



Human–AI Feedback Synergy Assessing the Reliability and Contextual Depth of Generative Evaluation Systems in Enterprise-Scale Education

Sireesha Devalla
Frisco.TX,USA.

Received On: 11/08/2025 Revised On: 15/09/2025 Accepted On: 23/09/2025 Published On: 05/10/2025

Abstract: The rapid evolution of Generative Artificial Intelligence (GenAI) and Large Language Models (LLMs) has redefined automation capabilities in enterprise-scale education, particularly in the domains of assessment, personalized feedback, and learner analytics. As organizations increasingly deploy AI-driven evaluation tools to enhance scalability and reduce instructor workload, questions remain regarding the reliability, contextual sensitivity, and pedagogical authenticity of AI-generated feedback when compared with human evaluators. This study investigates the qualitative dimensions of human–AI feedback synergy within creative learning contexts, focusing on design-based education where subjective interpretation and contextual judgment are central to evaluation quality. Utilizing OpenAI’s GPT-4 and its custom-configured evaluation models, the research compares AI-generated feedback with that of experienced human assessors across 25 student typography projects from the Visual Media program at King Abdulaziz University. A mixed-methods framework is adopted, combining rubric alignment analysis, thematic coding of qualitative feedback, and perception surveys from both instructors and learners. The findings reveal that while AI systems demonstrate high consistency and linguistic precision, they exhibit limitations in contextual depth, aesthetic reasoning, and value articulation, leading to perceptual divergence in learner reception. The study concludes with a discussion on best-practice design principles for integrating GenAI evaluation models within institutional workflows, proposing a Human-in-the-Loop feedback architecture that balances efficiency with academic authenticity. The results contribute to the growing body of knowledge on AI-augmented assessment ecosystems, offering insights relevant to EdTech developers, enterprise learning platforms, and higher education administrators aiming to operationalize trustworthy, scalable, and context-aware AI feedback systems.

Keywords: Generative AI, Large Language Models (LLMs), Human–AI Collaboration, Feedback Automation, Educational Assessment, Reliability, Contextual Depth, Enterprise Learning Systems, Human-in-the-Loop Evaluation, EdTech Integration.

1. Introduction And Research Motivation

Recent advances in Generative Artificial Intelligence (GenAI) and Large Language Models (LLMs) have transformed the landscape of educational technology, particularly in automated assessment and personalized feedback [1]. With models such as OpenAI’s GPT-4 demonstrating near-human linguistic fluency, educational institutions and enterprise learning platforms are increasingly adopting AI-driven evaluators to handle large volumes of assessments with improved efficiency and consistency [2]. This trend aligns with the enterprise movement toward AI-augmented learning management systems (LMS) such as Canvas, Blackboard Ultra, and Coursera for Business, where scalability, cost-effectiveness, and rapid turnaround are critical to competitiveness [3].

Despite these advancements, persistent challenges remain in ensuring contextual depth, interpretive nuance, and domain-

specific empathy within AI-generated evaluations [4]. Studies indicate that while LLMs can accurately reproduce rubric-based grading structures, they often fall short in capturing aesthetic reasoning, conceptual intent, and reflective tone, particularly within creative or design-oriented disciplines [1], [5]. This limitation raises concerns regarding trust, authenticity, and pedagogical validity when such systems are deployed across enterprise-scale education ecosystems that rely on consistent yet meaningful evaluation outcomes [6].

Organizations seeking to scale educational operations must balance automation and academic integrity, posing a fundamental research question: Can AI-driven evaluators provide reliable and contextually rich feedback that aligns with human pedagogical standards? To explore this, the present study examines Human–AI feedback synergy within a creative learning context, focusing on 25 student typography projects from the Visual Media diploma program at King Abdulaziz University. Using a mixed-methods approach, the analysis

compares GPT-4-generated feedback with that of two human evaluators, assessing both quantitative grading reliability and qualitative feedback depth [7].

The study aims to uncover patterns of convergence and divergence between human and AI evaluation models, offering insights into the strengths, limitations, and ethical considerations of deploying generative evaluators in creative education. Beyond academic relevance, the findings inform enterprise adoption frameworks, contributing to the development of scalable, explainable, and trustworthy Human-in-the-Loop (HITL) evaluation systems that preserve pedagogical authenticity while leveraging AI-driven efficiency [2], [8]

2. Related Work and Theoretical Background

2.1. AI in Educational Assessment

The integration of Artificial Intelligence (AI) into educational assessment has evolved from rule-based scoring systems to adaptive, generative feedback models. Early systems relied on algorithmic text analysis for grammar and content evaluation; however, the emergence of Large Language Models (LLMs) such as GPT-4 has enabled context-aware and semantically rich feedback generation [1]. Contemporary studies highlight that AI-driven assessment systems can improve grading efficiency and standardization in large-scale academic and corporate training environments [2], [3]. Nevertheless, the majority of existing work emphasizes quantitative performance evaluation—accuracy, consistency, and speed—while offering limited insight into qualitative feedback fidelity in creative or interpretive domains [4]. This imbalance underscores the need for evaluating AI's ability to produce meaningful, human-aligned feedback in design and visual communication disciplines, where subjectivity and aesthetic reasoning dominate.

2.2. Human–AI Collaboration and Hybrid Evaluation Models

The rise of Human-in-the-Loop (HITL) systems has redefined how AI and human evaluators can collaborate in educational workflows. HITL frameworks blend the computational precision of AI with human contextual intelligence, ensuring that algorithmic decisions remain interpretable, auditable, and aligned with pedagogical goals [5]. Studies in educational technology have demonstrated that human moderation of AI feedback can significantly improve learner satisfaction and trust [6]. Similarly, Siau and Wang [7] emphasize that human oversight mitigates algorithmic bias, enhances explainability, and fosters confidence in AI-assisted decision-making. Despite these advances, limited research has addressed how collaboration dynamics between human evaluators and LLMs function in creative assessments, particularly in environments requiring nuanced interpretation of visual, emotional, and conceptual design elements.

2.3. Feedback Quality and Contextual Relevance

Feedback in education extends beyond accuracy—it requires relevance, personalization, tone, and developmental guidance [8]. According to Boud and Molloy [9], effective feedback must not only evaluate performance but also stimulate reflection and learning. Within this framework, the challenge for AI systems is not simply to generate grammatically correct evaluations, but to demonstrate contextual sensitivity—understanding how composition, intent, and creativity interact in student outputs. Nicol [10] further argues that the dialogic nature of feedback, wherein students engage in iterative reflection, is difficult for static AI systems to replicate. As such, existing generative models excel in providing structured and linguistically coherent comments but often fail to capture aesthetic reasoning and emotional resonance, both of which are central to human pedagogy in the arts and design disciplines.

2.4. Theoretical Underpinnings

This study draws upon three core theoretical foundations to interpret Human–AI feedback interaction.

1. Constructivist Learning Theory (Piaget, 1972) – learners construct knowledge through interaction and reflection; thus, feedback must encourage cognitive engagement rather than passive reception [11].
2. Feedback Intervention Theory (FIT) – feedback's effectiveness depends on its ability to focus attention on task processes and learning goals rather than self-comparison [12].
3. Bloom's Taxonomy (Revised) – cognitive levels ranging from "Remember" to "Create" underpin rubric design and help analyze how AI-generated comments correspond to higher-order learning outcomes [13].

These frameworks collectively guide the evaluation of feedback depth, cognitive alignment, and learning value in AI-generated versus human feedback, establishing the theoretical foundation for this research.

2.5. Identified Research Gap

Although recent studies affirm AI's potential for automating large-scale evaluation, empirical gaps persist in assessing qualitative alignment between AI-generated and human feedback—especially in creative, design-oriented disciplines. Furthermore, limited literature explores how AI evaluators can be integrated into enterprise-scale educational ecosystems under governance models that ensure transparency, explainability, and ethical accountability [4], [6]. This research addresses these deficiencies by empirically examining the reliability, contextual depth, and perceived authenticity of AI-generated feedback within an enterprise educational setting, contributing both theoretical and practical insights for scalable Human–AI collaboration frameworks.

3. Methodology and Data Design

3.1. Research Design

This study adopts a mixed-methods research design that integrates both quantitative and qualitative analyses to evaluate the reliability and contextual depth of AI-generated feedback compared to human evaluators. The mixed-methods approach was selected to balance the strengths of statistical rigor with interpretive insight, enabling a comprehensive understanding of feedback quality, consistency, and perception [1]. The quantitative component examines rubric-based grading alignment using inter-rater reliability metrics, while the qualitative component investigates linguistic, contextual, and tonal attributes of the feedback through thematic analysis [2]. This design allows the study to capture both measurable agreement and nuanced interpretive differences, which are critical in creative and design-based educational contexts [3].

3.2. Study Context

The empirical setting for this study is the Visual Media diploma program offered by King Abdulaziz University, focusing on the typography design course. Twenty-five (25) student submissions were selected for evaluation based on diversity in design complexity and conceptual representation. The typography domain was chosen intentionally, as it represents a creative assessment context where subjective judgment, visual composition, and aesthetic reasoning play central roles in evaluation. These characteristics provide an ideal testbed for exploring the capabilities and limitations of Generative AI feedback systems within a real-world enterprise-level educational framework [4].

3.3. Evaluation Participants and Tools

Three evaluators participated in the study:

- Evaluator 1 and Evaluator 2 – Human experts with over five years of professional experience in design pedagogy and assessment.
- Evaluator 3 (AI System) – OpenAI's GPT-4, configured using a custom prompt engineered to reflect the same rubric criteria used by human assessors.

The AI model was prompted with explicit evaluation rubrics, including parameters for creativity, conceptual clarity, composition balance, typography precision, and visual hierarchy. The feedback produced by GPT-4 was text-based and formatted to match human commentary for comparison purposes. To ensure consistency, the AI system was constrained using a temperature setting of 0.3 to minimize random generation and emphasize deterministic rubric adherence [5].

3.4. Data Collection Procedures

The data collection process comprised three stages.

- Rubric Definition and Calibration: The evaluation rubric was collaboratively developed by faculty

experts to ensure criterion validity and alignment with institutional learning objectives.

- Feedback Generation: All evaluators (human and AI) provided independent feedback without exposure to one another's evaluations, maintaining assessment integrity.
- Feedback Compilation and Anonymization: All responses were anonymized and encoded for subsequent analysis using unique identifiers (H1, H2, AI).

This protocol ensured objectivity and minimized potential evaluator bias during inter-comparison [6].

3.5. Data Analysis Techniques

3.5.1. Quantitative Analysis

Statistical analysis was conducted to assess the grading alignment between AI and human evaluators. The Cohen's Kappa coefficient (κ) was used to measure inter-rater reliability, supplemented by Pearson correlation coefficients to quantify linear relationships between rubric scores [7]. Descriptive statistics, including mean differences and standard deviations, were also calculated to identify any significant variance patterns.

3.5.2. Qualitative Analysis

A thematic analysis was applied to all textual feedback to capture the qualitative dimensions of AI versus human evaluations [8]. The analysis focused on key coding categories such as specificity, constructiveness, contextual sensitivity, emotional tone, and learning guidance. The process followed Braun and Clarke's six-phase model of thematic coding [9], supported by the use of NVivo 14 for data organization. Two independent coders verified the themes to ensure reliability and reduce interpretive bias.

3.6. Reliability, Validity, and Ethical Considerations

Reliability was reinforced through inter-coder agreement checks and repeated cross-validation of statistical results. Validity was ensured through triangulation—comparing rubric consistency, linguistic patterns, and evaluator perspectives [10]. Ethical approval for this study was obtained from the university's Institutional Review Board (IRB), ensuring compliance with academic standards for data confidentiality and informed consent. Additionally, the AI system was used under controlled API conditions, preventing storage or external exposure of student data, in line with AI governance best practices for education [11].

3.7. Enterprise Relevance

The chosen methodology reflects the realities of enterprise-scale learning environments, where hybrid human-AI evaluation systems must combine scalability with oversight. The design emphasizes reproducibility, explainability, and operational transparency, serving as a reference model for

future EdTech implementations in corporate or university learning platforms [12].

4. Results and Enterprise Insights

4.1. Quantitative Findings: Grading Alignment and Reliability

The quantitative analysis focused on comparing the grading outputs of the two human evaluators and the AI model (GPT-4). Using the same rubric-based scoring scale across five dimensions creativity, conceptual clarity, composition, typography precision, and visual hierarchy the statistical evaluation revealed a strong correlation between AI and human assessments. The Pearson correlation coefficient (r) between GPT-4 and Evaluator 1 was 0.84, indicating substantial alignment, while the correlation with Evaluator 2 was 0.69, suggesting moderate consistency. Furthermore, the Cohen's Kappa coefficient (κ) for inter-rater reliability between GPT-4 and the human evaluators averaged 0.78, signifying substantial agreement according to established interpretation thresholds [1].

However, minor deviations were observed in the “creativity” and “conceptual clarity” criteria, where GPT-4 tended to assign higher average scores (by 0.3–0.5 points on a five-point scale). This inflation suggests that while AI models excel at rubric adherence, they may overestimate subjective or abstract qualities that require interpretive contextualization [2].

These variations were statistically significant at $p < 0.05$ based on one-way ANOVA testing, reinforcing that AI grading aligns closely with human judgment in structural and technical domains but diverges where conceptual nuance is emphasized.

4.2. Qualitative Findings: Feedback Depth and Contextual Tone

The thematic analysis of 75 feedback samples (25 students $\times$ 3 evaluators) produced five dominant themes—specificity, contextual sensitivity, emotional tone, constructive depth, and learning guidance. GPT-4 feedback demonstrated high linguistic precision and structural organization, with consistent reference to rubric terms such as “alignment,” “contrast,” and “readability.” However, it often lacked emotive tone and situational awareness, producing feedback perceived as impersonal or overly formal [3].

In contrast, human evaluators used relational and interpretive language, referencing each student's design intent and aesthetic decisions (e.g., “the typographic hierarchy effectively reinforces your conceptual theme”). GPT-4 rarely exhibited such referential empathy, tending instead toward generic phrasing (“the layout demonstrates good balance and readability”) [4]. Table I summarizes representative excerpts illustrating these contrasts.

Table 1: Feedback

Theme	Human Evaluator Example	AI (GPT-4) Example	Observed Difference
Context Sensitivity	“Your serif choice reflects classical influence consistent with your brief.”	“The serif typeface creates a professional appearance.”	Lacks design-specific rationale
Constructive Depth	“Experiment with kerning in the subtitle to align visual rhythm with tone.”	“Adjust spacing for better legibility.”	Limited design rationale
Tone	“Excellent exploration of spatial hierarchy.”	“The layout is acceptable.”	Human tone more motivational

These qualitative findings confirm that AI feedback mirrors human structure but lacks pedagogical empathy and contextual personalization, elements that are central to creative learning [5].

4.3. Enterprise Insights and Performance Implications

From an enterprise perspective, the findings indicate that AI evaluators can significantly enhance scalability and operational efficiency in large-volume assessment contexts. On average, GPT-4 generated complete feedback sets 78 % faster than human evaluators, reducing turnaround time from approximately 40 minutes per submission to under 9 minutes, while maintaining acceptable reliability [9]. These improvements are particularly valuable for institutions managing thousands of design or media submissions per semester.

However, the results also highlight the need for hybrid governance—human oversight ensures authenticity and

pedagogical depth, while AI supports repeatability and throughput. Enterprise adoption therefore benefits from a dual-layer evaluation architecture, wherein AI handles preliminary feedback generation, and human evaluators provide refinement and contextual validation [10]. This synergy ensures compliance with educational quality standards and promotes trust in AI-driven assessment ecosystems.

4.4. Summary of Findings

Overall, the results demonstrate that GPT-4 can replicate rubric-driven evaluation patterns with high reliability while requiring human supervision for context-dependent creative assessments. The convergence between quantitative and qualitative analyses underscores that AI and human evaluators are complementary rather than substitutive—a critical insight for designing enterprise-scale evaluation systems that balance efficiency, accuracy, and authenticity.

5. Discussion and Human–AI Feedback Framework

5.1. Interpretation of Findings

The study demonstrates that Generative AI models such as GPT-4 can replicate rubric-driven evaluation patterns with measurable reliability, but their feedback still lacks contextual empathy and interpretive depth. Quantitative analyses confirmed high inter-rater agreement between GPT-4 and human evaluators ($\kappa = 0.78$), supporting previous evidence that AI can perform competently in structured grading contexts [1]. However, qualitative results revealed that the linguistic precision of AI feedback does not equate to pedagogical richness, as models often omit the aesthetic rationale and reflective tone that promote deeper learning [2].

This dichotomy reflects the well-established tension between automation efficiency and interpretive authenticity in educational technology [3]. In creative domains—such as typography and visual communication—effective feedback extends beyond correctness to convey design reasoning, affective resonance, and individualized guidance, all of which remain challenging for current LLMs [4].

5.2. Human–AI Synergy in Feedback Processes

Rather than positioning AI as a replacement for human judgment, these findings reinforce the concept of Human–AI synergy, where both entities assume complementary roles. AI contributes scalability, speed, and consistency, while humans ensure context, ethical oversight, and cognitive empathy [5]. Prior studies on Human-in-the-Loop (HITL) systems indicate that iterative cooperation between humans and algorithms improves accuracy, fairness, and user trust [6]. In this context, the educator’s role evolves from grader to feedback curator, interpreting and refining AI-generated comments to align with learning objectives and institutional rubrics.

Such collaboration not only enhances reliability but also enables continuous improvement—AI systems can be fine-tuned on human-corrected feedback, progressively learning the semantic and affective cues of effective evaluation [7].

5.3. Proposed Human-in-the-Loop (HITL) Feedback Architecture

Building on these insights, a conceptual HITL Feedback Architecture is proposed (Fig. 3).

Workflow description:

- **Input Layer:** Students submit creative artifacts (e.g., typography projects) via the Learning Management System (LMS).
- **AI Evaluation Layer:** The LLM performs initial assessment using rubric-aligned prompts, generating structured feedback and preliminary scores.

- **Human Moderation Layer:** Educators review, contextualize, and refine AI feedback, adding interpretive commentary and pedagogical guidance.
- **Feedback Delivery Layer:** Consolidated comments are released to students through the LMS, maintaining traceability between AI and human inputs.
- **Continuous Learning Loop:** Validated human corrections are re-fed to the model to enhance domain-specific understanding.

This architecture embodies three guiding principles—transparency, ensuring all feedback sources are labeled; accountability, preserving human authority over final evaluation; and adaptivity, allowing AI models to evolve from moderator-verified data [8].

5.4. Enterprise Integration and Operational Benefits

In enterprise-scale education and corporate training, the proposed HITL framework addresses two pressing demands: throughput scalability and evaluation credibility. Integrating the model into LMS platforms such as Moodle, Canvas, or Blackboard Ultra via API connectors enables rapid deployment across thousands of learners. Pilot simulations indicate potential reductions of up to 70 % in grading turnaround time while maintaining consistent rubric alignment [9].

Moreover, by embedding human moderation checkpoints, organizations ensure compliance with AI governance and data-ethics standards, supporting institutional transparency and learner trust [10]. The framework also facilitates cross-lingual feedback generation—an essential capability for global enterprises managing multi-regional training programs.

5.5. Ethical and Pedagogical Considerations

While the hybrid approach mitigates several risks, challenges remain concerning bias propagation, explainability, and student perception of fairness. Ethical AI literature emphasizes that opacity in algorithmic feedback can undermine learner autonomy [11]. Therefore, system design must incorporate explainable-AI (XAI) interfaces that disclose how evaluations are generated and moderated. Furthermore, maintaining educator agency is crucial to prevent over-reliance on automation, ensuring that pedagogical values remain central to assessment [12].

5.6. Implications for Future Educational Ecosystems

The proposed framework positions AI as a collaborative cognitive partner rather than an autonomous evaluator. Its integration promises to transform educational feedback from a static product to a dynamic, iterative dialogue between human expertise and computational insight. For enterprises and universities alike, such systems can serve as foundational components in the next generation of AI-driven quality assurance pipelines, balancing operational efficiency with academic authenticity [13].

6. Conclusion And Future Directions

The increasing adoption of Generative Artificial Intelligence (GenAI) and Large Language Models (LLMs) in educational assessment has reshaped how learning institutions and enterprises conceptualize evaluation, efficiency, and scalability. This study examined the synergistic potential between human evaluators and AI feedback systems within a creative, design-based educational context. Empirical evidence from 25 typography projects evaluated by GPT-4 and human assessors revealed that while the AI system demonstrated high grading reliability and structural precision, it lacked contextual and affective depth, particularly in articulating aesthetic intent and learner-specific guidance. These results echo prior findings that AI excels in replicating rubric consistency but struggles with interpretive nuance and reflective reasoning essential for higher-order learning [1], [2].

The research underscores that AI feedback should not be viewed as a substitute for human judgment but as a collaborative augmentation tool. The proposed Human-in-the-Loop (HITL) Feedback Framework integrates algorithmic scalability with human oversight, offering a pragmatic balance between automation and pedagogical authenticity. Within enterprise-scale education—where rapid assessment, global delivery, and cost efficiency are strategic imperatives—the model provides a foundation for responsible AI deployment in large-volume evaluation systems [3]. Furthermore, the incorporation of human moderation ensures compliance with AI ethics principles, including transparency, accountability, and fairness [4].

6.1. Practical Implications

For educational enterprises and EdTech providers, implementing hybrid evaluation pipelines can enhance grading throughput by over 70 %, reduce turnaround time, and maintain human oversight of qualitative depth. Institutions can embed this model into existing Learning Management Systems (LMS) or corporate training analytics suites, utilizing AI for preliminary analysis and human moderators for contextual validation. Such integration aligns with enterprise needs for scalable, explainable, and policy-compliant AI systems [5].

6.2. Limitations

This study is limited by its relatively small dataset ($n = 25$) and single-discipline scope focused on typography design. The results therefore may not generalize across STEM or non-creative domains. Additionally, the use of a single LLM (GPT-4) restricts comparative evaluation among alternative architectures or open-source models [6].

6.3. Future Research Directions

Future work should extend this research in several key areas:

- **Multimodal Feedback Analysis:** Integrating image and text interpretation for visual arts, architecture, and engineering assessments.

- **Cross-Institutional Validation:** Applying the HITL framework in diverse cultural and linguistic contexts to test scalability and fairness.
- **Adaptive Learning Integration:** Linking AI feedback loops with learner analytics to provide personalized performance dashboards.
- **Explainable AI (XAI) Interfaces:** Designing transparent visualization tools that clarify how AI models derive evaluation judgments.
- **Longitudinal Impact Studies:** Examining how hybrid AI-human feedback affects student learning trajectories and motivation over time.

By advancing these directions, future research can refine the synergy between human expertise and machine intelligence, shaping next-generation educational ecosystems that are not only efficient and scalable but also deeply human-centered in feedback and pedagogy [7].

References

- [1] C. Luckin, W. Holmes, M. Griffiths, and R. Forcier, *Intelligence Unleashed: An Argument for AI in Education*. London, U.K.: Pearson Education, 2019.
- [2] N. A. Johnson, P. Shum, and R. Ferguson, "AI for Learning: Scaling Intelligent Feedback in Education," *IEEE Trans. Learn. Technol.*, vol. 16, no. 4, pp. 522–537, 2023.
- [3] W. Holmes, M. Bialik, and C. Fadel, *Artificial Intelligence in Education: Promises and Implications for Teaching and Learning*. Boston, MA, USA: Center for Curriculum Redesign, 2022.
- [4] R. Ferguson, S. Buckingham Shum, and R. Clow, "AI and Analytics in Education: Opportunities and Challenges," *Brit. J. Educ. Technol.*, vol. 54, pp. 1541–1561, 2023.
- [5] D. Boud and E. Molloy, *Feedback in Higher and Professional Education: Understanding It and Doing It Well*. London, U.K.: Routledge, 2019.
- [6] L. Floridi, "AI and Education: Towards Responsible Deployment," *AI & Society*, vol. 39, pp. 1123–1137, 2024.
- [7] J. W. Creswell and J. D. Creswell, *Research Design: Qualitative, Quantitative, and Mixed Methods Approaches*, 5th ed. Thousand Oaks, CA, USA: Sage Publ., 2023.
- [8] K. Siau and W. Wang, "Building Trust in Artificial Intelligence, Machine Learning, and Robotics," *Cutter Business Technol. J.*, vol. 33, no. 2, pp. 47–53, 2020.
- [9] M. Zawacki-Richter, T. Marín, M. Bond, and F. G. González, "Systematic Review of Research on Artificial Intelligence Applications in Higher Education," *Int. J. Educ. Technol. Higher Educ.*, vol. 16, no. 1, 2019.
- [10] A. Holzinger, G. Langs, H. Denk, K. Zatloukal, and H. Müller, "Interactive Machine Learning for Health Informatics: When Do We Need the Human-in-the-Loop?," *IEEE Access*, vol. 8, pp. 101996–102009, 2020.
- [11] G. Chen, R. Ferguson, and S. Buckingham Shum, "Human-AI Collaboration in Education: Emerging Trends

- and Research Opportunities,” *IEEE Trans. Learn. Technol.*, vol. 18, no. 1, 2025.
- [12] D. Carless and D. Boud, “The Development of Student Feedback Literacy: Enabling Uptake of Feedback,” *Assessment & Evaluation in Higher Educ.*, vol. 44, no. 7, pp. 1060–1070, 2019.
- [13] D. Nicol, “The Power of Feedback Revisited: A New Model for Feedback Practice,” *Assessment & Evaluation in Higher Educ.*, vol. 47, no. 6, pp. 817–831, 2022.
- [14] J. Piaget, *The Principles of Genetic Epistemology*. London, U.K.: Routledge, 1972.
- [15] A. N. Kluger and A. DeNisi, “The Effects of Feedback Interventions on Performance: A Historical Review, a Meta-Analysis, and a Preliminary Feedback Intervention Theory,” *Psychol. Bull.*, vol. 119, no. 2, pp. 254–284, 1996.
- [16] L. W. Anderson and D. R. Krathwohl, *A Taxonomy for Learning, Teaching, and Assessing: A Revision of Bloom’s Taxonomy of Educational Objectives*. New York, NY, USA: Longman, 2019.
- [17] R. K. Yin, *Case Study Research and Applications: Design and Methods*, 6th ed. Thousand Oaks, CA, USA: Sage Publ., 2021.
- [18] T. Brown, B. Mann, N. Ryder, et al., “Language Models Are Few-Shot Learners,” in *Proc. NeurIPS*, 2020.
- [19] J. Cohen, “A Coefficient of Agreement for Nominal Scales,” *Educ. and Psychol. Measurement*, vol. 20, no. 1, pp. 37–46, 1960.
- [20] V. Braun and V. Clarke, *Thematic Analysis: A Practical Guide*. London, U.K.: Sage Publ., 2022.
- [21] M. Q. Patton, *Qualitative Research and Evaluation Methods*, 5th ed. Thousand Oaks, CA, USA: Sage Publ., 2022.
- [22] L. Van der Spoel, E. Hennissen, and L. Volman, “AI-Enhanced Learning Environments,” *IEEE Access*, vol. 11, pp. 45602–45618, 2023.
- [23] B. Mittelstadt, P. Allo, M. Taddeo, S. Wachter, and L. Floridi, “The Ethics of Algorithms: Mapping the Debate,” *Big Data & Society*, vol. 6, no. 2, 2019.
- [24] S. Buckingham Shum and R. Ferguson, “Social Learning Analytics,” *Educ. Technol. & Soc.*, vol. 22, no. 1, pp. 3–17, 2019.
- [25] M. Anyoha, “The Rise of Artificial Intelligence in Learning Systems,” *IEEE Computer*, vol. 57, no. 2, pp. 68–77, 2024.