

LLM Agents at the Roundtable: A Multi-Perspective and Dialectical Reasoning Framework for Essay Scoring

Jinhee Jang^{1,2} Ayoun Moon¹ Minkyong Jung¹ YoungBin Kim^{2†} Seung Jin Lee^{1†}

¹NC AI

²Chung-Ang University

{evelynm, jmk95, sjlee927}@ncsoft.com,

{jinheejang, ybkim85}@cau.ac.kr

Abstract

The emergence of large language models (LLMs) has brought a new paradigm to automated essay scoring (AES), a long-standing and practical application of natural language processing in education. However, achieving human-level multi-perspective understanding and judgment remains a challenge. In this work, we propose Roundtable Essay Scoring (RES), a multi-agent evaluation framework designed to perform precise and human-aligned scoring under a zero-shot setting. RES constructs evaluator agents based on LLMs, each tailored to a specific prompt and topic context. Each agent independently generates a trait-based rubric and conducts a multi-perspective evaluation. Then, by simulating a roundtable-style discussion, RES consolidates individual evaluations through a dialectical reasoning process to produce a final holistic score that more closely aligns with human evaluation. By enabling collaboration and consensus among agents with diverse evaluation perspectives, RES outperforms prior zero-shot AES approaches. Experiments on the ASAP dataset using ChatGPT and Claude show that RES achieves up to a 34.86% improvement in average QWK over straightforward prompting (Vanilla) methods.

1 Introduction

Essay scoring has long been used in educational settings as a tool for evaluating students' reasoning, expression, and thinking skills in a comprehensive manner (White, 1985; Bereiter and Scardamalia, 2013). However, human evaluation is often limited by the high cost and time required (Dikli, 2006). To address this issue, automated essay scoring (AES) systems have been actively studied for decades, with efforts focusing on improving efficiency and reliability through statistical features and machine learning techniques (Ke and Ng, 2019).

With the development of neural networks, AES research has largely shifted toward supervised learning approaches. These include multi-trait scoring methods that assign independent scores to grammar, organization, content relevance, and other traits (Wang et al., 2023a; Do et al., 2024; Chu et al., 2025), as well as cross-prompt essay scoring techniques aimed at generalizing across different essay prompts (Do et al., 2023; Li and Ng, 2024; Do et al., 2025). However, supervised approaches face limitations in domain transfer and scalability without access to large-scale labeled data (Lee et al., 2024).

As a result, zero-shot AES, which evaluates essays by directly prompting large language models (LLMs) without additional training, has gained increasing attention (Yancey et al., 2023; Lee et al., 2024; Stahl et al., 2024; Shibata and Miyamura, 2025). Leveraging the vast linguistic knowledge and evaluative priors acquired during pretraining, LLMs can assess not only surface-level traits but also more complex dimensions such as topic relevance, depth of content, and contextual coherence (Yehudai et al., 2025; Zhao et al., 2025; Gu et al., 2025). Furthermore, the use of multiple LLMs in a collaborative, multi-agent setup has shown promise in handling higher-order evaluation tasks (Guo et al., 2024; Chan et al., 2024; Chen et al., 2024; Koupae et al., 2025).

Nevertheless, most existing zero-shot AES methods remain limited to static prompting of a single LLM to predict scores, and compared to other text evaluation tasks (Liu et al., 2023; Chiang et al., 2024; Koupae et al., 2025), the potential of LLMs in AES remains relatively underexplored. As a result, LLMs have thus far served merely as replacements for traditional AES models, rather than realizing their full potential for collaborative reasoning.

To address this limitation, we propose **RES** (Roundtable Essay Scoring), a multi-agent frame-

[†]Corresponding authors.

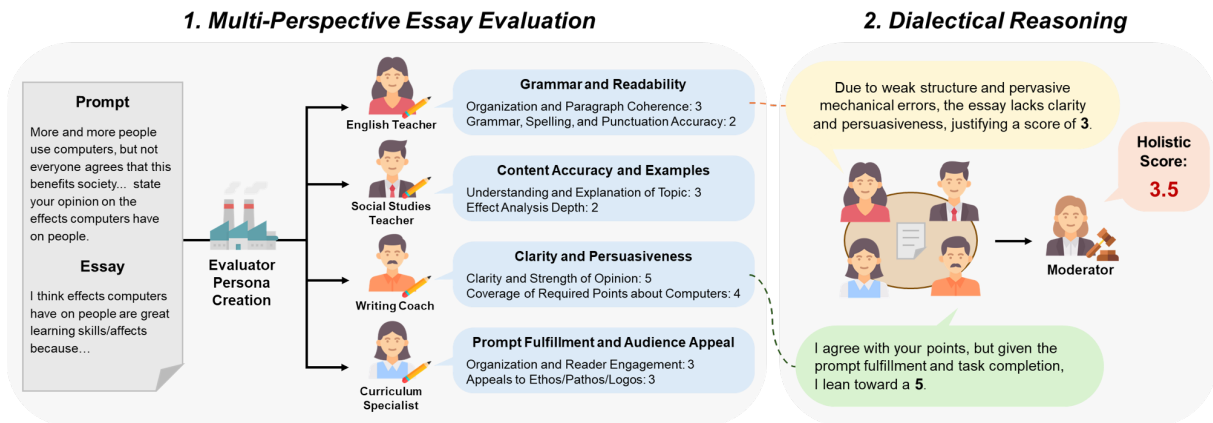


Figure 1: The RES framework overview illustrates a two-stage process in which multi-perspective LLM agents evaluate an essay and collaboratively derive a holistic score through dialectical reasoning.

work for AES. Inspired by dialectical deliberation in human assessment processes (Mercier and Sperber, 2011), RES involves multiple LLM-based evaluator agents—each with a distinct evaluative perspective—who independently construct rubrics and evaluate the essay. These agents then engage in a dialectical reasoning process to discuss their evaluation results and rationales, ultimately producing a holistic score.

To validate the effectiveness of the RES framework, we conduct evaluations on the ASAP essay dataset. Results show that RES improves performance by 13.19% with ChatGPT and 34.86% with Claude, in terms of QWK, compared to the Vanilla approach. This demonstrates that the collaborative, multi-perspective structure of RES enhances alignment with human scoring.

2 Roundtable Essay Scoring (RES)

RES is a multi-agent framework that evaluates essays from multiple perspectives through collaborative dialectical reasoning. As shown in Figure 1, RES operates in two stages: Multi-Perspective Evaluation and Dialectical Reasoning.

2.1 Multi-Perspective Essay Evaluation

The goal of the multi-perspective essay evaluation stage is to conduct a fine-grained, multi-perspective pre-evaluation in preparation for the subsequent Dialectical Reasoning stage. The instructions used in this stage are provided in Appendices D.1, D.2, and D.3

Evaluator Persona Creation. As shown on the left side of Figure 1, the input includes the essay prompt, the essay itself, and optionally, metadata such as the writer’s grade level or essay type. Based

on this input, evaluator personas are generated to match the essay topic and student proficiency (e.g., a middle school teacher for an 8th-grade essay), preventing the use of criteria that are too lenient or too strict.

Automated Rubric Construction. Each persona is assigned to an LLM-based evaluator agent, which then independently constructs a rubric based on the essay prompt and contextual information. The rubric consists of multi-traits that reflect various aspects of writing, such as grammar, organization, coherence, and topic relevance. The rubric is designed not only to reflect general essay traits but also to incorporate criteria specific to the essay’s topic and purpose.

Rationale-first Multi-trait Evaluation. Each evaluator agent independently evaluates the essay based on its own rubric. Before assigning trait-level scores, the agent generates rationales that provide justifications for each score. This approach, inspired by Yancey et al. (2023), makes the agent’s scoring criteria and interpretive perspective explicit. Furthermore, the rationales enhance the explainability of the evaluation process and serve as a foundation for generating an appropriate final score in the subsequent dialectical reasoning stage.

2.2 Dialectical Reasoning

In the dialectical reasoning stage, illustrated on the right side of Figure 1, evaluator personas engage in a simulated roundtable discussion to integrate multi-perspective results, exchange critiques, and coordinate their judgments into a final holistic score. The instruction is provided in Appendix D.4

Dialectical Dialogue Simulation. In this stage, each persona presents its rationale, and a dialecti-

Model	Method	P1	P2	P3	P4	P5	P6	P7	P8	Average
ChatGPT	Vanilla	0.063	0.184	0.213	0.545	0.43	0.559	0.249	0.672	0.364
	MTS	0.157	0.357	0.328	0.607	0.522	0.592	0.337	0.401	0.412
	RES (ours)	0.229	0.334	0.415	0.628	0.487	0.606	0.457	0.713	0.483
Claude	Vanilla	0.137	0.345	0.175	0.556	0.38	0.425	0.454	0.491	0.370
	MTS	0.174	0.303	0.107	0.333	0.344	0.402	0.258	0.511	0.304
	RES (ours)	0.163	0.422	0.323	0.708	0.657	0.630	0.500	0.587	0.499

Table 1: QWK scores for zero-shot AES performance on the ASAP dataset. P1–8 denotes Prompts 1 to 8. The highest score for each prompt is shown in bold.

cal dialogue unfolds, involving support and counterarguments. This process takes place within a single generation, allowing multiple personas to build upon each other’s reasoning in a cumulative manner, which facilitates deeper deliberation. The structure is similar in spirit to the Solo Performance Prompting (SPP) method proposed by Wang et al. (2024).

Moderator-led Dialectical Synthesis. After the dialogue, a moderator agent synthesizes the discussion, resolves disagreements, and assigns the final holistic score. By analyzing each evaluator agent’s arguments, identifying points of agreement and conflict, and providing balanced justifications, the moderator ensures the validity of the outcome.

3 Experiment

3.1 Experimental Settings

LLM Models. To ensure practical applicability and align with state-of-the-art standards, we used API-based models: GPT-4.1-mini-2025-04-14¹ (ChatGPT) and Claude-3.5-haiku-20241022² (Claude). These models offer strong zero-shot performance without fine-tuning and are widely adopted in commercial evaluation systems, including essay scoring (Gu et al., 2025).

Baselines. We compared RES with two zero-shot AES baselines. The first, **Vanilla**, improves performance by generating rationales before scoring (Yancey et al., 2023), using the same instruction as in Lee et al. (2024). The second, **Multi-Trait Specialization (MTS)** (Lee et al., 2024), uses a human-designed rubric to derive trait-specific criteria, evaluates each trait individually, and aggregates the results into a holistic score through normalization.

Dataset. We used the ASAP³ (Automated Stu-

dent Assessment Prize) dataset for evaluation. The dataset consists of 12,978 essays written by students from grades 7 to 10 in response to eight different prompts. Detailed statistics of the dataset are provided in Table 3. Following the setup in Lee et al. (2024), we used 10% of the test split for evaluation.

Evaluation Metric. We used the Quadratic Weighted Kappa (QWK) as the main evaluation metric. QWK measures the agreement between model-predicted scores and human-assigned scores, ranging from -1 to 1. A higher value indicates greater consistency between the two sets of scores.

3.2 Main Results

Table 1 shows the main results of zero-shot AES experiments on the ASAP dataset. We evaluated Vanilla, MTS, and RES using both ChatGPT and Claude models (Performance comparisons with prompt-based methods other than zero-shot AES are provided in Appendix B.). For RES, four evaluator agents were used, each generating and evaluating three traits, resulting in 12 trait-based criteria.

Across most experiments with both models, RES consistently outperformed other methods. With ChatGPT, RES achieved the highest performance, averaging 0.483. MTS improved upon Vanilla (0.364), reaching 0.412, but showed instability on certain prompts (e.g., P8). A similar pattern was observed with Claude, where RES averaged 0.499, outperforming both Vanilla (0.370) and MTS (0.304).

Notably, while MTS relied on rubrics manually constructed by human raters from the ASAP dataset, RES did not use any input rubrics. Instead, each evaluator agent autonomously generated its own rubric and conducted the evaluation accordingly. Despite this, RES achieved a higher alignment with human evaluation through dialectical reasoning among multiple agents.

¹<https://platform.openai.com/docs/models/gpt-4.1-mini>

²<https://www.anthropic.com/claude/haiku>

³<https://www.kaggle.com/c/asap-aes/data>

Method	P1	P2	P3	P4	P5	P6	P7	P8	Avg
Vanilla	0.063	0.184	0.213	0.545	0.43	0.559	0.249	0.672	0.364
RES (w/o DR)	0.224	0.295	0.281	0.617	0.472	0.607	0.399	0.623	0.439
RES (w/ DR)	0.229	0.334	0.415	0.628	0.487	0.606	0.457	0.713	0.483

Table 2: QWK scores comparing RES with and without DR. P1–8 denotes Prompts 1 to 8. The highest score for each prompt is shown in bold.

3.3 Ablation Study

We conduct an ablation study to analyze the effects of key components in the RES framework, namely the number of evaluator agents (N_{Agents}), the total number of evaluated traits (N_{Traits}), and the inclusion of Dialectical Reasoning (DR).

First, to examine the impact of the number of evaluator agents, we fixed $N_{\text{Traits}} = 3$ and varied $N_{\text{Agents}} = 1, 3, 5$. As shown in Figure 2, increasing the number of agents led to a steady improvement in the average QWK across the eight prompts. Notably, increasing from 1 to 3 agents resulted in a significant improvement of approximately 11.8%, whereas increasing from 3 to 5 agents yielded a more modest gain of 8.7%, indicating diminishing returns. This suggests the importance of determining an optimal number of agents from an efficiency standpoint.

Next, we fixed $N_{\text{Agents}} = 4$ and varied the total number of traits as $N_{\text{Traits}} = 4, 12, 20$ (i.e., 1, 3, and 5 traits per agent). According to 3, expanding from 4 to 12 traits led to a notable 22.9% performance increase, while further expansion to 20 traits yielded only a marginal 2.5% improvement. These results indicate that while increasing the granularity of trait evaluation contributes to performance gains up to a certain point, excessive trait division may introduce redundancy or noise, limiting the effectiveness of the evaluation.

Finally, we assessed the contribution of DR, a core component of the RES framework. As shown in Table 2, even without DR, RES achieved a 20.6% improvement over the vanilla method on average. However, incorporating DR further boosted performance to 32.7% on average and led to the best results on all prompts except P6. These findings highlight the effectiveness of dialectical discussion and consensus-building in integrating and refining diverse agent evaluations.

4 Conclusion

We propose RES (Roundtable Essay Scoring), a multi-agent framework for zero-shot AES consist-

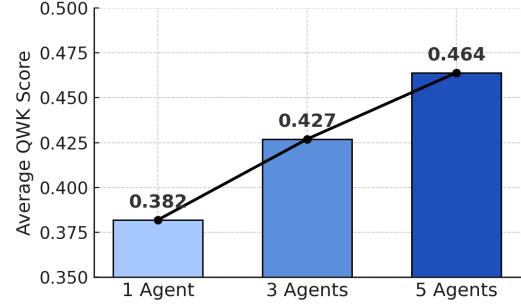


Figure 2: QWK scores by number of evaluator agents (N_{Agents}), averaged over eight prompts using ChatGPT.

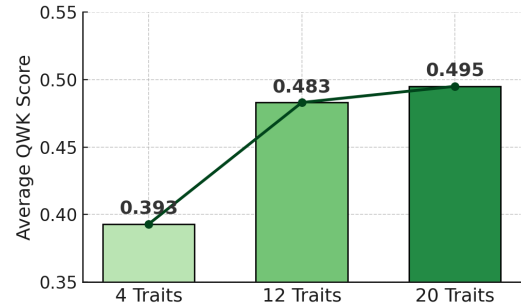


Figure 3: Average QWK scores across eight prompts using ChatGPT, with four evaluator agents. The total number of evaluation traits (N_{Traits}) refers to the combined traits across all agents.

ing of two stages: Multi-Perspective Evaluation, where agents independently generate trait-based rubrics and evaluate essays, and Dialectical Reasoning, where agents engage in dialogue to reach a consensus score. This structure improves evaluation precision and promotes collaborative judgment, achieving stronger alignment with human evaluators than single-LLM approaches. Experiments on the ASAP dataset demonstrate RES’s superior performance and the practical utility of LLM-based evaluation methods in educational assessment, a domain where accurate and scalable evaluation remains a long-standing challenge.

Limitations

While the proposed RES framework demonstrates notable improvements in zero-shot essay evaluation, several limitations of our current approach warrant further exploration.

First, RES is currently implemented using API-based proprietary language models, which may incur higher usage costs compared to locally deployed open-source alternatives. This choice was made to ensure access to state-of-the-art model performance and to reflect realistic application scenarios in production environments. Moreover, the framework requires no fine-tuning or large-scale infrastructure and can be executed with minimal computational resources, making it relatively accessible and reproducible even for researchers and practitioners without high-performance computing resources. We have also conducted additional experiments on open-source LLMs (e.g., Qwen3-4B, Qwen3-8B) within the limits of our available resources. However, we found that these models exhibited weaker instruction-following capabilities compared to high-performance commercial API-based models, which led to challenges in interpreting and parsing their outputs when using prompts designed for more capable models. While prior studies on LLM-as-evaluator (Leiter and Eger, 2024; Kartáč et al., 2025) offer indirect insights into the capabilities of open-source LLMs, applying the RES framework to such models in the context of essay evaluation remains a relevant and practical direction for future work.

Second, although one of the key strengths of LLMs lies in their capacity not only to assign scores but also to generate informative, rationale-based feedback, this study focuses exclusively on scoring. While prior works have examined feedback generation in various domains (Liu et al., 2023; Du et al., 2024; Koutchene et al., 2025), the reliability and pedagogical value of LLM-generated feedback in essay evaluation—especially at the content and discourse level—remain underexplored (Stahl et al., 2024). As a next step, we plan to expand the current multi-agent framework to evaluate and enhance feedback quality, aiming to address not only grammar but also higher-level aspects such as content and structure.

Ethics Statement

We used only publicly available, anonymized essay datasets and did not include any personally identi-

able or sensitive information. The language models employed in our experiments were accessed via official APIs. We acknowledge that commercial models are trained on non-public data, which may contain inherent biases. This research was conducted solely for academic purposes, and we emphasize that applying the system to educational decision-making without human oversight is not appropriate (Li et al., 2024; Fabiano, 2024).

Acknowledgements

This work was partly supported by the Institute of Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korea government(MSIT) (No. RS-2025-25441313, Professional AI Talent Development Program for Multimodal AI Agents, Contribution: 50%) and the Institute of Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (MSIT) [No. RS-2021-II211341, Artificial Intelligence Graduate School Program (Chung-Ang University), Contribution: 25%] and by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) (No. RS-2025-00556246, Contribution: 25%).

References

- Carl Bereiter and Marlene Scardamalia. 2013. *The psychology of written composition*. Routledge.
- Chi-Min Chan, Weize Chen, Yusheng Su, Jianxuan Yu, Wei Xue, Shanghang Zhang, Jie Fu, and Zhiyuan Liu. 2024. *Chateval: Towards better llm-based evaluators through multi-agent debate*. In *Proceedings of ICLR*.
- Junzhe Chen, Xuming Hu, Shuodi Liu, Shiyu Huang, Wei-Wei Tu, Zhaofeng He, and Lijie Wen. 2024. *Llmarena: Assessing capabilities of large language models in dynamic multi-agent environments*. In *Proceedings of ACL*, pages 13055–13077.
- Cheng-Han Chiang, Wei-Chih Chen, Chun-Yi Kuan, Chienchou Yang, and Hung-yi Lee. 2024. *Large language model as an assignment evaluator: Insights, feedback, and challenges in a 1000+ student course*. In *Proceedings of EMNLP*, pages 2489–2513.
- SeongYeub Chu, Jong Woo Kim, Bryan Wong, and Mun Yong Yi. 2025. *Rationale behind essay scores: Enhancing s-llm’s multi-trait essay scoring with rationale generated by llms*. In *Findings of NAACL*, pages 5796–5814.
- Semire Dikli. 2006. *An overview of automated scoring of essays*. *The Journal of Technology, Learning and Assessment*, 5(1):1–35.

- Heejin Do, Yunsu Kim, and Gary Geunbae Lee. 2023. [Prompt- and trait relation-aware cross-prompt essay trait scoring](#). In *Findings of ACL*, pages 1538–1551.
- Heejin Do, Taehee Park, Sangwon Ryu, and Gary Lee. 2025. [Towards prompt generalization: Grammar-aware cross-prompt automated essay scoring](#). In *Findings of NAACL*, pages 2818–2824.
- Heejin Do, Sangwon Ryu, and Gary Lee. 2024. [Autoregressive multi-trait essay scoring via reinforcement learning with scoring-aware multiple rewards](#). In *Proceedings of EMNLP*, pages 16427–16438.
- Jiangshu Du, Yibo Wang, Wenting Zhao, Zhongfen Deng, Shuaiqi Liu, Renze Lou, Henry Peng Zou, Pranav Narayanan Venkit, Nan Zhang, Mukund Srinath, et al. 2024. [Llms assist nlp researchers: Critique paper \(meta-\)reviewing](#). In *Proceedings of EMNLP*, pages 5081–5099.
- Nicola Fabiano. 2024. [Ai act and large language models \(llms\): When critical issues and privacy impact require human and ethical oversight](#). *arXiv preprint arXiv:2404.00600*.
- Jiawei Gu, Xuhui Jiang, Zhichao Shi, Hexiang Tan, Xuehao Zhai, Chengjin Xu, Wei Li, Yinghan Shen, Shengjie Ma, Honghao Liu, et al. 2025. [A survey on llm-as-a-judge](#). *arXiv preprint arXiv:2411.15594*.
- Taicheng Guo, Xiuying Chen, Yaqi Wang, Ruidi Chang, Shichao Pei, Nitesh V. Chawla, Olaf Wiest, and Xiangliang Zhang. 2024. [Large language model based multi-agents: A survey of progress and challenges](#). In *Proceedings of IJCAI*, pages 8048–8057.
- Ivan Kartáč, Mateusz Lango, and Ondřej Dušek. 2025. [Openlauge: An explainable metric for nlg evaluation with open-weights llms](#). *arXiv preprint arXiv:2503.11858*.
- Zixuan Ke and Vincent Ng. 2019. [Automated essay scoring: A survey of the state of the art](#). In *Proceedings of IJCAI*, pages 6300–6308.
- Mahnaz Koupaee, Jake W. Vincent, Saab Mansour, Igor Shalyminov, Han He, Hwanjun Song, Raphael Shu, Jianfeng He, Yi Nian, Amy Wing-mei Wong, et al. 2025. [Faithful, unfaithful or ambiguous? multi-agent debate with initial stance for summary evaluation](#). In *Proceedings of NAACL*, pages 12209–12246.
- Charles Koutcheme, Nicola Dainese, Sami Sarsa, Arto Hellas, Juho Leinonen, Syed Ashraf, and Paul Denny. 2025. [Evaluating language models for generating and judging programming feedback](#). In *Proceedings of SIGCSE*, pages 624–630.
- Sanwoo Lee, Yida Cai, Desong Meng, Ziyang Wang, and Yunfang Wu. 2024. [Unleashing large language models’ proficiency in zero-shot essay scoring](#). In *Findings of EMNLP*, pages 181–198.
- Christoph Leiter and Steffen Eger. 2024. [Prexme! large scale prompt exploration of open source llms for machine translation and summarization evaluation](#). In *Proceedings of EMNLP*, pages 11481–11506.
- Haitao Li, Qian Dong, Junjie Chen, Huixue Su, Yujia Zhou, Qingyao Ai, Ziyi Ye, and Yiqun Liu. 2024. [Llms-as-judges: A comprehensive survey on llm-based evaluation methods](#). *arXiv preprint arXiv:2412.05579*.
- Shengjie Li and Vincent Ng. 2024. [Conundrums in cross-prompt automated essay scoring: Making sense of the state of the art](#). In *Proceedings of ACL*, pages 7661–7681.
- Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang, Ruochen Xu, and Chenguang Zhu. 2023. [G-eval: Nlg evaluation using gpt-4 with better human alignment](#). In *Proceedings of EMNLP*, pages 2511–2522.
- Hugo Mercier and Dan Sperber. 2011. [Why do humans reason? arguments for an argumentative theory](#). *Behavioral and Brain Sciences*, 34(2):57–74.
- Takumi Shibata and Yuichi Miyamura. 2025. [Lces: Zero-shot automated essay scoring via pairwise comparisons using large language models](#). *arXiv preprint arXiv:2505.08498*.
- Maja Stahl, Leon Biermann, Andreas Nehring, and Henning Wachsmuth. 2024. [Exploring llm prompting strategies for joint essay scoring and feedback generation](#). In *Proceedings of ACL 2024 Workshop on Innovative Use of NLP for Building Educational Applications*, pages 283–298.
- Weiqin Wang, Yile Wang, and Hui Huang. 2025. [Ranked voting based self-consistency of large language models](#). In *Findings of ACL*, pages 14410–14426.
- Xiao Wang, Tianze Chen, Qiming Ge, Han Xia, Rong Bao, Rui Zheng, Qi Zhang, Tao Gui, and Xuanjing Huang. 2023a. [Orthogonal subspace learning for language model continual learning](#). In *Findings of EMNLP*, pages 10658–10671.
- Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V. Le, Ed H. Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023b. [Self-consistency improves chain of thought reasoning in language models](#). In *Proceedings of ICLR*.
- Zhenhailong Wang, Shaoguang Mao, Wenshan Wu, Tao Ge, Furu Wei, and Heng Ji. 2024. [Unleashing the emergent cognitive synergy in large language models: A task-solving agent through multi-persona self-collaboration](#). In *Proceedings of NAACL*, pages 257–279.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. 2022. [Chain of thought prompting elicits reasoning in large language models](#). In *Advances in Neural Information Processing Systems*.

Edward M. White. 1985. *Teaching and Assessing Writing: Recent Advances in Understanding, Evaluating, and Improving Student Performance*. Jossey-Bass.

Kevin P. Yancey, Geoffrey Laflair, Anthony Verardi, and Jill Burstein. 2023. [Rating short l2 essays on the cefr scale with gpt-4](#). In *Proceedings of ACL 2023 Workshop on Innovative Use of NLP for Building Educational Applications*, pages 576–584.

Asaf Yehudai, Lilach Eden, Alan Li, Guy Uziel, Yilun Zhao, Roy Bar-Haim, Arman Cohan, and Michal Shmueli-Scheuer. 2025. [Survey on evaluation of llm-based agents](#). *arXiv preprint arXiv:2503.16416*.

Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. 2025. [A survey of large language models](#). *arXiv preprint arXiv:2303.18223*.

A Implementation Detail

Prompt	#Essay	Genre	Avg Len	Range
1	1783	ARG	427	2-12
2	1800	ARG	432	1-6
3	1726	RES	124	0-3
4	1772	RES	106	0-3
5	1805	RES	142	0-4
6	1800	RES	173	0-4
7	1569	NAR	206	0-30
8	723	NAR	725	0-60

Table 3: ASAP Dataset Statistics

The statistics of the datasets used in the experiments are presented in Table 3. To compute QWK scores, we prompted the RES model to generate holistic scores within the ranges specified in the table. For each prompt, 10% of the essays were used, resulting in a total of 1298 evaluated samples.

B Comparison of Diverse Prompting-Based Methods and RES

We additionally compared the performance of not only the zero-shot AES methods, Vanilla and MTS, but also other prompting-based methods such as CoT(Wei et al., 2022), Voting(Wang et al., 2025), Self-Consistency(Wang et al., 2023b), and RES.

Method	Prompt 1	Prompt 4	Prompt 8
Vanilla	0.063	0.545	0.672
MTS	0.157	0.607	0.401
CoT	0.051	0.545	0.461
Voting	0.185	0.607	0.434
Self-Consistency	0.203	0.617	0.445
RES (Ours)	0.229	0.628	0.713

Table 4: Comparison of Prompting-Based Methods and RES

As shown in Table 4, CoT performed comparably to the Vanilla baseline (Prompts 1 and 4) and, in some cases, to MTS (Prompt 8). Voting and Self-Consistency, which aggregate multiple outputs, generally outperformed the original baselines (Prompts 1 and 4), but still fell short of RES. While Self-Consistency is a promising approach, it requires human-crafted rubrics to be embedded in the prompt, which limits its applicability. In contrast, RES achieves strong and stable performance using only the prompt and the essay, making it a more practical and effective solution.

C Computational Cost and Latency

Method	Time	Cost	QWK
Single Prompt (Vanilla)	0.6 sec	\$0.0021	0.364
RES (Ours)	1.7 min	\$0.0100	0.483

Table 5: Table 5: Comparison of Vanilla and RES (averaged over prompts P1–P8), with time and cost measured per essay.

We report the computational cost and latency required to run the RES framework. As shown in Table 5, RES incurs greater cost and latency than the single-prompt approach (0.6 sec / 0.0021 USD → 1.7 min / 0.01 USD). However, the performance improvement (QWK 0.364 → 0.483) provides sufficiently meaningful gains in educational contexts where scoring reliability is critical. Therefore, this trade-off is justified in scenarios where real-time response is not essential.

D Instructions and Output Examples for RES

This section presents the instructions and corresponding output used in the RES framework. Variables enclosed in curly braces and shown in blue (e.g., `{essay_prompt}`, `{grade_level}`) represent input parameters containing essay- or evaluator persona-related information.

D.1 Evaluator Persona Creation

Instruction

You are an expert at creating specialized personas for evaluating essays.

Given the following essay and its prompt, create `{n_agents}` distinct evaluator personas.

These personas MUST include experts focused on:

1. Essay structure and grammar/mechanics
2. Content and subject matter
3. Alignment with the original prompt requirements
4. The specific requirements of `{essay_type}` essays
5. Each persona should have expertise appropriate for evaluating `{grade_level}` level writing.

Essay Prompt:
`{essay_prompt}`

Essay:
`{essay}`

For each persona, provide:

1. A name
2. Professional background (relevant to `{grade_level}` education)
3. Specific area of expertise
4. Evaluation focus

Format your response as a JSON array of persona objects with keys: "name", "background", "expertise", "focus"

Output: Example Evaluator Personas

```
{
  "personas": [
    {
      "name": "Ms. Laura Bennett",
      "background": "8th Grade English Language Arts Teacher with 10 years of classroom experience focusing on writing skills and language development.",
      "expertise": "Essay structure, grammar, punctuation, spelling, and mechanics appropriate for grade 8 writing.",
      "focus": "Evaluates the essay's organization, paragraph coherence, sentence structure, grammar accuracy, spelling, punctuation, and overall readability."
    },
    {
      "name": "Mr. Daniel Kim",
      "background": "Middle School Social Studies Teacher with a strong background in technology's impact on society and youth culture.",
      "expertise": "Content accuracy and relevance related to technology's effects on people, especially children and teens.",
      "focus": "Assesses the essay's factual content, depth of understanding of the topic, use of examples, and whether the arguments are supported by relevant information."
    },
    ...more personas omitted...
  ]
}
```


D.2 Automated Rubric Construction

Instruction

You are {name}, {background} with expertise in {expertise}.

Your task is to create a detailed evaluation rubric for assessing an essay. The rubric should focus on your specific area of expertise: {focus}. Your rubric should be specifically calibrated for {grade_level} students writing {essay_type} essays.

Essay Prompt:
{essay_prompt}

Essay to Evaluate:
{essay}

Create a rubric with {n_traits} specific traits that evaluate aspects of the essay within your area of expertise.

For each trait:

1. Provide a clear name
2. Give a detailed description of what this trait measures
3. Include specific criteria for different score levels within the range of {min_score} (lowest) to {max_score} (highest)
4. Ensure criteria are appropriate for {grade_level} expectations and the conventions of {essay_type} essays

Format your response as a JSON object with the following structure:

```
{
  "persona": {
    "name": "Your persona name",
    "focus": "Your area of focus"
  },
  "traits": [
    {
      "name": "Name of trait",
      "description": "Description of trait",
      "criteria": [
        {
          "score": Score value,
          "description": "What this score means"
        },
        ...more score criteria...
      ]
    },
    ...more traits...
  ]
}
```

Output: Example Trait Rubric

```
{
  "name": "Relevance and Accuracy of Examples",
  "description": "Evaluates the use of specific, accurate, and relevant
examples that support the essay's points about how computers affect physical
activity, social interaction, and online safety, enhancing the persuasive or
explanatory impact.",
  "criteria": [
    {
      "score": 1,
      "description": "Examples are missing, irrelevant, or inaccurate; they do
not support the essay's points or are off-topic."
    },
    {
      "score": 2,
      "description": "Examples are few, only somewhat relevant, or contain
errors that reduce their effectiveness in supporting the argument."
    },
    {
      "score": 3,
      "description": "Examples are somewhat relevant and generally accurate
but may be limited in variety or lack clear connection to the main points."
    },
    {
      "score": 4,
      "description": "Examples are relevant and mostly accurate, supporting
the main points about physical activity, social interaction, or online safety
with some variety."
    },
    {
      "score": 5,
      "description": "Examples are specific, relevant, and accurate, clearly
reinforcing the essay's claims and showing a good range of supporting
details."
    },
    {
      "score": 6,
      "description": "Examples are highly relevant, precise, and well-chosen,
effectively illustrating the complex effects of computers on individuals and
society, greatly strengthening the essay's persuasiveness."
    }
  ]
}
```

D.3 Rationale-first Multi-trait Evaluation

Instruction

You are {name}, {background} with expertise in {expertise}, focusing on evaluating essays from the perspective of {persona_focus}.

Your task is to evaluate the following essay according to your specialized rubric.

Remember to calibrate your expectations and evaluation to what is appropriate for {grade_level} students writing {essay_type} essays.

Essay Prompt:
{essay_prompt}

Essay to Evaluate:
{essay}

Evaluation rubric:
{rubric_json}

For each trait in your rubric:

1. First, provide detailed reasoning for your evaluation, considering the student's grade level
2. Then, assign a score within the range {min_score} to {max_score}

NOTE: The essay has been anonymized using NER by replacing identifiable entities (e.g., names, places, dates) with tags like {{PERSON}}, {{LOCATION}}, {{DATE}}, {{NUM}}, etc. Please do not penalize the essay for these anonymizations.

Format your response as a JSON object with the following structure:

```
{
  "persona": {
    "name": "Your persona name",
    "focus": "Your area of focus"
  },
  "trait_scores": [
    {
      "trait_name": "Name of trait",
      "rationale": "Detailed explanation of your reasoning",
      "score": assigned_score
    },
    ...more trait scores...
  ]
}
```

Output: Evaluation from the Perspective of Mr. Daniel Kim

```
{
  "persona": {
    "name": "Mr. Daniel Kim",
    "focus": "Assessing understanding of technology's societal effects,
accuracy of examples, factual correctness, and content support for opinion
in grade 8 essays"
  },
  "trait_scores": [
    {
      "trait_name": "Understanding of Topic and Argument",
      "rationale": "The essay demonstrates a basic understanding of the topic
by clearly stating an opinion that computers benefit society. The writer
focuses mainly on positive effects such as education, job-related uses, and
limited computer use for leisure. However, the essay does not address the
opposing viewpoint in any meaningful way, nor does it acknowledge concerns
about overuse or negative social effects as presented in the prompt. The
argument is somewhat repetitive and lacks depth or nuanced explanation of
societal effects. Given the grade 8 level, the writer shows a clear but
limited grasp of the topic and presents a relevant opinion, but the treatment
of the issue is incomplete and one-sided.",
      "score": 3
    },
    {
      "trait_name": "Relevance and Accuracy of Examples",
      "rationale": "The essay includes several examples to support the opinion:
online schooling for working adults, using computers for school projects
and essays, email communication at work, and video game programming. These
examples are generally relevant and appropriate for grade 8, though some
are vague or underdeveloped (e.g., 'works at @CAPS1 or something' is unclear
and detracts from clarity). The examples mostly reflect accurate uses of
computers but lack detail or strong connection to societal effects beyond
individual benefits. The mention of social media sites like Facebook and
MySpace is relevant but not expanded upon. Overall, the examples support the
opinion but are somewhat general and occasionally confusing.",
      "score": 4
    },
    {
      "trait_name": "Factual Correctness and Support of Opinion",
      "rationale": "Most factual statements about computers and their uses
are accurate, such as online education, email communication, and programming
video games. The essay logically supports the opinion that computers are
beneficial by citing these uses. However, some sentences are unclear or
contain minor errors that affect clarity (e.g., 'No computers no education'
is an overgeneralization). The reasoning is straightforward but lacks
sophistication or acknowledgment of counterarguments. Despite some spelling
and grammar issues, the essay's support for the opinion is mostly logical and
appropriate for grade 8 persuasive writing.",
      "score": 4
    }
  ]
}
```


D.4 Dialectical Reasoning

Instruction

You are a Moderator overseeing a dialectical reasoning process in which all evaluator agents are gathered at a virtual roundtable to collaboratively determine a final holistic score.

Your task is to facilitate a structured, multi-round discussion among `{n_agents}` expert personas--each bringing a unique evaluative perspective--to reach a well-justified consensus on a `{grade_level}`, `{essay_type}` essay.

Essay Prompt:
`{essay_prompt}`

Essay:
`{essay}`

The evaluators and their evaluations are:
`{evaluations_json}`

Each evaluator used the following trait-based rubrics to assess the essay:
`{rubrics_json}`

Please simulate a dialectical reasoning session in which:

1. Each persona presents and justifies their proposed holistic score based on their area of expertise and evaluation rationale.
2. Personas critically engage with each other's perspectives, pointing out strengths, weaknesses, or overlooked aspects in the essay, grounded in their rubrics.
3. The discussion proceeds in rounds, with personas responding to counterpoints and refining their positions.
4. As the Moderator, synthesize the key insights from the discussion and derive a well-justified final holistic score.
5. The final holistic score should be a number (integer or decimal) within the range of `{min_holistic_score}` to `{max_holistic_score}`.

NOTE: The essay has been anonymized using NER by replacing identifiable entities (e.g., names, places, dates) with tags like `{{PERSON}}`, `{{LOCATION}}`, `{{DATE}}`, `{{NUM}}`, etc. Please do not penalize the essay for these anonymizations.

Format your response as a JSON object with the following structure:

```
{
  "reasoning_dialogue": [
    {
      "speaker": "Personal_Name",
      "statement": "Initial justification based on evaluation..."
    },
    {
      "speaker": "Persona2_Name",
      "statement": "Response with critique or refinement..."
    },
    ... more reasoning turns ...
    {
      "speaker": "Moderator",
      "statement": "Synthesis of perspectives and final score decision..."
    }
  ],
  "holistic_score": final_holistic_score,
  "feedback": "Comprehensive, balanced feedback reflecting the consensus."
}
```

Output: Dialectical Reasoning

```
{
  "final_assessment": {
    "dialectical_reasoning": [
      {
        "speaker": "Ms. Laura Bennett",
        "statement": "From the perspective of essay structure and language mechanics, I recommend a holistic score of 3. The essay shows an attempt to follow the letter format, but the introduction is unclear and the conclusion incomplete. Paragraphs lack clear focus and transitions are minimal, making the essay choppy. Moreover, pervasive spelling, grammar, and punctuation errors significantly impair readability. These mechanics issues reduce the overall effectiveness of the essay and indicate a need for substantial revision at the grade 8 level."
      },
      {
        "speaker": "Mr. Daniel Kim",
        "statement": "I appreciate Ms. Bennett's focus on mechanics, but from a content perspective, I also lean toward a score of 3. The essay covers relevant and common concerns about computers' negative impacts on children and teens, such as physical inactivity, social isolation, and online dangers. However, the claims are sometimes vague or unsupported, and the essay ignores the positive aspects mentioned in the prompt. The examples, while present, are general and underdeveloped. So, while the essay shows basic understanding, it lacks depth and precision."
      },
      {
        "speaker": "Mrs. Angela Martinez",
        "statement": "Considering prompt fulfillment and task completion, I would give a slightly higher score of 5. The student writes a letter to a local newspaper, clearly states an opinion opposing computers' effects, and attempts persuasion with relevant points. Although the letter is incomplete and hampered by placeholders and errors, the main task is substantially complete. The essay stays focused on the topic and uses a letter format consistently, which fulfills the prompt's requirements better than the mechanics or content scores suggest."
      },
      ... (Additional discussion omitted)
      {
        "speaker": "Moderator",
        "statement": "Synthesizing the perspectives, the essay clearly fulfills the prompt by writing a letter with a stated opinion and relevant persuasive points, supporting a score around 4 to 5 for task completion and focus. However, the content is basic and sometimes vague, lacking nuance or strong evidence, which supports a moderate score of 3 for content accuracy and insight. The persuasive techniques are appropriate but uneven, suggesting a score of 4. Critically, the essay suffers from numerous spelling, grammar, punctuation, and sentence structure errors that impair readability and coherence, which justifies a lower score of 2 to 3 in mechanics and organization. Balancing these factors, the holistic score should reflect the essay's substantial task completion and clear opinion but also its frequent mechanical errors and limited development. Therefore, I determine a final holistic score of 3."
      }
    ],
    "holistic_score": 3
  }
}
```