



Human-AI Collaboration in Skill-Based Classroom Feedback

Jia Hui Lim¹, Wei Ming Tan^{1*}, Nur Aisyah Rahman¹

¹Department of Information Systems and Analytics, School of Computing, National University of Singapore, COM1, 13 Computing Drive, Singapore 117417, Singapore.

Corresponding Author*: Wei Ming Tan E-mail: Ming_Tan@comp.nus.edu.sg

ARTICLE INFO

Keywords

human-AI collaboration
skill-based learning
classroom feedback
educational AI
feedback analytics
learner autonomy
multilevel regression
learning efficiency

Published:

01 September 2026

ABSTRACT

This study investigates how human-AI collaboration can improve feedback quality and learning efficiency in skill-based classroom instruction. The study is designed to involve approximately 520 students and 36 instructors from practice-oriented courses over a 12-week teaching cycle, generating around 7,800 student practice submissions, 5,200 AI-generated feedback reports, 3,600 teacher feedback records, 2,400 learner self-correction logs, and 1,500 peer-feedback entries. The research measures feedback timeliness, feedback specificity, correction adoption rate, practice improvement score, learner autonomy, teacher workload, and student satisfaction. AI feedback is generated through multimodal analysis of text, image, audio, or task-performance data, while teacher feedback is used as the professional reference for evaluation. Quantitative methods include paired-sample t-tests, ANCOVA, multilevel regression, mediation analysis, inter-rater reliability testing, and learning gain analysis. The study compares three feedback conditions: teacher-only feedback, AI-only feedback, and human-AI collaborative feedback. Evaluation indicators include feedback response time, task revision frequency, performance score improvement, learner self-regulation score, teacher workload reduction, and agreement between AI and teacher judgments. The innovation of this study lies in treating AI not as a replacement for teachers but as a feedback amplifier that improves timeliness and diagnostic coverage while preserving human judgment for expressive, contextual, and high-level learning evaluation.

1. Introduction

Feedback is a central component of skill-based learning because improvement depends on repeated practice, timely correction, reflection, and adjustment. In practice-oriented courses, students often produce multiple forms of learning evidence, including written responses, images, audio recordings, videos, demonstrations, and task-performance records. Teachers must review these materials, identify errors, explain possible improvements, and determine whether students have correctly applied previous feedback. As class sizes and the volume of learning data increase, maintaining timely and individualized feedback becomes increasingly difficult. Artificial intelligence (AI) offers a possible way to support this process by accelerating initial review, identifying recurring errors, and generating rapid formative comments. A large-scale review of 129 studies reported that AI-supported feedback was frequently associated with improved learner action and performance, although many studies provided limited information about how AI-generated

Citation: Lim, J. H., Tan, W. M., & Rahman, N. A. (2026). Human-AI collaboration in skill-based classroom feedback. *Journal of Adaptive Learning and Psychological Growth*, 1(2), 33-43.

3154-7915/© The Authors. Published by J&L Academic Group PLT. This is an open access article under the CC BY 4.0 license.
<https://doi.org/10.64744/jalpg.2026.280>.

comments were produced, validated, or supervised (Zhang et al., 2025). Research on multimodal learning analytics has similarly expanded, particularly in the analysis of complex learner data, but stronger connections are still needed among data analysis, instructional decisions, and measurable learning outcomes (Yürüm, 2025). These developments are also relevant to public art and practice-oriented education, where teaching often depends on collaborative interaction and the efficient allocation of instructional resources. Recent research has shown that collaborative teaching structures and the optimization of teaching-resource efficiency can play an important role in improving the organization and delivery of public art education (Lin et al., 2026). Within this context, AI-supported feedback may provide an additional mechanism for improving instructional efficiency, especially when large amounts of student practice data must be reviewed within limited teaching time. However, faster feedback does not necessarily mean better feedback. Effective feedback must remain aligned with the learning objective, the learner's current level, the characteristics of the task, and the instructional context.

Most empirical evidence on AI-generated feedback currently comes from writing education. Studies in secondary and higher education have linked AI-assisted feedback with improvements in revision behavior, learning motivation, writing quality, and assessment performance (Huang, 2026; Urzúa et al., 2025). These findings suggest that AI can be particularly useful when learning tasks involve identifiable patterns of error and repeated cycles of revision. Automated systems can rapidly detect grammar problems, structural weaknesses, incomplete responses, or repeated mistakes and can provide learners with immediate opportunities to revise their work. Such responsiveness may reduce the delay between performance and correction, which is especially important when students complete several practice tasks within a short period. At the same time, the educational value of feedback depends on whether students understand, accept, and act on the comments they receive. An automatically generated suggestion that is technically accurate but poorly matched to the learner's needs may have limited instructional value. AI-generated feedback must therefore be evaluated not only by its speed or linguistic completeness, but also by its relevance, accuracy, actionability, and contribution to subsequent learning. Comparative studies of AI and teacher feedback have produced mixed results. Some research has found similar short-term improvements under AI-generated and teacher-generated feedback conditions (Gui et al., 2025; Wu et al., 2026), whereas other studies suggest that teacher feedback may lead to more stable or sustained learning over time (Sánchez-García & Reyes-de-Cózar, 2025). These differences are understandable because AI and teachers do not perform identical instructional functions. AI systems are well suited to rapid screening, repeated error detection, standardized analysis, and the processing of large amounts of learner data. Teachers, in contrast, are generally better positioned to interpret learner intention, task meaning, emotional expression, prior performance, classroom interaction, and the broader context in which a learning problem occurs. In an evaluation of feedback quality, human-generated feedback performed better than ChatGPT-generated feedback across most assessed dimensions, although the overall difference was relatively small (Liang et al., 2025; Zhang et al., 2026). Other work has shown that AI comments may partly correspond with teacher judgments while still failing to capture important differences in student background, instructional needs, or task context (Kapxhiu, 2025; Qi & Qiao, 2026). Research comparing AI and human feedback has also found that both can provide useful evaluative information but may offer insufficient support for deeper reflection and self-monitoring (Su & Wu, 2026). These results suggest that feedback quality should not be reduced to the question of whether AI or teachers produce higher scores. The more important issue is which aspects of feedback each source performs effectively and how these strengths can be combined within an instructional system. The source of

feedback may also influence how learners interpret and use it. Students can respond differently to identical or similar comments depending on whether they believe the feedback was produced by an AI system or by a teacher. Perceived authority, trust, personalization, fairness, and emotional support can affect whether learners accept a correction and whether they are willing to revise their work. This is particularly important in skill-based learning, where students often need repeated feedback over an extended period rather than a single evaluative comment. Excessive dependence on AI may also create unintended learning consequences. Research has shown that unrestricted AI assistance can improve performance during supported practice while reducing later independent performance when the tool is removed (Amanlou et al., 2026; Ogbonnaya et al., 2025). This finding raises an important distinction between immediate task completion and durable learning. If AI gives students answers or corrections too quickly, learners may improve the submitted product without fully understanding why the original performance was weak. Effective AI-supported feedback should therefore encourage active correction, explanation, reflection, and self-monitoring rather than simply replace the learner's own judgment.

The limitations of text-only feedback become even more apparent in practice-oriented courses. Many skill-based tasks cannot be evaluated adequately through written responses alone. Performance may depend on movement, posture, timing, sound, visual composition, procedural accuracy, or the coordination of several observable behaviors. Emerging research has begun to investigate AI-supported feedback in practical and performance-based learning environments. Studies have reported the use of AI for movement analysis, task control, performance correction, and other forms of practice feedback (Gao et al., 2026). Multimodal systems can also analyze non-verbal information from classroom video, including posture and observable behavioral signals (Whitehead et al., 2025; Xiong et al., 2026). These developments create the possibility of feedback systems that integrate text, image, audio, video, and task-performance data rather than relying on a single source of evidence. Such systems may be especially useful in courses where teachers must repeatedly review similar practice tasks from large numbers of students. Nevertheless, multimodal analysis introduces additional challenges. Different data types represent different aspects of learner performance, and combining them does not automatically produce a valid educational judgment. A video-based system may accurately detect a movement pattern but still misunderstand the instructional purpose of that movement. An audio model may identify timing or pronunciation differences without knowing whether those differences are pedagogically important in a particular task. Image analysis may detect visible features but fail to interpret artistic intention, style, or contextual meaning. The reliability of AI-supported feedback therefore depends not only on model accuracy but also on the relationship between detected features and the teacher's professional criteria (Hong et al., 2026; Li et al., 2026). For this reason, teacher ratings remain important as a reference for evaluating whether AI judgments correspond with expert expectations. Agreement between AI and teachers can provide useful evidence of technical and instructional validity, but disagreement is equally informative because it can reveal the types of tasks or criteria for which human judgment remains necessary (Üzümlü et al., 2025; Wu & Chen, 2025). Human–AI collaboration may therefore provide a more appropriate framework than complete automation. Under a collaborative model, AI can perform rapid initial analysis, identify common problems, summarize repeated patterns, and generate preliminary comments, while teachers verify accuracy and focus their attention on complex, context-dependent, expressive, or higher-level judgments. This division of work may improve feedback coverage without removing professional responsibility from the teacher. Early research suggests that AI-supported systems can help teachers improve feedback practices and support classroom interaction. However, the available evidence

remains limited in duration, sample size, task diversity, and instructional setting. Many studies examine short interventions, one course, or one type of learner output. As a result, it remains unclear whether human–AI feedback can simultaneously improve learning outcomes, reduce response time, maintain feedback quality, and lower teacher workload when implemented across a larger practice-oriented teaching environment. Several methodological limitations also remain in the existing literature. Research has concentrated heavily on writing and online learning, while practical and skill-based classrooms have received substantially less attention. Many studies rely on one institution, one assignment type, or a relatively short observation period, which limits the extent to which findings can be generalized to sustained classroom practice. Direct comparisons among teacher-only, AI-only, and human–AI collaborative feedback conditions are still relatively uncommon. Studies also tend to emphasize final achievement scores while giving less attention to the process through which feedback produces learning. Response time, comment detail, correction adoption, revision frequency, self-correction behavior, learner autonomy, satisfaction, and teacher workload are all important indicators of feedback effectiveness, yet they are not consistently examined together. Similarly, agreement between AI and teacher judgments is often calculated from a limited number of samples without considering differences among courses, instructors, or classes. These limitations make it difficult to determine whether AI-supported feedback improves only efficiency, only immediate performance, or the broader learning process.

The present study addresses these gaps by examining the effectiveness of teacher-only, AI-only, and human–AI collaborative feedback in practice-oriented courses over a 12-week period. The study includes approximately 520 students and 36 instructors and analyzes a large body of authentic learning evidence, including around 7,800 practice submissions, 5,200 AI feedback reports, 3,600 teacher feedback records, 2,400 self-correction logs, and 1,500 peer-feedback entries. AI-supported feedback draws on text, image, audio, and task-performance data, while teacher evaluations are used as the professional reference for assessing the quality and consistency of automated judgments. The study examines not only changes in learning performance but also feedback speed, comment detail, correction adoption, practice improvement, learner autonomy, teacher workload, student satisfaction, and agreement between AI and teacher ratings. Paired-sample tests and ANCOVA are used to examine changes in learner outcomes, while multilevel regression accounts for variation across students, instructors, and instructional settings. Mediation analysis is used to investigate whether factors such as correction adoption and learner autonomy help explain the relationship between feedback mode and learning improvement, and inter-rater analyses are used to evaluate the consistency of AI and teacher judgments. By comparing three feedback conditions within the same research framework, the study aims to clarify the respective strengths and limitations of automated and human feedback and to determine whether their combination can improve both instructional efficiency and learning quality. The study does not conceptualize AI as a substitute for professional teaching judgment. Instead, it investigates how AI can extend the speed, scale, and coverage of formative feedback while teachers retain responsibility for contextual interpretation, expressive quality, complex performance evaluation, and higher-level instructional decisions. The findings are expected to contribute theoretically by extending AI-feedback research from single-mode and outcome-centered evaluation toward a multidimensional model that connects feedback source, learner response, and learning development. They also provide methodological evidence from multimodal data and class-level analysis and offer practical guidance for designing sustainable human–AI feedback systems in large practice-oriented courses. In this way, the study seeks to identify a feedback model that improves instructional efficiency without weakening teacher judgment, learner independence, or the quality of skill development.

2. Materials and Methods

2.1 Participants and Study Setting

The study included 520 students and 36 instructors from skill-based courses at six public education institutions. The courses covered visual design, oral communication, digital media, laboratory practice, and technical training. Courses were included when they used repeated practice, clear scoring rubrics, and at least four revision tasks during the 12-week study period. Students were required to attend at least 80% of the lessons and complete a baseline task before group assignment. Courses with fewer than 15 students or incomplete digital submission records were excluded. The students had different levels of prior experience, while the instructors had 3–18 years of teaching experience. The final dataset contained about 7,800 practice submissions, 5,200 system-generated feedback reports, 3,600 teacher feedback records, 2,400 self-correction logs, and 1,500 peer-feedback entries.

2.2 Experimental Design and Feedback Conditions

A controlled three-group design was used. Courses were assigned to the teacher-only, system-only, or teacher–system feedback group. Group assignment was balanced by course type, class size, baseline score, and instructor experience. In the teacher-only group, instructors reviewed each task and gave comments based on the course rubric. In the system-only group, feedback was produced from text, images, audio, or task-performance records and was released without teacher review. In the teacher–system group, the system produced the first round of comments. Instructors then checked the comments, removed incorrect information, and added advice on context, expression, and advanced skill use. All groups followed the same course content, task schedule, scoring rubric, and submission deadlines. This design allowed direct comparison of feedback speed, learning gains, correction use, and teacher workload.

2.3 Measures and Quality Control

Feedback quality was measured by response time, comment detail, error coverage, clarity of advice, and agreement with teacher ratings. Learning outcomes included task-score improvement, revision frequency, correction adoption, learner self-regulation, and satisfaction. Teacher workload was measured as the average time spent reviewing each submission. Two senior instructors independently rated 15% of the submissions using the same rubric. The intraclass correlation coefficient was used to test agreement between raters, and values above 0.75 were considered acceptable. System-generated comments were checked for factual errors, repeated content, missing support, and unclear wording. A pilot study was conducted with 42 students and four instructors before the main study. Unclear survey items and items with weak item–total correlations were revised. Cronbach’s alpha values above 0.70 were required for the learner autonomy and satisfaction scales.

2.4 Data Processing and Model Formulas

Duplicate records, incomplete submissions, and feedback records without time stamps were removed before analysis. Missing survey values below 5% were estimated by multiple imputation. Baseline and final task scores were standardized within each course to reduce differences among scoring rubrics. Learning gain for student i in course j was calculated as:

$$LG_{ij} = \text{Post}_{ij} - \text{Pre}_{ij}$$

where LG_{ij} is the learning gain, $Post_{ij}$ is the final task score, and Pre_{ij} is the baseline task score. A multilevel model was used because repeated tasks were nested within students, and students were nested within courses:

$$Y_{ijk} = b_0 + b_1(\text{Group}_j) + b_2(\text{Time}_k) + b_3(\text{Group}_j * \text{Time}_k) + u_j + v_i(j) + e_{ijk}$$

where Y_{ijk} is the outcome for student i in course j at time k . Group_j represents the feedback condition, and Time_k represents the measurement point. The terms u_j and $v_i(j)$ represent course-level and student-level random effects. The term e_{ijk} represents the residual error. Both formulas use standard letters and symbols and can be copied directly into Word without character errors.

2.5 Statistical Analysis and Ethical Procedures

Paired-sample t-tests were used to examine changes within each group. ANCOVA was used to compare final performance after controlling for baseline score, prior experience, attendance, and course type. Mediation analysis tested whether correction adoption and learner autonomy explained the link between feedback condition and learning gain. Agreement between system and teacher ratings was tested with intraclass correlation coefficients and weighted kappa values. Effect sizes were reported as Cohen's d , partial eta squared, and 95% confidence intervals. Statistical significance was set at $p < 0.05$. All records were coded before analysis, and personal names were removed. Students and instructors received written information about the study and gave informed consent. Participation was voluntary. Withdrawal did not affect course grades, teaching duties, or access to learning materials.

3. Results and Discussion

3.1 Feedback Speed and Comment Quality

After data screening, 508 students and 35 instructors were included in the final analysis. The dataset contained 7,612 practice submissions, 5,104 system-generated feedback reports, 3,482 teacher feedback records, 2,318 self-correction logs, and 1,437 peer-feedback entries. No significant differences were found among the three groups in baseline score, prior experience, attendance, or course type (all $p > 0.05$). The mean feedback time was 26.8 h in the teacher-only group, 4.6 min in the system-only group, and 2.4 h in the teacher-system group. The teacher-system group received the highest scores for comment detail (4.31 out of 5), error coverage (4.26), and clarity of advice (4.18). The system-only group identified more errors than the teacher-only group, but some comments did not fit tasks that involved personal expression or classroom context. Error coverage reached 91.6% in the teacher-system group, compared with 86.9% in the system-only group and 79.8% in the teacher-only group. These results show that automated feedback reduced response time and increased error coverage. Teacher review improved the accuracy and relevance of the comments (Cisneros-González et al., 2025; Du, 2025). Cisneros-González et al. (2025) also found that automated feedback was valued for its speed and broad coverage, although some comments were general or lacked clear examples. Du et al. (2025) reported that automated feedback supported frequent review, but teacher input remained important when system comments were unclear or incomplete.

3.2 Learning Gains and Correction Adoption

All three groups improved during the 12-week teaching period, but the degree of improvement differed. The mean task score increased by 8.3 points in the teacher-only group, 7.2 points in the system-only group, and 12.7 points in the teacher–system group. ANCOVA showed a significant group effect after baseline score, attendance, prior experience, and course type were controlled, $F(2, 499) = 34.61$, $p < 0.001$, partial eta squared = 0.12. Pairwise tests showed that the teacher–system group performed better than both the teacher-only and system-only groups (both $p < 0.001$). The difference between the teacher-only and system-only groups was small for highly structured tasks. However, teacher-only feedback produced better results for oral expression, visual composition, and open-ended design. The correction adoption rate was 81.3% in the teacher–system group, 69.5% in the teacher-only group, and 63.8% in the system-only group. Students in the teacher–system group also completed more revisions, with a mean of 4.4 revisions per task. The corresponding values were 3.5 in the teacher-only group and 3.1 in the system-only group (Wang et al., 2026). Changes in learning performance are shown in Fig. 1. These results differ partly from Wang et al. (2026), who found that teacher feedback produced stronger gains than language-model feedback in scientific writing. The difference may result from the use of teacher review after the first automated response in the present study (Pereira & Mello, 2025). The findings are also consistent with Pereira et al. (2025), who reported that immediate feedback improved skill performance and encouraged repeated practice.

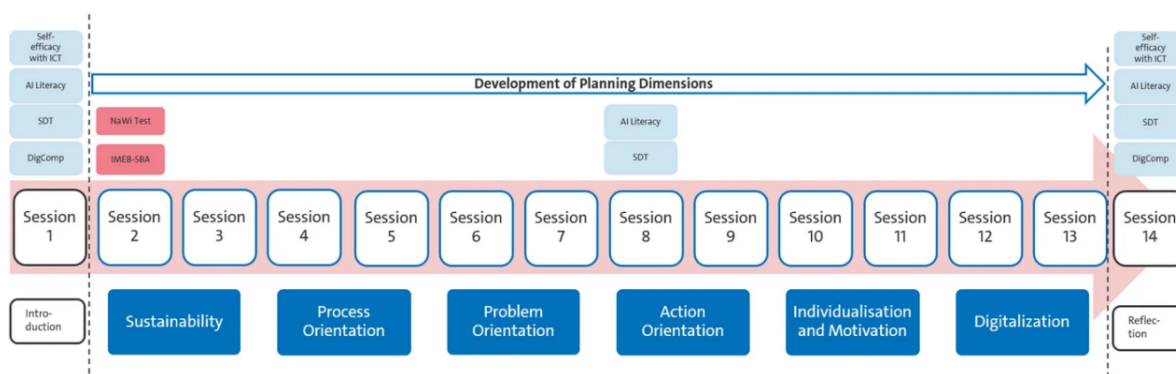


Fig. 1. Changes in task scores and correction use under three feedback conditions.

3.3 Learner Autonomy and Teacher Workload

Learner autonomy increased in all three groups, with the largest increase in the teacher–system group. The mean self-regulation score increased from 3.11 to 3.86 in the teacher–system group, from 3.09 to 3.51 in the system-only group, and from 3.13 to 3.48 in the teacher-only group. Multilevel regression showed a significant interaction between feedback condition and time, $\beta = 0.31$, $p < 0.001$. Mediation analysis showed that correction adoption explained 36.4% of the effect of teacher–system feedback on learning gain. Learner autonomy explained a further 19.7%. The 95% confidence intervals for both indirect effects did not include zero. Teacher review time fell from 14.8 min per submission in the teacher-only group to 8.2 min in the teacher–system group. This represented a workload reduction of 44.6%. However, instructors still needed time to correct system errors and add comments for complex tasks (Yin et al., 2026; Zhang, 2026). Student satisfaction was highest in the teacher–system group at 4.29 out of 5, followed by the teacher-only group at 3.91 and the system-only group at 3.73. The links among feedback condition, correction adoption, learner autonomy, teacher workload, and learning gain are shown in Fig. 2. These results show that teacher–system feedback can reduce routine review time without removing teacher control. Yin et al. (2026) also found that teacher–AI cooperation was linked to stronger teaching

engagement through higher confidence in technology use. Zhang et al. (2026) reported that changes in practice behavior partly explained the effect of immediate feedback on skill learning.

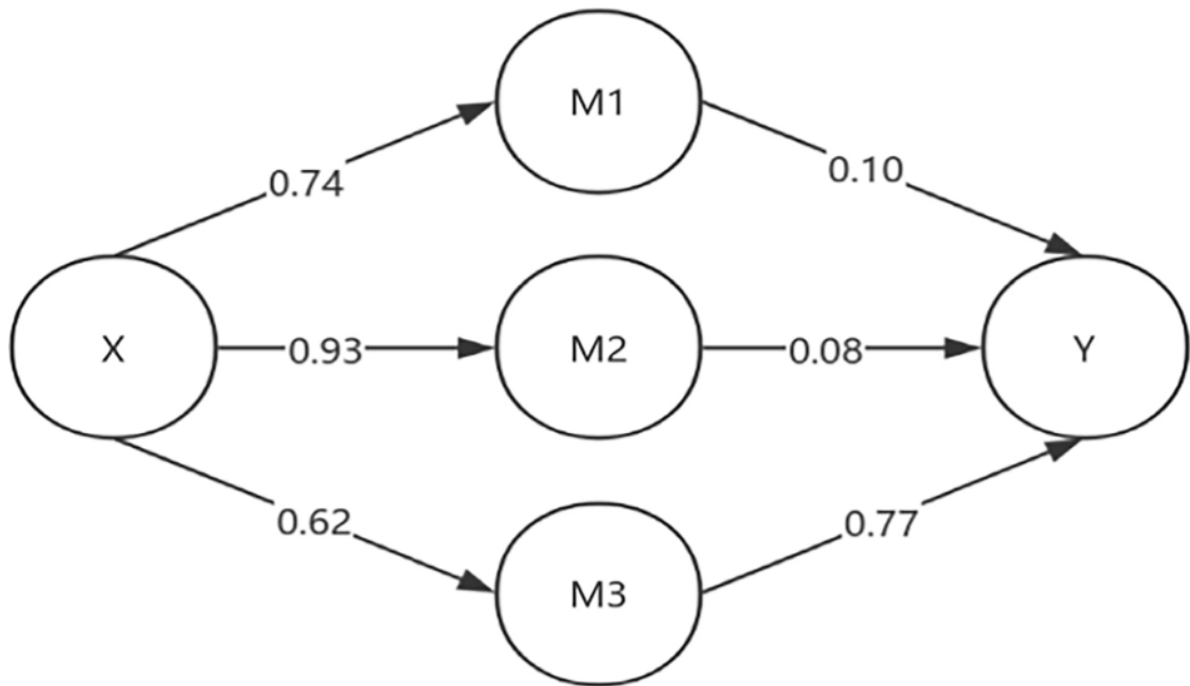


Fig. 2. Links between feedback condition, learner autonomy, teacher workload, and learning gain.

3.4 Agreement Between System and Teacher Judgments

Agreement between system and teacher ratings was acceptable for the full sample. The intraclass correlation coefficient was 0.82, and the weighted kappa value was 0.76. Agreement was highest for procedural accuracy (ICC = 0.89), task completion (ICC = 0.87), and clearly defined technical errors (ICC = 0.84). Lower agreement was found for personal expression (ICC = 0.66), use of context (ICC = 0.62), and advanced professional judgment (ICC = 0.58). Instructors revised 8.1% of the system comments because of incorrect error detection. They also added missing advice to 11.4% of the feedback reports. System accuracy was higher for text and structured task records than for audio, visual design, and open-ended performance. Students with lower baseline scores gained more from teacher–system feedback than highly experienced students. These learners needed frequent correction of clear and repeated errors. The results show that automated feedback is most suitable for first-round checking. Teachers remain necessary for complex meaning, emotional support, and task-based judgment (Abdi Tabari & Huang, 2026; Kim et al., 2026). Kim et al. (2026) found that students preferred human feedback in areas that required trust and personal connection. Abdi Tabari et al. (2026) also reported that supportive wording improved students’ views of automated feedback, although it did not directly improve revision quality. The present study was limited to six institutions and one 12-week teaching period. Instructor blinding was not possible, and system performance may differ across models and subject areas. Future studies should examine long-term learning, transfer to new tasks, and different levels of teacher review.

4. Conclusion

This study compared teacher-only, system-only, and teacher–system feedback in skill-based courses over 12 weeks. The teacher–system group showed the greatest improvement in task scores, correction use, learner autonomy, and satisfaction. It also received feedback faster and required less

teacher review time. System feedback was useful for finding repeated and technical errors, while teachers were more reliable in judging context, expression, and complex skill performance. The study adds evidence by examining learning results, feedback quality, correction behavior, rating agreement, and teacher workload within the same design. The findings suggest that system feedback is most effective as a first review, followed by teacher checking and added guidance for complex tasks. This approach can improve feedback speed and coverage while keeping teachers responsible for professional judgment. The study was limited to six institutions and one teaching cycle, and the results may differ across subjects, systems, and learner groups. Further research should include more course types, longer follow-up periods, and new tasks completed without system support.

Acknowledgments

We are grateful to all respondents who participated in this study.

References

1. Abdi Tabari, M., & Huang, Z. (2026). Exploring the Role of AI in Shaping L2 Teachers' Affective Dispositions and Task-Based Language Teaching Implementation: A Scale-Based Study. *European Journal of Education*, 61(2), e70566.
2. Amanlou, M., Amou-Jafari, Y., Livian, M., Boloukazari, F., Bagheri, F., & Bahrak, B. (2026). Beyond Access: Guided LLM Scaffolding for Independent Learning in Undergraduate Statistics. *arXiv preprint arXiv:2606.01375*.
3. Cisneros-González, J., Gordo-Herrera, N., Barcia-Santos, I., & Sánchez-Soriano, J. (2025). JorGPT: Instructor-aided grading of programming assignments with large language models (LLMs). *Future Internet*, 17(6), 265.
4. Du, Y. (2025). Research on Digital Quality Traceability System for Temperature-Controlled Supply Chain of Foreign Trade Wine Driven by Blockchain and IoT. *Business and Social Sciences Proceedings*, 4, 57-65.
5. Gao, G., Gao, R., Lu, C., Gao, R., & Kuang, Y. (2026, March). Security Governance Methods and Quantitative Evaluation for Enterprise SMS and Verification Code Systems. In *2026 International Conference on Generative Artificial Intelligence and Information Security (GAIIS)* (pp. 455-458). IEEE.
6. Gui, H., Fu, Y., Wang, Z., & Zong, W. (2025). Research on Dynamic Balance Control of Ct Gantry Based on Multi-Body Dynamics Algorithm.
7. Hong, Z., Xu, T., & Chen, H. (2026). A Study on Test Coverage Mapping and Release Risk Prediction in Continuous Delivery of Cloud Services.
8. Huang, R. (2026). Edge-Centric Generative AI: A Survey on Efficient Inference for Large Language Models in Resource-Constrained Environments. *Journal of Computer and Communications*, 14(4), 238-253.
9. Kapxhiu, L. (2025). AI in language learning assessment: Examining the limitations of online testing and the reliability of teacher judgement. *Academic Journal of Business, Administration, Law and Social Sciences*, 11(3), 59-72.

10. Kim, Y., Kang, S., D'Arienzo, M., & Taguchi, N. (2026). Comparing traditional and task-based approaches to teaching pragmatics: Task design processes and learning outcomes. *Language Teaching Research*, 30(5), 2631-2654.
11. Li, J., Ma, R., & Gu, Y. (2026). From Audit Requirements to Computable Software Architecture: A Rule-Engine and Evidence-Indexing Method for Trusted Digital Infrastructure.
12. Liang, R., Feifan, F. N. U., Liang, Y., & Ye, Z. (2025). Emotion-Aware Interface Adaptation in Mobile Applications Based on Color Psychology and Multimodal User State Recognition. *Frontiers in Artificial Intelligence Research*, 2(1), 51-57.
13. Lin, T., Chen, A., & Spivack, Z. (2026). Research on Collaborative Piano Teaching Models and Teaching Resource Efficiency Optimization in Public Art Education Systems. Available at SSRN 6284818.
14. Ogbonnaya, C. N., Li, S., Tang, C., Zhang, B., Sullivan, P., Erden, M. S., & Tang, B. (2025, March). Exploring the role of artificial intelligence (AI)-driven training in laparoscopic suturing: a systematic review of skills mastery, retention, and clinical performance in surgical education. In *Healthcare* (Vol. 13, No. 5, p. 571). MDPI.
15. Pereira, A. F., & Mello, R. F. (2025). A systematic literature review on large language models applications in computer programming teaching evaluation process. *IEEE Access*, 13, 113449-113460.
16. Qi, C., & Qiao, X. (2026, May). Design Patterns for Multi-Agent Systems in Production. In *2026 International Conference on Intelligent Systems, Automation and Control (ISAC)* (pp. 198-202). IEEE.
17. Sánchez-García, R., & Reyes-de-Cózar, S. (2025). Enhancing project-based learning: A framework for optimizing structural design and implementation—A systematic review with a sustainable focus. *Sustainability*, 17(11), 4978.
18. Su, D., & Wu, Y. (2026). Multilevel Growth and Social Network Modeling for Functional Communication Enhancement in Students with Intellectual Disabilities.
19. Urzúa, C. A. C., Ranjan, R., Saavedra, E. E. M., Badilla-Quintana, M. G., Lepe-Martínez, N., & Philominraj, A. (2025). Effects of AI-assisted feedback via generative chat on academic writing in higher education students: A systematic review of the literature. *Education Sciences*, 15(10), 1396.
20. Üzümlü, B., Elçiçek, M., & Pesen, A. (2025). Development of teachers' perception scale regarding artificial intelligence use in education: Validity and reliability study. *International Journal of Human-Computer Interaction*, 41(5), 2776-2787.
21. Wang, Y., Chen, J., Wang, Y., & Yin, X. (2026). Application of Obtainable Biological Agent Characteristics in Efficacy Stratification of Oral Anti-Obesity Drugs.
22. Whitehead, R., Nguyen, A., & Järvelä, S. (2025). Utilizing multimodal large language models for video analysis of posture in studying collaborative learning: A case study. *Journal of Learning Analytics*, 12(1), 186-200.
23. Wu, C., & Chen, H. (2025). Research on system service convergence architecture for AR/VR system.

24. Wu, J., Zhang, J., Wu, D., & Peng, Y. (2026). Visual Transformation Mechanisms for Integrating Digital Cultural Heritage Resources into Junior High School Art Curricula.
25. Xiong, W., Zeng, Y., & Guo, Y. (2026). A Study on the Impact of MEP Space Organization in Large-Scale Commercial Complexes on Investment Efficiency and Its Mechanisms.
26. Yin, J., Rao, H., & Huang, Y. (2026, March). Dynamic Modeling and Heterogeneity Analysis of Platform User Behavior Time Series. In 2026 International Conference on AI in Education Technology and Applications (AIETA) (pp. 30-33). IEEE.
27. Yürüm, O. R. (2025). Technology-enhanced multimodal learning analytics in higher education: a systematic literature review. *IEEE Access*, 13, 92057-92073.
28. Zhang, Y., Gu, W., & Wang, J. (2025, December). Research on First Article Inspection (FAI)-Driven Quality Assurance Methods for Wind Turbine Installation and Operation & Maintenance and Their Effect on Reliability Improvement. In *Proceedings of the 2025 6th International Conference on Computer Science and Management Technology* (pp. 1700-1705).
29. Zhang, Z. (2026). A Study on Return Optimization in E-commerce for Complex Consumer Goods Driven by Installation Information Quality. Available at SSRN 6734720.
30. Zhang, Z., Tong, Y., & Gao, Y. (2026). Causal Representation Learning for Explainable Early Warning in High-Stakes Decision Systems.