning rates is fundamental; while LoRA allows for reduced iteration cycles due to its inherent computational efficiency, maintaining stable gradient dynamics through adaptive learning rates is essential for balancing convergence speed with model stability. Central to this process is the configuration of LoRA-specific parameters, where empirical evidence suggests that utilizing higher ranks, such as 16, 32, or 64; alongside a consistent alpha-to-rank ratio of 4:1 significantly enhances performance in structured text tasks. Advanced variants, including QLoRA and SBoRA, further refine this process by dynamically allocating computational resources across model layers, thereby mitigating memory overhead while preserving factual accuracy and fluency. Ultimately, the successful deployment of LLMs for domain-specific question-answering tasks, such as those in cybersecurity or biomedical research, relies on this synthesis of high-quality dataset curation and the precise tuning of low-rank adapters to achieve high-fidelity outputs in resource-constrained environments.

While the optimal number of epochs varies by task, dataset, and computational constraints, most studies suggest that fine-tuning within 3–5 epochs is effective for question-response datasets. This range balances performance, efficiency, and privacy considerations. For tasks requiring higher precision or domain-specific adaptation, additional experimentation may be necessary to identify the ideal epoch count. Adjusting the learning rate during LLM fine-tuning is a pivotal factor that influences model performance, convergence speed, and computational efficiency. By employing adaptive and component-specific learning rates, practitioners can optimize fine-tuning outcomes for diverse tasks while mitigating challenges like overfitting and inefficiency.

C. *Multi-Source Evaluation:*

The "LLM-as-a-Judge" framework represents a paradigm shift in automated evaluation, utilizing LLMs to assess the outputs of other generative systems with a degree of scalability and efficiency that human review cannot match. This approach is increasingly employed to evaluate linguistic coherence, legal reasoning, and factual integrity; however, its implementation is frequently complicated by inherent algorithmic biases and a susceptibility to prompt sensitivity. While LLMs can align closely with general human preferences, they often

lack the specialized domain expertise required for nuanced fields such as medicine or law, where hallucinations and a lack of procedural transparency pose significant ethical risks. To mitigate these limitations, emerging methodologies including collaborative criteria refinement tools like MetricMate and the integration of crowd-based reasoning were aimed to enhance the reliability of these automated judges. Nevertheless, while domain-specific fine-tuning offers a pathway toward improved accuracy, rigorous human oversight remains indispensable to ensure fairness and mitigate the risk of bias amplification in high-stakes evaluative contexts.

AWS SageMaker provides a comprehensive and scalable platform for hyperparameter tuning, leveraging advanced optimization techniques and cloud resources. Its features, such as AMT, early stopping, and warm-starting, make it a powerful tool for improving model performance across various domains. However, users must carefully manage computational costs and consider dataset characteristics to maximize the benefits of hyperparameter tuning.

IV. FEATURED CONTRIBUTIONS AND ACHIEVEMENTS

This collection features the most successful methodologies from the league's finalists. The POLYCC LLM League 2025 is structured into two primary phases: a virtual preliminary round focused on skill acquisition and a high-stakes, on-site Grand Finale. The preliminary phase commences with a formal registration period during the month June 2025, followed by an online qualifying round beginning on the 1st July 2025. This stage serves as a technical filter to identify high-potential talent within the POLYCC community. To facilitate knowledge transfer, participants are provided with a comprehensive support ecosystem, including virtual briefings and an Online Coaching Clinic via the YouTube EDUTV@POLYCC platform. These sessions focus on critical competencies, including navigating Amazon SageMaker, dataset synthesis, and the fine-tuning of Large Language Models (LLMs) via hyperparameter optimization. Furthermore, an intensive "A+ Bootcamp" held on 27th July allows for direct technical consultation with AWS industry experts. This phase concludes on 1st August 2025, with the top six teams from both the student and staff categories advancing to the final round.

The Grand Finale, held on 17th August 2025, at the Shah Alam Convention Centre (SACC), transitions the competition into a rigorous, face-to-face format designed to simulate real-world AI deployment pressures. The final consists of seven rounds for each category, where teams must demonstrate technical precision and tactical coordination. To ensure individual accountability, only one representative per team is permitted on stage at a time, with substitutions allowed only after the third round. Each round follows a strict four-minute sequence: a 30-second question to be prompt is read by MC, a 90-second execution window for the participant to plan and strategize model configuration and prompt engineering, and a 120-second live evaluation. Performance is measured through a tripartite scoring model: an AI jury assessment (40%), a professional jury evaluation (40%), and spectator engagement via Telegram (20%) who are watching through Live YouTube Streaming from all over Malaysia. This multi-modal evaluation ensures that finalists possess not only technical mastery but also the professional advocacy skills necessary for industry-ready AI development. The analysis of hyperparameter configurations for the PolyCC LLM League reveals that superior win rates achieved by the student team are primarily driven by a strategic emphasis on extended training cycles and high-quality synthetic data refinement. Top-performing institutions, such as Politeknik Kota Bharu (PKB) and Politeknik Mersing Johor (PMJ), achieved optimal results by utilizing 10 to 14 epochs, a significantly more intensive approach than standard fine-tuning protocols, coupled with a stable learning rate of 0.0001. LoRA with a Rank of 16 and an Alpha of 64 provided the necessary model capacity to internalize domain-specific nuances without overfitting. Furthermore, the data suggests that dataset refinement is more critical than scale; for instance, Politeknik Sultan Sallehuddin maintained competitiveness with only 276 highly curated entries, while PKB leveraged multi-model generation strategies to produce between 1,200 and 4,000 refined rows. Collectively, these findings indicate that the "winning formula" integrates aggressive temporal training with a curated, multi-model synthetic data strategy to maximize LLMs performance in competitive contextual environments.

On the other hand, lecturer team compete with each other with different set of dataset theme which based on Technology Enabled Classroom (TECC) and Maker Market guideline. Both guidelines are used reference to guide lecturers in utilizing learning space to enhance learning experiences and spark creativity. The following two paragraphs provide a formal synthesis of the findings for the analysis and discussion section, incorporating the contributions of all listed institutions. The comparative analysis reveals that performance in the leaderboard was driven primarily by dataset fidelity rather than sheer volume. PKB secured the premier position by prioritizing dataset quality and alignment over size, utilizing a robust configuration of rank equals to 64 and a high alpha scaling factor value between 512-1024. This strategic approach outperformed pks, which followed in second place. In contrast, Politeknik Muadzam Shah (PMS) demonstrated a critical ceiling for data utility, noting that increasing the dataset beyond a 1,000-line threshold resulted in diminishing returns. These findings suggest that for LoRA-based adaptation, the information density of the training set is a more potent predictor of success than the total number of training instances. Further examination of the lower-ranking institutions highlights the risks of over-parameterization and marginal engineering gains. Politeknik Banting Selangor (PBS) achieved the fourth position, reporting that dataset engineering techniques provided only marginal

improvements within their tested scope, despite experimenting with varied learning rates with values between 0.00001 and 0.0005) and low rank value is 8. Most notably, Kolej Komuniti Temerloh (KKTM) occupied the fifth position despite employing the most aggressive parameters, including the highest rank equals to 256 and the largest dataset of 2,700 lines. This discrepancy suggests that excessive rank dimensions and unrefined data volume may introduce noise, ultimately hindering the model's ability to generalize effectively compared to the leaner, alignment-focused strategies of top-tier performers like PKB.

V. CONCLUDING REMARKS

We are grateful to the organizers, participants, and reviewers whose dedication made this event a success. The findings presented in this special issue indicate that when fine-tuning strategies are resource-optimized and data-driven, LLMs can effectively serve complex domain-specific needs. We trust these papers will inspire a broad array of future research in the pursuit of more reliable and context-aware AI systems.

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.

REFERENCES

- [1] K. M. Kelly, B. Wu, and C. Liebe, "Quack the Code: A Computer Game Show Offers Learning Through Teaching AI in Undergraduate Software Engineering," in Proc. Annu. Conf. Innovation and Technol. Comput. Sci. Educ. (ITiCSE), 2025.
- [2] S. Murthy, "Celebrating the final AWS DeepRacer League championship and road ahead," AWS Machine Learning Blog, Aug. 29, 2024. [Online]. Available: <https://aws.amazon.com/blogs/machine-learning/celebrating-the-final-aws-deepracer-league-championship-and-road-ahead/> (accessed May 12, 2026).
- [3] Amazon Web Services, "AWS collaborates with Malaysia Ministry of Higher Education's JPPKK to build next-generation AI workforce," AWS Public Sector Blog, Oct. 1, 2024. [Online]. Available: <https://aws.amazon.com/blogs/publicsector/aws-collaborates-with-malaysia-ministry-of-higher-educations-jppkk-to-build-next-generation-ai-workforce/>. [Accessed: May 15, 2026].
- [4] Ministry of Digital, "Two years of the Ministry of Digital: Driving the national digitalisation agenda towards an AI nation by 2030," Dec. 12, 2025. [Online]. Available: <https://www.digital.gov.my/en-GB/siaran/Dua-Tahun-Kementerian-Digital:Memacu-Agenda-Pendigitalan-Negara,-Ke-Arah-Negara-AI-2030>. [Accessed: May 15, 2026].
- [5] V. Oh and N. K. Idries, "Racing beyond DeepRacer: Debut of the AWS LLM League," AWS Machine Learning Blog, Apr. 11, 2025. [Online]. Available: <https://aws.amazon.com/blogs/machine-learning/racing-beyond-deepracer-debut-of-the-aws-llm-league/> (accessed May 12, 2026).
- [6] Nadăș, M., Dioșan, L., & Tomescu, A. (2025). Synthetic data generation using large language models: Advances in text and code. IEEE Access.
- [7] Goyal, M., & Mahmoud, Q. H. (2024). A systematic review of synthetic data generation techniques using generative AI. Electronics, 13(17), 3509.
- [8] Fedorov, V., Bilyy, R., Toroshanko, O., & Zhykhareva, Y. (2025, September). The Use of LLMs for Synthesis Data for Small Datasets. In 2025 15th International Conference on Advanced Computer Information Technologies (ACIT) (pp. 959-964). IEEE.