

Vietnamese EFL Teachers' Perceptions of ChatGPT-Generated Feedback: A Cross-Sectional Survey Study

Nguyen Huu Hoang
(hoangnguyen@hanu.edu.vn)
Hanoi University, Vietnam

Abstract

This cross-sectional survey examined Vietnamese EFL teachers' perceptions of ChatGPT-generated feedback quality and integration willingness factors. A total of 188 university teachers from six Vietnamese cities completed a questionnaire addressing feedback evaluations and Technology Acceptance Model constructs. Results revealed hierarchical perception patterns where teachers rated content feedback most favourably ($M = 3.67$), followed by organization feedback ($M = 3.45$), with language feedback receiving lowest evaluations ($M = 3.21$). This pattern contradicts traditional automated writing evaluation capabilities, suggesting that large language models reverse conventional performance hierarchies by excelling in higher-order rhetorical tasks while generating scepticism regarding surface-level corrections. Multiple regression analysis demonstrated that perceived usefulness ($\beta = .47$) and institutional support ($\beta = .34$) emerged as primary predictors of integration willingness, explaining 58.3% of variance. The prominence of institutional support challenges Technology Acceptance Model universality, revealing how hierarchical decision-making structures in Vietnamese contexts modify traditional acceptance pathways. Teachers' sophisticated discrimination in evaluating AI feedback indicates successful implementation requires acknowledgment of pedagogical complexities beyond technological capabilities. These findings establish foundational knowledge for evidence-based AI integration strategies while contributing theoretical insights regarding cultural adaptation of technology acceptance frameworks in Southeast Asian educational environments.

Share and cite

Nguyen, H. (2026). Vietnamese EFL Teachers' Perceptions of ChatGPT-Generated Feedback: A Cross-Sectional Survey Study. *Electronic Journal of Foreign Language Teaching [e-FLT]*, 23(1), 78–93. <https://doi.org/10.56040/hhng2311>

1 Introduction

The emergence of ChatGPT in late 2022 has precipitated unprecedented discussions about artificial intelligence's transformative potential in educational contexts, particularly within language teaching and learning environments (Peters et al., 2023). This large language model, developed by OpenAI, represents a significant advancement over traditional automated writing evaluation systems, demonstrating sophisticated capabilities in understanding, analysing, and responding to student writing in ways that approximate human-like feedback provision (Bucol & Sangkawong, 2024; Imran & Almusharraf, 2023). Unlike conventional automated essay evaluation tools that primarily focus on surface-level linguistic features, ChatGPT exhibits enhanced capacity for addressing higher-order

writing concerns including content development, argumentation quality, and organizational coherence, thereby challenging established assumptions about the limitations of machine-generated feedback (Rudolph et al., 2023).

Within English as a Foreign Language (EFL) contexts, the provision of effective writing feedback has long constituted one of the most labour-intensive and pedagogically complex aspects of instruction (Li & Zhang, 2021). EFL teachers, particularly those working in large-enrolment university settings, face substantial challenges in delivering timely, detailed, and constructive feedback that addresses both linguistic accuracy and rhetorical effectiveness (Wu & Luo, 2022). Quality feedback significantly influences writing development, yet time constraints often result in delayed or superficial responses that fail to promote meaningful revision (Wu & Schunn, 2020; Zhai & Ma, 2021). These challenges become particularly acute in contexts where teachers manage heavy course loads and large class sizes, conditions that characterize much of higher education in developing nations.

Vietnamese higher education exemplifies these challenges while presenting distinctive cultural and institutional factors that influence technology adoption patterns. Rapid expansion of English education in Vietnam has created substantial demand while institutional resources lag, with teacher-student ratios inhibiting individualized feedback (Ngo & Tran, 2023; Nguyen, 2024). Traditional Vietnamese educational approaches emphasize teacher authority and direct instruction, creating potential tensions with innovative pedagogical technologies that redistribute instructional responsibilities (Earl, 2024). Additionally, hierarchical Vietnamese institutions mean that technology adoption depends on administrative support and alignment with established values rather than individual preferences (Vo & Mu, 2020).

Despite the potential benefits of AI-enhanced feedback systems, current research on teacher perceptions of ChatGPT in educational contexts remains limited in two specific respects. First, existing studies are predominantly conducted in Western educational settings where cultural, institutional, and technological contexts differ markedly from those found in Southeast Asian universities (García-Peñalvo, 2023; Zhai, 2022), so that findings on teacher acceptance from individualist, low power-distance settings cannot be assumed to hold in hierarchical Asian institutions where decisions about technology adoption are filtered through administrative approval rather than individual choice (Alvi & Khechine, 2025; Hofstede, 2001). Second, within the limited research conducted in Asian and Southeast Asian contexts, attention has centred on student perspectives or on the technical accuracy of AI output (Soori et al., 2025; Xu et al., 2025), leaving the teacher as evaluator and gatekeeper of feedback quality comparatively unexamined. This combination creates a clear opening for an inquiry that asks how teachers themselves, rather than students or algorithms, judge AI-generated feedback within a hierarchical, non-Western institutional setting, and that, in doing so, tests whether a model built primarily on Western, individualist assumptions about user acceptance extends to a collectivist, hierarchical setting where adoption decisions are rarely made by individuals alone.

This study addresses this combination of gaps by investigating Vietnamese EFL teachers' perceptions of ChatGPT feedback quality and factors influencing their integration willingness, extending the Technology Acceptance Model to a setting where institutional hierarchy, rather than individual preference, governs adoption, and complementing prior work on AI feedback accuracy with evidence on how the teachers responsible for acting on that feedback evaluate it. Drawing upon the Technology Acceptance Model, whose core constructs of perceived usefulness and perceived ease of use were developed to explain how evaluators of a new tool form judgments that precede adoption, this research examines how perceived usefulness, ease of use, and contextual factors shape teacher attitudes toward AI-enhanced feedback. Two research questions guide this investigation:

RQ1: What are Vietnamese EFL teachers' perceptions of ChatGPT feedback quality across different writing assessment dimensions?

RQ2: What factors influence Vietnamese EFL teachers' willingness to integrate ChatGPT feedback into their writing instruction practices?

2 Literature review

2.1 Theoretical framework

Understanding teacher adoption of educational technologies requires theoretical grounding that explains the decision-making processes underlying technology integration behaviours. The Technology Acceptance Model, originally conceptualized by Davis (1989), provides the most extensively validated framework for examining individual acceptance of information technologies across diverse contexts. This model posits that two primary beliefs—perceived usefulness and perceived ease of use—serve as fundamental determinants of behavioural intention to use technology, which subsequently influences actual usage behaviour. Perceived usefulness refers to the degree to which individuals believe that using a particular technology will enhance their job performance, while perceived ease of use encompasses the extent to which individuals expect the technology to be free from effort (Venkatesh & Davis, 2000).

Subsequent developments have expanded TAM to accommodate educational environments, recognizing that teacher adoption involves additional considerations. UTAUT, developed by Venkatesh et al. (2003), incorporates facilitating conditions, social influence, and moderating factors to provide more explanatory power. Educational technology research demonstrates that institutional support and peer acceptance significantly influence adoption, particularly in hierarchical systems where organizational factors constrain individual autonomy (Teo, 2011). Cross-cultural studies reveal that cultural dimensions like power distance moderate relationships between acceptance beliefs and behavioural intentions, requiring contextualization when applying technology adoption models in non-Western educational settings (Srite & Karahanna, 2006). TAM, rather than UTAUT, is retained as the organizing framework for the present investigation because its two core constructs, perceived usefulness and perceived ease of use, capture the evaluative judgment that precedes any adoption decision, whereas UTAUT's additional constructs describe the conditions surrounding that decision. UTAUT's facilitating conditions and social influence are accordingly treated as relevant moderators within this TAM-centred framework rather than adopted as a parallel model. Applying TAM here also extends its reach beyond the individualist, low power-distance settings in which the model was originally validated, testing whether perceived usefulness and perceived ease of use retain their explanatory power among teachers whose acceptance beliefs are shaped by institutional and collective expectations, a condition the original formulation does not specify.

2.2 Automated writing evaluation and ChatGPT

Automated writing evaluation has progressed from simple rule-based systems to sophisticated AI applications capable of approximating human-like text analysis (Singh & Namin, 2025). Early systems like Project Essay Grade and E-rater relied on surface-level linguistic features to generate scores correlating with human raters (Shermis & Burstein, 2013). While demonstrating acceptable reliability for large-scale assessment, their feedback remained limited to basic error identification, failing to address higher-order concerns like argumentation quality and rhetorical effectiveness (Stevenson & Phakiti, 2014).

The emergence of large language models represents a paradigmatic shift in automated writing evaluation capabilities, with ChatGPT demonstrating unprecedented proficiency in understanding textual meaning, generating contextually appropriate responses, and providing detailed explanatory feedback across multiple writing dimensions (Faiz et al., 2023). Unlike traditional tools operating through predetermined algorithms and scoring rubrics, ChatGPT employs transformer-based neural

networks trained on vast textual corpora to generate responses exhibiting apparent understanding of rhetorical conventions, argumentative structures, and disciplinary expectations (Haber et al., 2025).

Recent empirical investigations reveal both promising capabilities and significant limitations. Alnemrat et al. (2025) found no statistically significant difference between AI-generated and teacher-generated feedback in enhancing EFL argumentative writing performance overall, yet dimensional analysis revealed divergent effectiveness patterns. Soori et al. (2025) demonstrated that AI feedback excelled in improving lexical resources, while hybrid approaches combining AI and teacher feedback yielded superior outcomes for task achievement, coherence, and grammatical accuracy. These findings challenge assumptions about uniform AI effectiveness across writing components and highlight the importance of matching feedback modality to specific pedagogical objectives. This technological advancement introduces complexities regarding feedback accuracy, pedagogical appropriateness, and the potential for generating plausible but incorrect suggestions that could misdirect student revision efforts (Sallam, 2023).

2.3 Teacher perceptions of AI-generated feedback

Educator evaluations of AI-generated feedback quality across writing components directly influence adoption patterns and implementation outcomes. Research on automated writing evaluation reveals significant variability in teacher acceptance, driven by perceived threats to professional autonomy, pedagogical authenticity concerns, and uncertainty about learning outcomes (Zawacki-Richter et al., 2019). Teachers consistently express tension between efficiency gains and pedagogical quality, with many questioning whether AI systems deliver meaningful feedback supporting student learning or merely identify surface-level errors (Xu et al., 2025).

Quality evaluations vary substantially depending on which writing component AI feedback targets. Soori et al. (2025) found educators rated AI feedback highest for lexical resources, where students achieved superior vocabulary performance compared to teacher-only feedback groups. This positive assessment reflects AI's corpus-driven lexical analysis capabilities, providing extensive language database access for precise vocabulary recommendations. However, teachers evaluated AI feedback more sceptically for higher-order concerns—particularly task achievement and coherence—where contextual understanding becomes essential. Hybrid approaches combining AI and teacher feedback received highest quality ratings for these complex components, indicating educators recognize AI's limitations in delivering feedback requiring interpretive sophistication.

ChatGPT's emergence has intensified these quality debates. Alm and Ohashi (2024) documented that language educators worldwide, including those in Asian contexts, demonstrated moderate ChatGPT awareness but minimal educational application following its release. Only 16% employed ChatGPT for teaching, primarily for resource creation rather than feedback provision, reflecting quality concerns about assessment capabilities. Teachers with direct experience held more favourable quality perceptions, suggesting practical engagement reduces scepticism. Alnemrat et al. (2025) found that while AI and teacher feedback produced equivalent overall writing improvements, teachers' willingness to employ ChatGPT varied considerably by task: high for creating reading materials, low for providing feedback or assessment. This selective pattern indicates educators differentiate between AI applications augmenting their work versus those potentially supplanting pedagogical judgment, highlighting the need for implementation strategies aligned with pedagogical values and cultural expectations.

2.4 Teacher perceptions and AI integration in educational practice

Teacher attitudes toward AI integration reflect interactions between technological capabilities, pedagogical beliefs, and institutional constraints, with autonomy threats and authenticity concerns

influencing adoption (Zawacki-Richter et al., 2019). Studies investigating automated writing evaluation systems have consistently identified tensions between efficiency gains and pedagogical quality, with many educators expressing scepticism about AI systems' ability to provide meaningful feedback supporting student learning rather than merely identifying surface-level errors (Xu et al., 2025).

ChatGPT has intensified these debates, with educators expressing enthusiasm for workload reduction alongside apprehension about quality variations and diminished pedagogical relationships (Rudolph et al., 2023; Zeevy-Solovey, 2024). Cultural factors appear to moderate these perception patterns significantly, with research indicating that teachers in collectivist educational cultures prioritize institutional support and peer acceptance more heavily than counterparts in individualist contexts (Johnson et al., 2023). Additionally, disciplinary differences influence acceptance patterns, with language teachers demonstrating particular sensitivity to feedback quality given the central importance of written communication skills in their pedagogical objectives (Al-Khresheh, 2024). These findings suggest that AI integration success depends not merely on technological capabilities but on careful consideration of educator concerns and development of implementation strategies aligning with established pedagogical values and cultural expectations.

2.5 EFL writing feedback practices in Vietnamese higher education contexts

AI integration research has emerged predominantly from Western settings, where infrastructure, pedagogical autonomy, and policies differ markedly from Asian systems (Mustroph & Steinbock, 2024; Shin et al., 2021). This geographical concentration limits generalizability to contexts where hierarchical decision-making and centralized structures regulate pedagogical innovation differently. Asian educational systems exhibit higher power distance in teacher-student relationships, collectivist values shaping technology adoption, and administrative structures mediating individual teacher autonomy (Alvi & Khechine, 2025; Hofstede, 2001). These distinctive characteristics necessitate context-specific investigations to understand how AI-generated feedback tools function within educational environments structured around fundamentally different cultural assumptions.

Vietnamese higher education exemplifies these broader Asian patterns while presenting additional contextual factors. Rapid English program expansion has created enrolment pressures, with teacher-student ratios often exceeding 1:50 in writing courses, constraining individualized feedback (Le & Barnard, 2009). Traditional educational approaches emphasize teacher authority and student deference, creating classroom dynamics where students rarely challenge feedback or engage in collaborative revision (Nguyen, 2024). Institutional hierarchies mean technology adoption depends substantially on administrative support and alignment with established values rather than individual teacher preferences (Vo & Mu, 2020).

Current feedback practices reveal heavy reliance on error correction and surface-level linguistic commentary, with limited attention to higher-order rhetorical concerns (Nguyen, 2022). This pattern reflects time constraints and teacher preparation limitations, as many instructors lack extensive process-oriented writing pedagogy training (Pham & Truong, 2021). Technology integration remains uneven, with urban universities demonstrating greater digital capacity while rural institutions face connectivity barriers (Nguyen et al., 2025). Generational differences affect adoption, with younger teachers showing greater acceptance of digital innovations compared to experienced educators maintaining traditional approaches (Vo & Mu, 2020). Institutional policies requiring administrative approval create implementation gaps between technological availability and classroom utilization (Nguyen, 2024). These conditions, heavy feedback loads, uneven technological capacity, and approval-dependent adoption, define the institutional setting into which ChatGPT-generated feedback would be introduced, yet none of the teacher perception or AI integration research reviewed above was conducted within this setting.

2.6 Research gaps

The convergence of artificial intelligence capabilities, teacher acceptance research, and contextual educational factors reveals two gaps that the studies reviewed above leave open. The first is geographical. The teacher acceptance evidence cited in this review, including Johnson et al. (2023) on power distance and Mustroph and Steinbock (2024) and Shin et al. (2021) on institutional autonomy, was generated in contexts where individual teachers, not administrative hierarchies, drive adoption decisions. Whether the same acceptance constructs hold once perceived institutional and collective expectations, rather than individual preference alone, shape the underlying beliefs, as established for Vietnamese higher education above (Vo & Mu, 2020), remains untested.

The second gap is perspectival. ChatGPT's rapid emergence has outpaced empirical research on the teacher specifically. The studies reviewed above that examine ChatGPT feedback directly evaluate quality by writing component (Soori et al., 2025) or measure student outcomes and task-level willingness to use ChatGPT (Xu et al., 2025; Alm & Ohashi, 2024), but none asks teachers to weigh that component-level quality against their decision to integrate ChatGPT feedback into instruction. That evaluative step, the teacher's quality judgment that precedes any adoption decision, remains the missing link in the literature on ChatGPT-generated feedback. The absence of validated instruments compounds these gaps, particularly in non-Western contexts where cultural factors may moderate acceptance relationships differently than established models predict (Hofstede, 2001).

These limitations necessitate research examining teacher perceptions of AI-generated feedback within specific cultural and institutional contexts while developing measurement approaches appropriate for emerging technologies. The present investigation addresses the geographical gap through RQ1, which asks how Vietnamese EFL teachers perceive ChatGPT feedback quality across writing dimensions rather than assuming Western dimensional patterns transfer directly, and addresses the perspectival gap through RQ2, which asks what factors shape these teachers' willingness to integrate that feedback into instruction. This focus, building on the TAM rationale outlined above, offers theoretical insight into cultural moderation of acceptance constructs alongside practical knowledge for implementation in resource-constrained, hierarchically governed educational environments.

3 Methodology

3.1 Research design

This investigation employed a cross-sectional survey design to examine Vietnamese EFL teachers' perceptions of ChatGPT-generated feedback quality and integration determinants. The design selection was theoretically grounded in the Technology Acceptance Model's emphasis on measuring attitudinal precursors to behavioural intention (Davis, 1989). Given ChatGPT's nascent status in Vietnamese educational contexts, capturing baseline perceptions provides essential data for understanding potential adoption trajectories. Cross-sectional methodology proves particularly effective for technology acceptance research when investigating initial attitudes toward emerging educational technologies, as demonstrated in studies examining teacher adoption of digital tools (Scherer et al., 2019; Teo, 2011). The quantitative approach enables standardized measurement across participants while facilitating statistical analysis of variable relationships, which proves essential for empirically testing technology acceptance frameworks (Venkatesh & Davis, 2000). This design acknowledges the exploratory nature of ChatGPT research while providing methodologically sound foundations for establishing baseline knowledge in Vietnamese higher education contexts.

3.2 Participants

The study recruited 188 Vietnamese EFL teachers from public and private universities across six major cities representing northern, central, and southern regions. Participants were selected through

stratified purposive sampling to ensure balanced representation across key demographic variables including institutional type, teaching experience, and educational background. While this urban sampling ensures technological infrastructure necessary for the study, it limits generalizability to rural Vietnamese educational contexts where infrastructure and technology adoption patterns may differ significantly. This geographic distribution captured institutional diversity characteristic of Vietnamese urban higher education.

The sample comprised 112 female and 76 male teachers with teaching experience ranging from 2 to 28 years ($M = 11.7$, $SD = 6.3$). Educational qualifications included 34% holding master's degrees and 66% possessing doctoral qualifications, with institutional affiliation of 58% from public universities and 42% from private institutions. Age distribution spanned 26 to 54 years ($M = 38.2$, $SD = 8.1$), capturing perspectives across different career stages. Regarding prior ChatGPT exposure, 72% of participants reported some familiarity with the technology, though only 31% had used it for educational purposes. All participants possessed minimum B2-level English proficiency according to the Common European Framework of Reference and had taught academic writing courses within the previous two years. This demographic profile ensures adequate representation of the urban Vietnamese EFL teacher population while maintaining sufficient sample size for statistical analysis according to established power analysis guidelines (Cohen, 1992).

3.3 Instruments

The research instrument comprised a 32-item questionnaire specifically designed to address the absence of validated measures for assessing teacher perceptions of AI-generated feedback in Vietnamese EFL contexts. The questionnaire was originally developed in English and subsequently translated into Vietnamese by two independent bilingual experts, followed by back-translation validation to ensure cultural and linguistic appropriateness. Translation accuracy was verified through expert review and pilot testing to confirm item clarity and cultural relevance. The questionnaire incorporated four distinct sections addressing demographic characteristics, ChatGPT feedback quality perceptions, technology acceptance factors, and behavioural intentions. The demographic component included prior ChatGPT exposure levels measured through general usage frequency and educational-specific application to establish baseline technological experience, addressing variables that influence technology adoption patterns among educators (Teo, 2011).

The primary measurement component consisted of 27 items organized into theoretically grounded constructs: content feedback quality (5 items), organization feedback quality (5 items), and language feedback quality (5 items) based on established writing evaluation taxonomies (Biber et al., 2011), and Technology Acceptance Model constructs including perceived usefulness (5 items), perceived ease of use (4 items), institutional support (4 items), facilitating conditions (4 items), and behavioural intention (4 items) adapted from Davis (1989) and Venkatesh et al. (2003). All items employed five-point Likert scales, with reliability analysis revealing acceptable internal consistency across constructs (α ranging from .85 to .92), supporting the instrument's psychometric properties for Vietnamese EFL contexts.

3.4 Data collection procedure

Data collection occurred between March and May 2025 through an online survey platform administered via Google Forms. This timeframe was strategically selected to avoid academic calendar disruptions while ensuring adequate response rates during the spring semester when university faculty maintain regular schedules. The digital approach was necessitated by geographic distribution requirements across multiple Vietnamese provinces and practical considerations related to accessing diverse institutional contexts efficiently. Participant recruitment employed a multi-stage strategy beginning with identification of potential respondents through professional networks, academic conferences, and institutional contacts established through prior research collaborations. Initial contact

messages included detailed study descriptions, participation requirements, and estimated completion times to facilitate informed decision-making. The survey link was distributed through educational forums, professional associations, and direct institutional channels to maximize reach while maintaining sampling integrity. Reminder messages were sent at two-week intervals to enhance response rates without creating participant burden. Data collection adhered to ethical guidelines established by the APA (2017), with informed consent obtained prior to survey access and voluntary participation emphasized throughout the process. Participants received explicit assurance regarding anonymity, data confidentiality, and withdrawal rights at any stage without penalty. Response monitoring indicated completion rates of approximately 73%, with the final sample representing adequate geographic and institutional diversity to support the research objectives while maintaining statistical power requirements for planned analyses.

3.5 Data analysis

Data analysis utilized SPSS version 27.0 to address both research questions through distinct analytical procedures. Preliminary analysis included descriptive statistics and reliability assessment using Cronbach's alpha coefficients to evaluate internal consistency across measurement scales. The first research question examining teacher perceptions of ChatGPT feedback quality was addressed through descriptive statistics and one-way ANOVA to identify differences in perception patterns based on demographic variables including teaching experience, institutional type, and prior ChatGPT exposure. The second research question investigating factors influencing integration willingness employed Pearson correlation analysis to establish relationships between Technology Acceptance Model constructs and behavioural intentions. Multiple regression analysis determined the relative predictive power of perceived usefulness, ease of use, institutional support, and facilitating conditions on integration willingness while controlling for demographic variables (Field, 2018). Effect sizes were calculated using Cohen's conventions to assess practical significance beyond statistical significance. Normality assumptions were verified through Shapiro-Wilk tests, and missing data patterns justified listwise deletion procedures (Tabachnick & Fidell, 2019).

4 Results

4.1 Teacher perceptions of ChatGPT feedback quality across writing assessment dimensions

Vietnamese EFL teachers in this study demonstrated hierarchical perceptions of ChatGPT feedback quality across writing assessment dimensions, revealing statistically significant differences in evaluation patterns. As presented in Table 1, participants rated content feedback quality most favourably ($M = 3.67$, $SD = 0.82$), followed by organization feedback quality ($M = 3.45$, $SD = 0.91$), with language feedback quality receiving the lowest evaluations ($M = 3.21$, $SD = 0.95$). One-way repeated measures ANOVA confirmed these differences as statistically significant, $F(2, 374) = 12.34$, $p < .001$, $\eta^2 = .042$, indicating that participants in this study perceive ChatGPT's competencies as domain-specific rather than uniformly distributed across feedback categories.

Table 1. Teacher Perceptions of ChatGPT Feedback Quality Across Writing Assessment Dimensions (N=188)

Dimension	M	SD	95% CI	Positive Rat-ings (≥4)	Strong Agree-ment (5)	Strong Disagree-ment (1)
Content Feed-back	3.67	0.82	[3.55, 3.79]	68%	23%	8%
Organization Feedback	3.45	0.91	[3.32, 3.58]	52%	18%	12%
Language Feed-back	3.21	0.95	[3.07, 3.35]	38%	15%	28%

Content feedback quality received the highest ratings, with 68% of participants expressing positive perceptions (ratings ≥4) and 23% indicating strong agreement (rating 5). Only 8% of participants expressed strong disagreement with ChatGPT’s content feedback capabilities. The moderate mean score (M = 3.67) indicates that while a majority of teachers viewed content feedback positively, 32% remained neutral or sceptical about ChatGPT’s capacity to provide appropriate content guidance.

Organization feedback perceptions revealed intermediate acceptance levels, with 52% of participants rating ChatGPT’s organizational feedback capabilities positively. Strong agreement was expressed by 18% of participants, while 12% indicated strong disagreement. The moderate standard deviation (SD = 0.91) suggests considerable variability in teacher assessments across the sample.

Language feedback evaluations showed the most polarized response pattern, with only 38% of participants expressing positive perceptions. The substantial standard deviation (SD = 0.95) reflects this polarization, with 28% strongly disagreeing with ChatGPT’s language feedback effectiveness while 15% endorsed its capabilities with strong agreement. This distribution indicates the most divided opinions among the three feedback dimensions assessed.

Demographic analysis revealed experience-dependent perception patterns across all feedback dimensions. Teachers with extensive experience (>15 years) demonstrated significantly lower acceptance across all dimensions compared to early-career colleagues (≤5 years), $F(2, 185) = 8.72$, $p < .001$, $\eta^2 = .086$. This pattern indicates differential evaluation criteria based on teaching experience levels within the sample.

4.2 Factors influencing Vietnamese EFL teachers’ ChatGPT integration willingness

Multiple regression analysis demonstrated that the combined Technology Acceptance Model framework explained 58.3% of variance in behavioural intention among participants in this study ($R^2 = .583$, $F(5, 182) = 50.73$, $p < .001$), indicating substantial predictive power for technology adoption decisions within this sample. Correlation analysis revealed statistically significant relationships between all predictor variables, as shown in Table 2, with perceived usefulness correlating most strongly with institutional support ($r = .52$, $p < .001$), indicating interdependent rather than independent operation of these factors within this sample.

Table 2. Correlation Matrix and Regression Results for ChatGPT Integration Predictors (N = 188)

Variable	1	2	3	4	β	t
1. Perceived Usefulness	---				.47***	7.23
2. Institutional Support	.52***	---			.34***	5.18
3. Perceived Ease of Use	.38***	.31***	---		.23**	3.44
4. Facilitating Conditions	.29***	.41***	.35***	---	.19*	2.67

Note. * $p < .05$, ** $p < .01$, *** $p < .001$. $R^2 = .583$.

The regression analysis identified perceived usefulness as the strongest predictor of integration willingness ($\beta = .47, p < .001$), with teachers who recognized ChatGPT's potential to enhance feedback efficiency and quality expressing significantly greater intention to incorporate the technology into their pedagogical practices. Institutional support emerged as the second strongest predictor ($\beta = .34, p < .001$), demonstrating that administrative endorsement and organizational backing significantly influenced individual teacher behaviours within this sample. The strong correlation between these primary predictors ($r = .52$) suggests that organizational support enhances teachers' recognition of technological utility.

Secondary factors contributed meaningfully but with reduced predictive power compared to utility and institutional considerations. Perceived ease of use showed moderate influence ($\beta = .23, p < .01$), indicating that participants prioritized technological benefits and organizational support over user-friendliness concerns. Facilitating conditions contributed significantly to the model ($\beta = .19, p < .05$), with teachers expressing greater integration willingness when perceiving adequate technical infrastructure and training support. Demographic analysis revealed that teaching experience demonstrated a negative relationship with integration willingness ($\beta = -.16, p < .05$), suggesting that veteran teachers harboured greater scepticism toward AI integration compared to early-career colleagues, while institutional type showed no significant predictive value.

5 Discussion

5.1 Teacher perceptions of ChatGPT feedback quality across writing assessment dimensions

Vietnamese EFL teachers demonstrated hierarchical perceptions reversing conventional automated writing evaluation patterns: content feedback ($M = 3.67$), organization ($M = 3.45$), then language ($M = 3.21$). This contradicts Western studies where language feedback receives highest ratings (Al-Khresheh, 2024; Zeevy-Solovey, 2024), suggesting that judgments of feedback quality are not stable across educational contexts in the way Western-derived automated writing evaluation research has generally assumed. Prior work has shown that AI feedback quality varies by writing component when measured against student outcomes (Soori et al., 2025); what this pattern adds is evidence that teachers themselves, not only outcome measures, rank these components differently once the setting shifts outside Western, individualist contexts, and rank them in the reverse of the order typically reported. Similar patterns emerge across Asian contexts, with Japanese teachers demonstrating parallel adoption hesitancy at 16% usage (Alm & Ohashi, 2024), raising questions about whether dimensional hierarchies reflect shared Confucian-heritage educational values or distinct national factors. The cross-sectional design cannot establish which explanation holds; it can only show that Vietnamese teachers' ordering is the reverse of the pattern reported elsewhere (Stevenson & Phakiti, 2014), a finding consistent with either sophisticated pedagogical reasoning about AI's cultural-rhetorical limitations or with theory-based evaluation grounded in limited practical experience (31% educational use), without the data distinguishing between the two.

Superior content ratings reflect recognition of ChatGPT's corpus-driven rhetorical pattern identification capabilities (Haber et al., 2025), diverging from Western scepticism about cultural appropriateness (Zawacki-Richter et al., 2019). Alnemrat et al. (2025) found AI and teacher feedback produced equivalent task achievement improvements in Saudi contexts, supporting content feedback viability. The 32% Vietnamese scepticism is plausibly related to awareness that students employ indirect Confucian-influenced argumentation patterns (Nguyen, 2024; Pham & Truong, 2021), though the survey did not ask teachers directly why they distrusted ChatGPT's content judgments, so whether this scepticism reflects genuine detection of content weaknesses or concern that ChatGPT imposes Western rhetorical norms on Vietnamese academic discourse cannot be determined from the present data and is offered here as an interpretation rather than a finding.

Moderate organizational acceptance ($M = 3.45$) reveals sophisticated cost-benefit analysis rather than simple approval. Vietnamese learners navigate conflicting rhetorical paradigms—implicit logical connections valued in Vietnamese discourse versus explicit transitions expected in English academic writing (Nguyen, 2024; Pham & Truong, 2021). Teachers' intermediate ratings suggest recognition that while ChatGPT identifies surface organizational issues, it cannot address underlying cross-cultural rhetorical transfer requiring pedagogical expertise. Soori et al. (2025) demonstrated hybrid feedback combining AI and teacher input achieved superior coherence outcomes, empirically validating teachers' intuition that AI organizational guidance requires human mediation.

Language scepticism ($M = 3.21$) contradicts Western patterns where automated correction receives highest acceptance (Xu et al., 2025; Shermis & Burstein, 2013), revealing nuanced understanding that effective error correction requires systematic knowledge of L1-to-L2 transfer patterns and developmental sequences (Li & Zhang, 2021)—precisely the contextual linguistic knowledge large language models lack. However, Soori et al. (2025) found AI significantly enhanced lexical resources, suggesting teachers may conflate distinct language feedback types. The vocabulary-grammar distinction proves consequential: lexical enhancement involves corpus-based suggestions where AI excels, while grammatical accuracy requires understanding error sources where AI provides de-contextualized corrections. Response polarization (28% strongly disagreed; 15% strongly agreed) indicates this conceptual confusion spans the sample, with some teachers recognizing AI's vocabulary capabilities while others generalize language scepticism across all linguistic features.

Experience-dependent patterns merit separate analysis given their magnitude. Teachers with extensive experience (>15 years) demonstrated significantly lower acceptance than early-career colleagues (≤ 5 years) across all feedback dimensions, $F(2, 185) = 8.72$, $p < .001$, $\eta^2 = .086$. This aligns with professional development research (Molefi et al., 2024; Scherer et al., 2019) yet exceeds Western effect sizes (Johnson et al., 2023), suggesting cultural amplification of experience effects. In hierarchical Vietnamese academic systems where pedagogical authority accumulates with seniority (Vo & Mu, 2020), veteran teachers' extensive tacit knowledge about Vietnamese learner challenges functions simultaneously as an asset—enabling sophisticated AI evaluation grounded in practical expertise—and a liability—increasing scepticism about AI replicating contextualized pedagogical understanding that required decades to develop.

5.2 Factors influencing Vietnamese EFL teachers' ChatGPT integration willingness

Perceived usefulness emerged as the strongest integration predictor ($\beta = .47$), aligning with Western TAM studies where utility perceptions consistently dominate (Scherer et al., 2019; Teo, 2011), yet transcending simple replication. Vietnamese EFL teachers face substantial workload pressures due to large classes and extensive feedback responsibilities (Nguyen, 2024), creating conditions where efficiency gains assume heightened salience. The observed effect ($\beta = .47$) exceeds typical Western sizes ($\beta = .30-.40$), indicating that resource constraints intensify utility-focused decision-making (Nguyen et al., 2025). This amplification suggests that while utility perceptions operate universally, their relative importance varies with institutional resource availability—a relationship warranting theoretical attention in frameworks applied across development contexts.

Institutional support as the second strongest predictor ($\beta = .34$) fundamentally challenges Western TAM research where organizational factors typically function as facilitating conditions rather than primary determinants (Venkatesh & Davis, 2000). This finding illuminates how hierarchical decision-making structures reshape acceptance pathways: administrative endorsement significantly influences individual behaviours where centralized processes govern pedagogical innovation (Vo & Mu, 2020). Alm and Ohashi (2024) documented parallel patterns across Asian contexts, suggesting regional rather than nation-specific phenomena. The strong correlation between perceived usefulness and institutional support ($r = .52$) reveals bidirectional influence absent in Western studies where these factors operate independently (Johnson et al., 2023)—organizational endorsement not merely enables adoption but shapes utility perceptions themselves, indicating teachers may evaluate

technological benefits partially through institutional validation rather than solely through independent professional judgment.

Perceived ease of use showed attenuated influence ($\beta = .23$), contradicting Western patterns where user-friendliness assumes greater importance (Venkatesh et al., 2003). This departure likely reflects highly educated faculty possessing sufficient technical competence to overcome learning curves when motivated by utility. The relative de-emphasis appears more pronounced than Western studies reporting stronger effects ($\beta = .35-.45$) (Al-Khresheh, 2024). Alnemrat et al. (2025) similarly found both groups demonstrated significant improvements regardless of technological complexity, indicating that educators prioritize pedagogical value over implementation ease.

Facilitating conditions contributed significantly ($\beta = .19$), underscoring how material constraints create implementation barriers operating independently of attitudinal factors (Nguyen, 2024). This aligns with research documenting uneven integration, where urban universities demonstrate greater capacity while rural institutions face connectivity barriers (Nguyen et al., 2025). The modest effect size relative to perceived usefulness and institutional support suggests pragmatic adaptation: teachers demonstrate willingness to adopt ChatGPT even when conditions remain suboptimal, provided they perceive utility and receive administrative endorsement. This pragmatic orientation reflects accommodations to resource-constrained environments where awaiting ideal conditions would delay innovation.

Teaching experience showed negative relationships with integration willingness ($\beta = -.16$, $p < .05$), while institutional type showed no predictive value. This pattern aligns with perception analysis where extensive experience (>15 years) correlated with significantly lower acceptance than early-career colleagues (≤ 5 years), $F(2, 185) = 8.72$, $p < .001$, $\eta^2 = .086$. Consistency across perception and willingness measures suggests resistance reflects fundamentally different evaluation frameworks rather than technological unfamiliarity. Alm and Ohashi (2024) documented that direct experience correlated with more positive perceptions, yet this cross-sectional design cannot determine whether experience causes attitude change or positive attitudes prompt exploration.

5.3 Theoretical and practical implications

These findings necessitate reconsideration of Technology Acceptance Model applications in hierarchical educational contexts where collective decision-making supersedes individual autonomy. Institutional support's prominence challenges the model's emphasis on individual beliefs as adoption determinants, revealing that power distance significantly moderates acceptance pathways in non-Western environments. This demands theoretical expansion accommodating organizational factors as co-determinants rather than facilitating conditions, particularly where administrative approval carries professional implications. The interdependence between perceived usefulness and institutional support ($r = .52$) indicates acceptance models should account for bidirectional relationships where organizational endorsement shapes individual perception formation. This reconceptualization is supported here for Vietnamese higher education specifically; whether it extends to collectivist educational systems across Southeast Asia more broadly is a claim this single-country sample cannot confirm and is offered as a direction for the comparative research outlined below rather than as an established finding.

These theoretical insights generate practical implications for AI implementation in hierarchical systems. Because perceived usefulness and institutional support were the two strongest predictors of integration willingness, successful ChatGPT integration is more likely to follow from coordinated institutional initiatives that pair administrative endorsement with demonstrations of efficiency gain, rather than from individual teacher enthusiasm or technical training alone. Implementation strategies should emphasize efficiency gains and pedagogical benefits over technical sophistication, acknowledging that adoption decisions emerge from professional necessity. Given the hierarchical perception pattern identified above, training should address content and organizational feedback as

ChatGPT's stronger applications while treating language-level correction as the area requiring the most teacher oversight, rather than presenting all three dimensions as equally ready for use. Given ease of use's attenuated influence, implementation should prioritize institutional support development over user-friendliness concerns for highly educated faculty. Universities can maximize adoption through clear administrative policies, adequate infrastructure, and professional development aligning with pragmatic pedagogical orientations. Experience-dependent patterns necessitate differentiated approaches acknowledging varying sophistication and scepticism across career stages where hierarchical structures amplify generational differences.

6 Conclusion

This investigation reveals that Vietnamese EFL teachers demonstrate sophisticated discrimination in evaluating ChatGPT feedback quality, prioritizing content and organizational guidance over traditional language correction capabilities. The hierarchical perception pattern challenges established assumptions about automated writing evaluation by demonstrating that large language models reverse conventional performance hierarchies, excelling in higher-order rhetorical tasks while generating scepticism regarding surface-level linguistic corrections. These findings suggest that teachers possess refined understanding of the pedagogical complexities inherent in effective feedback provision, recognizing that contextually appropriate guidance requires cultural sensitivity and developmental awareness that current AI systems cannot fully replicate.

The Technology Acceptance Model requires substantial modification when applied to hierarchical educational contexts beyond Vietnam, where institutional support emerges as a co-determinant of adoption willingness rather than a facilitating condition. The prominence of organizational factors reflects hierarchical decision-making structures characteristic of collectivist higher education systems, where individual teacher autonomy operates within collective approval processes that significantly influence technology integration outcomes. The strong predictive power of perceived usefulness indicates that educators in resource-constrained contexts approach AI integration through pragmatic lenses shaped by professional necessity rather than technological enthusiasm, prioritizing efficiency gains over innovation adoption. These findings contribute valuable theoretical insights for understanding technology acceptance in hierarchical educational systems across Southeast Asian and similar cultural contexts, extending beyond Vietnamese-specific applications to inform AI integration frameworks in collectivist educational environments.

Several limitations constrain the generalizability and depth of these findings. The cross-sectional design captures perceptions at a single time point, preventing analysis of attitude changes over extended AI exposure periods. An additional constraint concerns participants' limited practical experience with ChatGPT for educational purposes. While 72% reported general familiarity, only 31% had used it for educational applications. This limited hands-on experience may affect perception accuracy, as teachers with minimal practical experience potentially base evaluations on theoretical understanding rather than direct pedagogical application. The urban sampling focus limits applicability to rural educational contexts where infrastructure and technology adoption patterns may differ significantly. The quantitative approach, while providing robust statistical relationships, lacks the qualitative depth necessary to fully understand the underlying motivations and reasoning processes that shape teacher perceptions of AI feedback quality. Additionally, the study examines intentions rather than actual implementation behaviours, creating potential gaps between expressed willingness and real-world adoption patterns.

Longitudinal investigations tracking actual ChatGPT adoption and implementation outcomes among teachers in hierarchical educational systems represent the most important next step, validating whether the hierarchical perception patterns and institutional support dominance identified in this study accurately predict real-world implementation outcomes. Complementary qualitative re-

search through interviews and focus groups could provide deeper insights into the cultural and pedagogical motivations underlying these patterns. The cultural moderation effects identified warrant investigation in other Southeast Asian educational contexts to determine whether hierarchical influence patterns generalize across similar cultural environments, ultimately contributing to more culturally responsive technology acceptance frameworks for AI integration in diverse higher education systems.

Data availability statement

Data not available – participant consent

References

- Al-Khresheh, M. H. (2024). Bridging technology and pedagogy from a global lens: Teachers' perspectives on integrating ChatGPT in English language teaching. *Computers and Education Artificial Intelligence*, 6, 1–12. <https://doi.org/10.1016/j.caeai.2024.100218>
- Alm, A., & Ohashi, L. (2024). A worldwide study on language educators' initial response to ChatGPT. *Technology in Language Teaching & Learning*, 6(1), 1–23. <https://doi.org/10.29140/tltl.v6n1.1141>
- Alnemrat, A., Aldamen, H., Almashour, M., Al-Deaibes, M., & AlSharefeen, R. (2025). AI vs. teacher feedback on EFL argumentative writing: a quantitative study. *Frontiers in Education*, 10, 1–11. <https://doi.org/10.3389/feduc.2025.1614673>
- Alvi, I., & Khechine, H. (2025). Effect of cultural values on students' adoption of social media for collaborative learning. *Journal of Computer Assisted Learning*, 41(3), 1–19. <https://doi.org/10.1111/jcal.70051>
- APA. (2017). Ethical principles of psychologists and code of conduct. *American Psychological Association*. <https://www.apa.org/ethics/code/>
- Biber, D., Nekrasova, T., & Horn, B. (2011). *The effectiveness of feedback for L1-English and L2-writing development: A meta-analysis*. Educational Testing Service.
- Bucol, J. L., & Sangkawong, N. (2024). Exploring ChatGPT as a writing assessment tool. *Innovations in Education and Teaching International*, 1–16. <https://doi.org/10.1080/14703297.2024.2363901>
- Cohen, J. (1992). A power primer. *Psychological Bulletin*, 112(1), 155–159. <https://doi.org/10.1037/0033-2909.112.1.155>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Earl, C. (2024). Educator-led pedagogical transformations in Vietnam: An introduction. In *Vietnam's creativity agenda. Global Vietnam: Across time, space and community* (pp. 1–12). Springer.
- Faiz, R., Bilal, H. A., Asghar, D. I., & Safdar, A. (2023). Optimizing ChatGPT as a writing aid for EFL learners: balancing assistance and skill development in writing proficiency. *Linguistic Forum - A Journal of Linguistics*, 5(3), 24–37. <https://doi.org/10.5281/zenodo.14756539>
- Field, A. (2018). *Discovering statistics using IBM SPSS Statistics*. SAGE Publications.
- García-Peñalvo, F. J. (2023). The perception of artificial intelligence in educational contexts after the launch of ChatGPT: Disruption or panic? *Education in the Knowledge Society (EKS)*, 24, 1–9. <https://doi.org/10.14201/eks.31279>
- Haber, E., Jemielniak, D., Kurasiński, A., & Przeglądnińska, A. (2025). *Using AI in academic writing and research: A Complete Guide to Effective and Ethical Academic AI*. Springer.
- Hofstede, G. (2001). *Culture's consequences: Comparing values, behaviors, institutions and organizations across nations*. SAGE Publications.
- Imran, M., & Almusharraf, N. (2023). Analyzing the role of ChatGPT as a writing assistant at higher education level: A systematic review of the literature. *Contemporary Educational Technology*, 15(4), 1–14. <https://doi.org/10.30935/cedtech/13605>
- Johnson, R. G., Allen, K., & Cordoba, B. G. (2023). Where does culture belong at school? Exploring the role of individualism and power distance in school belonging across cultures. *Current Psychology*, 43(15), 13492–13527. <https://doi.org/10.1007/s12144-023-05280-y>
- Le, V. C., & Barnard, R. (2009). Curricular innovation behind closed classroom doors: A Vietnamese case study. *Prospect*, 24, 20–33.
- Li, D., & Zhang, L. (2021). Contextualizing feedback in L2 writing: The role of teacher scaffolding. *Language Awareness*, 31(3), 328–350. <https://doi.org/10.1080/09658416.2021.1931261>
- Molefi, R. R., Ayanwale, M. A., Kurata, L., & Chere-Masopha, J. (2024). Do in-service teachers accept artificial intelligence-driven technology? The mediating role of school support and resources. *Computers and Education Open*, 6, 1–16. <https://doi.org/10.1016/j.caeo.2024.100191>
- Mustroph, C., & Steinbock, J. (2024). ChatGPT in foreign language education - friend or foe? A quantitative study on pre-service teachers' beliefs. *Technology in Language Teaching & Learning*, 6(1), 1–17. <https://doi.org/10.29140/tltl.v6n1.1133>
- Ngo, M. T., & Tran, L. T. (2023). Current English education in Vietnam: Policy, practices, and challenges. In *English language education for graduate employability in Vietnam* (pp. 49–69). Springer.
- Nguyen, N. N. H. (2024). Critical review of common issues in classroom procedure of Vietnamese classroom in the new era. *International Journal of Linguistics Literature & Translation*, 7(8), 240–243. <https://doi.org/10.32996/ijllt.2024.7.8.27>

- Nguyen, P. H. (2024). Challenges Vietnamese EFL students faced when learning and meeting the requirements of B1 English language proficiency. *Journal of Language Teaching and Research*, 15(3), 893–901. <https://doi.org/10.17507/jltr.1503.22>
- Nguyen, T. H. N., Pham, T. K., Mai, Q. K., Tran, T. T., & Ta, D. P. (2025). Digital transformation in Vietnam's education: Opportunities, challenges, and development strategies. *Multidisciplinary Reviews*, 8(9), 1–11. <https://doi.org/10.31893/multirev.2025282>
- Nguyen, T. K. C. (2022). Teaching English writing at the secondary level in Vietnam: Policy intentions and classroom practice. *VNU Journal of Science Education Research*, 1, 71–81. <https://doi.org/10.25073/2588-1159/vnuer.4560>
- Nguyen, T. L. (2024). Analysis of teachers' written feedback on writing skills of third-year English-major students at a Vietnamese university. *VNU Journal of Foreign Studies*, 40(1), 103–123. <https://doi.org/10.63023/2525-2445/jfs.ulis.5188>
- Peters, M. A., Jackson, L., Papastephanou, M., Jandrić, P., Lazaroiu, G., Evers, C. W., Cope, B., Kalantzis, M., Araya, D., Tesar, M., Mika, C., Chen, L., Wang, C., Sturm, S., Rider, S., & Fuller, S. (2023). AI and the future of humanity: ChatGPT-4, philosophy and education – Critical responses. *Educational Philosophy and Theory*, 56(9), 828–862. <https://doi.org/10.1080/00131857.2023.2213437>
- Pham, V. P. H., & Truong, M. H. (2021). Teaching writing in Vietnam's secondary and high schools. *Education Sciences*, 11(10), 1–23. <https://doi.org/10.3390/educsci11100632>
- Rudolph, J., Tan, S., & Tan, S. (2023). ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? *Journal of Applied Learning & Teaching*, 6(1), 342–363. <https://doi.org/10.37074/jalt.2023.6.1.9>
- Sallam, M. (2023). ChatGPT utility in healthcare education, research, and practice: Systematic review on the promising perspectives and valid concerns. *Healthcare*, 11(6), 1–20. <https://doi.org/10.3390/healthcare11060887>
- Scherer, R., Siddiq, F., & Tondeur, J. (2019). The technology acceptance model (TAM): A meta-analytic structural equation modeling approach to explaining teachers' adoption of digital technology in education. *Computers & Education*, 128, 13–35. <https://doi.org/10.1016/j.compedu.2018.09.009>
- Shermis, M. D., & Burstein, J. (2013). *Handbook of automated essay evaluation: Current applications and new directions*. Routledge.
- Shin, J. C., Li, X., Nam, I., & Byun, B. (2021). Institutional autonomy and capacity of higher education governance in South Asia: A comparative perspective. *Higher Education Policy*, 35(2), 414–438. <https://doi.org/10.1057/s41307-020-00220-y>
- Singh, S. U., & Namin, A. S. (2025). A survey on chatbots and large language models: Testing and evaluation techniques. *Natural Language Processing Journal*, 1–18. <https://doi.org/10.1016/j.nlp.2025.100128>
- Soori, A., Khojasteh, L., & Javed, F. (2025). Comparing teacher E-feedback, AI feedback, and hybrid feedback in enhancing EFL writing skills. *Technology in Language Teaching & Learning*, 7(3), 1–20. <https://doi.org/10.29140/tltl.v7n3.102626>
- Srite, N., & Karahanna, N. (2006). The role of espoused national cultural values in technology acceptance. *MIS Quarterly*, 30(3), 679–704. <https://doi.org/10.2307/25148745>
- Stevenson, M., & Phakiti, A. (2014). The effects of computer-generated feedback on the quality of writing. *Assessing Writing*, 19, 51–65. <https://doi.org/10.1016/j.asw.2013.11.007>
- Tabachnick, B. G., & Fidell, L. S. (2019). *Using multivariate statistics*. Pearson.
- Teo, T. (2011). Factors influencing teachers' intention to use technology: Model development and test. *Computers & Education*, 57(4), 2432–2440. <https://doi.org/10.1016/j.compedu.2011.06.008>
- Venkatesh, N., Morris, N., Davis, N., & Davis, N. (2003). User acceptance of information Technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186–204. <https://doi.org/10.1287/mnsc.46.2.186.11926>
- Vo, N. H., & Mu, G. M. (2020). Perceived teacher support and students' acceptance of mobile-assisted language learning: Evidence from Vietnamese higher education context. *British Journal of Educational Technology*, 52(2), 879–898. <https://doi.org/10.1111/bjet.13044>
- Wu, H., & Luo, S. (2022). Integrating MOOCs in an undergraduate English course: Students' and teachers' perceptions of blended learning. *SAGE Open*, 12(2), 1–15. <https://doi.org/10.1177/21582440221093035>
- Wu, Y., & Schunn, C. D. (2020). The effects of providing and receiving peer feedback on writing performance and learning of secondary school students. *American Educational Research Journal*, 58(3), 492–526. <https://doi.org/10.3102/0002831220945266>
- Xu, S., Su, Y., & Liu, K. (2025). Investigating student engagement with AI-driven feedback in translation revision: A mixed-methods study. *Education and Information Technologies*, 1–27. <https://doi.org/10.1007/s10639-025-13457-0>
- Zawacki-Richter, O., Marin, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education – where are the educators? *International Journal of Educational Technology in Higher Education*, 16(1), 1–27. <https://doi.org/10.1186/s41239-019-0171-0>
- Zeevy-Solovey, O. (2024). Comparing peer, ChatGPT, and teacher corrective feedback in EFL writing: Students' perceptions and preferences. *Technology in Language Teaching & Learning*, 6(3), 1–23. <https://doi.org/10.29140/tltl.v6n3.1482>
- Zhai, N., & Ma, X. (2021). Automated writing evaluation (AWE) feedback: A systematic investigation of college students' acceptance. *Computer Assisted Language Learning*, 35(9), 2817–2842. <https://doi.org/10.1080/09588221.2021.1897019>
- Zhai, X. (2022). ChatGPT user experience: Implications for education. *SSRN Electronic Journal*, 1–18. <https://doi.org/10.2139/ssrn.4312418>

About the Author

Nguyen Huu Hoang (<https://orcid.org/0009-0005-8131-0434>) holds a PhD in English Studies from Hanoi University, where he currently works as a lecturer. With over ten years of teaching experience, his research focuses on CALL in tertiary EFL education, translation studies, and EMI.