

LLM-BASED INTELLIGENT SYSTEM FOR QUALITY CONTROL OF STUDENT ASSIGNMENTS

Dmitry Pakutnik

Department of Intelligent Information Technologies,
Belarusian State University of Informatics and Radioelectronics
Minsk, Belarus
E-mail: d.pakutnik@bsuir.by

This work explores the potential of large language models for automated quality control of student assignments through a human-in-the-loop approach. The proposed LLM-based intelligent system uses the generative abilities of language models to interact with experts and refine evaluation criteria dynamically. This adaptive process enhances grading accuracy, consistency, and interpretability, bringing automated assessment closer to human-level quality.

INTRODUCTION

The rapid development of large language models has transformed many areas of education, including automated assessment and feedback generation. Traditional grading systems often rely on predefined rules or static algorithms, which limits their ability to evaluate open-ended and concept-based student responses. Such approaches frequently lack flexibility, interpretability, and alignment with pedagogical standards.

To overcome these limitations, this research explores the integration of large language models with expert feedback within a human-in-the-loop framework [1]. The approach seeks to create an adaptive mechanism for refining evaluation criteria and ensuring transparent, consistent grading decisions. By combining the reasoning capacity of language models with human expertise, the study aims to enhance both the automation and pedagogical quality of assessment, contributing to the advancement of intelligent educational technologies.

I. RELATED WORK

Automated assessment technologies have progressed from early rule-based and feature-engineered systems to modern transformer architectures capable of semantic understanding. Models such as BERT and RoBERTa improved the accuracy of short-answer grading but still rely on predefined rubrics and lack adaptability when facing diverse student responses [2]. These limitations restrict their application in real educational environments where interpretability and consistency are essential.

The emergence of large language models has significantly expanded the possibilities of automated evaluation. LLMs demonstrate advanced reasoning and contextual comprehension, enabling more flexible grading that considers both the content and intent of student responses [3]. However, fully automated systems often misinterpret ambiguous instructions or domain-specific criteria, resulting in inconsistent evaluations.

To address these challenges, recent research explores **human-in-the-loop (HITL)** approaches

that integrate expert feedback into the grading process [4]. Studies such as *GradeHITL* show that human-guided refinement of grading rubrics increases reliability, transparency, and alignment with pedagogical goals. These developments highlight the potential of combining LLM-based reasoning with expert guidance to create adaptive and interpretable assessment frameworks — an idea further explored in the following section.

II. PROPOSED SYSTEM

The proposed LLM-Based Intelligent System for Quality Control of Student Assignments extends the human-in-the-loop framework to educational quality control by integrating LLM-based semantic reasoning, expert knowledge, and reinforcement learning for adaptive rubric optimization. This design enables both high-level automation and continuous refinement of grading criteria based on expert feedback.

The automated evaluation process is formulated as a classification function:

$$\hat{y} = F_{\theta}(G \parallel a),$$

where G is the grading rubric, a is the student's answer, θ are the LLM parameters, and $\parallel$ denotes text concatenation. The system assigns each answer a to one of the categorical score classes $\{c_i \mid i = 1, \dots, C\}$, where C represents the total number of categories.

Unlike traditional fine-tuning of model parameters θ , which is computationally expensive and requires large labeled datasets, this approach focuses on **rubric optimization** through human feedback. The rubric G is iteratively refined into a more accurate and interpretable version G^* , defined as:

$$G^* = \underset{G}{\operatorname{argmax}} \mathbb{E}_{a \in A} [\operatorname{Eval}(F_{\theta}(G \parallel a), y)]$$

where *Eval* measures the agreement between model predictions and expert-assigned scores. An overview of the framework is illustrated in Figure 1.

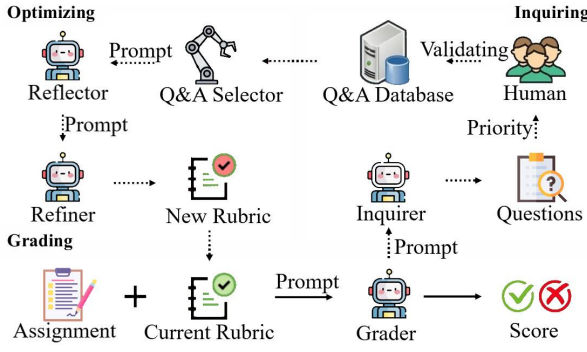


Figure 1 – Automated Grading with Human-in-the-Loop

Model Architecture The system comprises three interacting components that form a continuous feedback loop:

1. **Grader** – the LLM-based core evaluator responsible for analyzing student responses according to the current rubric. It employs *chain-of-thought* (CoT) reasoning to produce both intermediate explanations and final grades. This improves transparency and identifies ambiguities or inconsistencies within the rubric.
2. **Inquirer** – a generative component that detects rubric ambiguities or low-confidence grading cases and formulates targeted clarification questions. These questions are ranked using an uncertainty metric and filtered through a reinforcement learning policy π_θ that optimizes the selection of the most informative expert interactions according to the reward function:

$$r = R(F(G, x, h) \mid x) = \text{Eval}(\hat{y}, y)$$

where h is a question–answer pair from the expert database. A higher reward r indicates improved grading accuracy and rubric consistency. The reinforcement learning workflow used for Q&A selection and feedback optimization is illustrated in Figure 2.

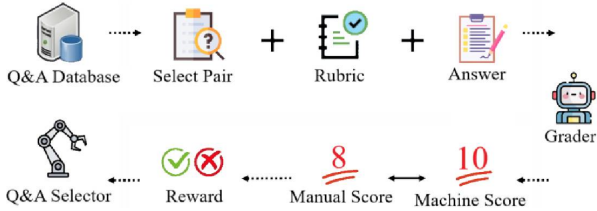


Figure 2 – Illustration of the RL-training process for Q&A selection and feedback optimization

3. **Optimizer** – refines the rubric by incorporating validated human–AI interactions. It operates as a multi-agent mechanism consisting of the *Retriever*, *Reflector*, and *Refiner*, which iteratively enhance the rubric G_t

through gradient-like updates:

$$G_t = G_{t-1} + \Delta G$$

where $\Delta G = f(\text{Expert Q\&A}, A_{\text{error}})$

This iterative optimization process ensures that each refinement step integrates expert insights into the grading logic, improving rubric clarity and adaptability.

Through the interaction of these components, the system dynamically aligns computational evaluation with expert judgment. The result is a grading framework that combines automation and human oversight to achieve consistent, interpretable, and pedagogically sound quality control of student assignments.

III. CONCLUSION

This work outlines an approach to integrating large language models with expert feedback to enhance the quality and consistency of automated assessment. The proposed human-in-the-loop mechanism is intended to support dynamic refinement of grading rubrics, improving transparency and alignment with pedagogical standards [5].

By continuously incorporating expert insights, such a framework could evolve from a static evaluator into an adaptive learning system capable of resolving ambiguities and maintaining reliable performance across diverse educational contexts. The approach contributes to advancing automated assessment toward greater interpretability, fairness, and instructional value.

IV. REFERENCES

1. Chu Y., Li H., Yang K., Shomer H., Liu H., Copur-Gençtürk Y., Tang J. A LLM-powered automatic grading framework with human-level guidelines optimization [Electronic resource] // arXiv. – 2024. – arXiv:2410.02165. – Mode of access: <https://arxiv.org/abs/2410.02165>. – Date of access: 23.10.2025
2. Devlin J. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding [Electronic resource] // arXiv, 2018. – Mode of access: <https://arxiv.org/abs/1810.04805>. – Date of access: 23.10.2025
3. Xie W., et al. Grade like a human: Rethinking automated assessment with large language models [Electronic resource] // arXiv. – 2024. – arXiv:2405.19694. – Mode of access: <https://arxiv.org/abs/2405.19694>. – Date of access: 23.10.2025
4. Wu X., Xiao L., Sun Y., Zhang J., Ma T., He L. A survey of human-in-the-loop for machine learning // Future Generation Computer Systems. – 2022. – Vol. 135. – P. 364–381
5. an L., Sha L., Zhao L., Li Y., Martinez-Maldonado R., Chen G., Li X., Jin Y., Gašević D. Practical and ethical challenges of large language models in education: a systematic scoping review // British Journal of Educational Technology. – 2024. – Vol. 55, № 1. – P. 90–112. – Mode of access: <https://bera-journals.onlinelibrary.wiley.com/journal/14678535>. – Date of access: 23.10.2025