

Vietnamese University Students' Use and Perceptions of Artificial Intelligence Tools in English Language Learning: A Cross-Sectional Survey

Nguyen Hoang Giang

Faculty of Foreign Languages, Ho Chi Minh University of Banking, Ho Chi Minh City, Vietnam, Corporate Relationship Manager at the Vietnam International Commercial Joint Stock Bank (VIB) Headquarters

Abstract—This study documents how Vietnamese undergraduate students use artificial intelligence (AI) tools for English language learning and examines the relationships between usage characteristics, perceived effectiveness, and learner autonomy. A cross-sectional survey of 150 undergraduates from five universities in Ho Chi Minh City collected data on usage patterns together with two multi-item Likert scales measuring perceived effectiveness (10 items) and learner autonomy (20 items). ELSA Speak was the most widely used application (62.7% of respondents), followed by Duolingo (56.7%), ChatGPT (45.3%), and Grammarly (20.0%), and pronunciation was the most frequently cited learning purpose, suggesting that students direct AI tools toward speaking skills that are underserved in Vietnamese classrooms. Both scales showed severely inadequate internal consistency (Cronbach's $\alpha = .07$ and $-.09$), so all scale-based results are interpreted with explicit caution. Perceived effectiveness differed significantly across duration-of-use groups, $F(3, 146) = 3.68$, $p = .014$, $\eta^2 = .07$, with users of less than one month reporting higher perceived effectiveness than users of six to twelve months, a pattern consistent with a novelty effect. No other group differences, correlations, or regression paths reached significance ($R^2 = .03$). The findings offer a locally grounded snapshot of AI adoption among Vietnamese EFL learners and a transparent illustration of how measurement quality constrains inference, underscoring the need for validated, translated, and piloted instruments in AI-in-education survey research.

Index Terms—artificial intelligence, English language learning, learner autonomy, Vietnamese EFL learners

I. INTRODUCTION

Artificial intelligence (AI) applications have moved rapidly from research laboratories into the everyday study habits of language learners. Pronunciation coaches built on automatic speech recognition (ASR), adaptive vocabulary platforms, generative chatbots, and intelligent writing assistants are now available on any smartphone, and reviews of the field document both their expanding capabilities and the unevenness of the evidence base concerning their educational value (Chen et al., 2020; Godwin-Jones, 2022; Kohnke et al., 2023).

Vietnam is a particularly relevant setting for studying this adoption. English proficiency carries high stakes for university graduation and employment, yet instruction remains largely exam-oriented, classes are large, and opportunities for individualized speaking practice are limited (Nguyen, 2021). Pronunciation is a well-documented difficulty for Vietnamese learners, particularly with respect to final consonants and vowel contrasts, and ASR-based applications are precisely the category of tool designed to address it (Mc Crocklin, 2016). At the same time, most locally available evidence on AI tool use is anecdotal or based on small samples, and little is known about which tools Vietnamese students actually use, how intensively, for what purposes, and with what perceived benefit.

The present study addresses this gap with a cross-sectional survey of 150 undergraduates from five universities in Ho Chi Minh City. It contributes a descriptive account of AI tool adoption in this population, a test of the associations between usage

characteristics and two perceptual outcomes perceived effectiveness and learner autonomy and, importantly, a transparent report of the psychometric performance of the survey instrument itself. The scales employed failed conventional reliability standards, and rather than suppressing this result, the study treats it as a substantive methodological finding with direct implications for the growing body of questionnaire-based research on AI in language education (Dörnyei & Taguchi, 2010; Taber, 2018).

II. OBJECTIVES

The study pursued four research objectives: (1) to describe the patterns of AI tool use applications adopted, frequency, session length, duration of use, and primary learning purposes among Vietnamese undergraduate English learners; (2) to examine students' perceived effectiveness of AI tools and whether it differs by gender, usage frequency, duration of use, or learning purpose; (3) to examine the relationships between usage characteristics and self-reported learner autonomy; and (4) to evaluate the internal consistency of the perceived-effectiveness and learner-autonomy scales and the implications of their psychometric performance for interpreting the substantive results.

III. LITERATURE REVIEW

Three theoretical perspectives frame the study. Sociocultural theory locates learning in socially mediated activity within the learner's zone of proximal development, and AI tools can be read as a new class of mediating artefact offering scaffolded, immediate feedback (Vygotsky, 1978). Connectivism argues that in digitally networked environments, learning increasingly consists of forming and navigating connections across distributed sources of knowledge, a description that fits multi-tool learners who combine applications for different skills (Siemens, 2005; see Goldie, 2016, for a critical review). Learner autonomy research, finally, conceptualises autonomy as the capacity to take control of one's own learning and stresses that technology affords, but does not guarantee, autonomous behaviour (Benson & Lamb, 2020).

Empirical work on AI in language learning spans several strands. Early chatbot studies identified both the motivational appeal of conversational agents and its fragility: Fryer and Carpenter (2006) catalogued the pedagogical potential of bots as low-anxiety practice partners, while the experimental comparison by Fryer et al. (2017) found that interest in chatbot partners declined significantly over time whereas interest in human partners did not the clearest demonstration of a novelty effect in this literature. Work on ASR shows that dictation-based pronunciation practice can foster autonomous self-correction (McCrocklin, 2016). More recent scholarship has examined generative AI: Kohnke et al. (2023) analyse the affordances and risks of ChatGPT for language teaching, Godwin-Jones (2022) surveys intelligent writing assistance, and Kohnke et al. (2024) document the technostress that rapid AI adoption imposes on English teachers. Mobile-assisted language learning research in Asian contexts reports generally positive but methodologically uneven effects on proficiency (Wang & Gunaban, 2023), and Vietnamese studies emphasise infrastructural and pedagogical constraints on technology integration (Nguyen, 2021).

A recurring weakness across this literature is measurement. Questionnaire-based studies frequently deploy ad hoc Likert scales without piloting, translation checks, or reported reliability, despite long-standing guidance that Cronbach's alpha must be evaluated and interpreted cautiously whenever multi-item composites are used (Cortina, 1993; Dörnyei & Taguchi, 2010; Taber, 2018). The present study contributes to this methodological conversation by reporting its own instrument's performance in full.

IV. CONCEPTUAL FRAMEWORK

Figure 1 summarises the variable structure examined. Usage characteristics (frequency of use, duration of use, session length) and gender were treated as independent variables; perceived effectiveness and learner autonomy, each measured with a multi-item Likert scale, were the dependent variables. Consistent with connectivist and autonomy-oriented accounts (Benson & Lamb, 2020; Siemens, 2005), the working expectation was that heavier and longer engagement with AI tools would be associated with more positive perceptions and higher self-reported autonomy.

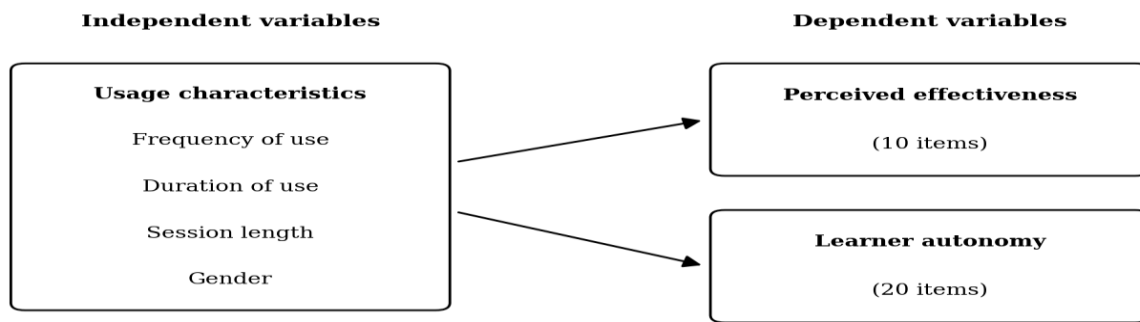


Figure 1 Conceptual framework

V. RESEARCH METHODOLOGY

1. Population and samples

The target population comprised undergraduate students at universities in Ho Chi Minh City who had used at least one AI-based application for English learning. A convenience sample of 150 students was recruited through course groups and social networks; no probability sampling frame was available, and the sample is therefore not statistically representative of the population (Creswell & Creswell, 2018). Respondents came from five institutions HUIT ($n = 44$, 29.3%), University of Science ($n = 32$, 21.3%), UFLO ($n = 27$, 18.0%), University of Economics ($n = 24$, 16.0%), and Van Lang University ($n = 23$, 15.3%) and were nearly balanced by gender (76 female, 50.7%; 74 males, 49.3%). Ages ranged from 18 to 24 years ($M = 21.00$, $SD = 1.92$), and respondents reported between 5 and 12 years of prior English study ($M = 8.34$, $SD = 2.41$).

2. Research instrument

Data were collected with a 40-item structured questionnaire in four sections: (a) demographic information (5 items); (b) AI tool usage patterns (5 items covering tools used, frequency, session length, duration of use, and primary learning purpose); (c) perceived effectiveness of AI tools (10 items, numbered 11–20); and (d) learner autonomy (20 items, numbered 21–40), with the autonomy items written to reflect the conceptualization of autonomy as control over one's learning in Benson and Lamb (2020). Sections (c) and (d) used five-point Likert response scales (1 = strongly disagree, 5 = strongly agree). The

questionnaire was administered in English without back-translation, and no pilot study was conducted before full administration two decisions that departed from recommended questionnaire practice (Dörnyei & Taguchi, 2010) and that are directly relevant to the reliability results reported below.

3. Collection of data

Responses were collected online over a two-week window in June 2026. Participation was voluntary and anonymous; respondents saw an informed-consent statement describing the study's purpose, the approximate completion time (10–15 minutes), and their right to skip questions or withdraw, and consented before the questionnaire items were displayed. No personally identifying information was collected. The study was conducted as university coursework, and formal institutional ethics review was not obtained prior to data collection; an ethics application was prepared subsequently to support any extension of the research, and this sequencing is acknowledged as a limitation. All 150 returned questionnaires were complete and retained for analysis. During data screening, one tool name recorded inconsistently in the raw data ("ELLA Speak") was recoded to ELSA Speak.

4. Data analysis

Analyses proceeded in three stages using statistical software. First, descriptive statistics summarised the sample and usage patterns. Second, the internal consistency of the two multi-item scales was evaluated with Cronbach's alpha and corrected item–total correlations (Cortina, 1993; Taber, 2018). Third,

inferential analyses examined group differences and associations: independent-samples *t* tests (gender), one-way analyses of variance with Tukey HSD follow-up comparisons (usage frequency, duration of use, learning purpose), Pearson correlations, and a multiple linear regression predicting learner autonomy from usage frequency, duration of use, session length, perceived effectiveness, and gender. All tests were two-tailed with $\alpha = .05$, and effect sizes are reported as Cohen's *d* or η^2 (Cohen, 1988). Given the reliability results, all inferential findings involving the composite scales are interpreted with explicit caution.

VI. RESULTS

Results are reported in the order of the research objectives, with one deliberate re-ordering: the reliability analysis (Objective 4) is presented immediately after the usage descriptives, because its outcome conditions how the scale-based results for Objectives 2 and 3 may legitimately be interpreted. All statistics were computed from the full dataset of 150 responses.

Patterns of AI tool use

Table 1 presents the usage results. ELSA Speak was the most widely used application (62.7% of respondents), followed by Duolingo (56.7%), ChatGPT (45.3%), and Grammarly (20.0%); respondents used 1.85 tools on average, confirming that multi-tool use is the norm rather than the exception. Usage frequency was spread across the three categories, with weekly-but-not-daily use most common (38.7% used the tools once or twice a week) and under a third of the sample (29.3%) reporting daily use. Sessions averaged 37.5 minutes ($SD = 13.2$, range 15–60). Duration of use was similarly dispersed: roughly half the sample (51.3%) had been using AI tools for six months or less, while only 19.3% reported more than a year of use. Primary learning purposes were remarkably evenly distributed, with pronunciation (23.3%) and grammar (21.3%) leading narrowly over vocabulary, writing, and conversation. The prominence of ELSA Speak and of pronunciation as a stated purpose suggests that students are directing AI tools toward the skill area where classroom

provision is weakest (McCrocklin, 2016; Nguyen, 2021).

Table 1 Patterns of AI Tool Use for English Learning
($N = 150$)

Usage characteristic	n	%
Tools used (multiple selection)		
ELSA Speak	94	62.7
Duolingo	85	56.7
ChatGPT	68	45.3
Grammarly	30	20.0
Frequency of use		
1–2 times/week	58	38.7
3–5 times/week	48	32.0
Daily	44	29.3
Duration of use		
Less than 1 month	39	26.0
1–6 months	38	25.3
6–12 months	44	29.3
More than 1 year	29	19.3
Primary purpose		
Pronunciation	35	23.3
Grammar	32	21.3
Vocabulary	28	18.7
Writing	28	18.7
Conversation	27	18.0

Note. Tool percentages sum to more than 100 because respondents could select multiple tools ($M = 1.85$ tools per respondent). Average session length was 37.5 minutes ($SD = 13.2$, range 15–60).

Reliability of the measurement scales

Table 2 reports the internal-consistency analysis. Both scales fell dramatically short of the conventional $\alpha \geq .70$ benchmark (Taber, 2018): Cronbach's α was .07 for the ten-item perceived-effectiveness scale and $-.09$ for the twenty-item learner-autonomy scale. Corrected item-total correlations were correspondingly negligible, ranging from $-.06$ to $.10$ for the effectiveness items and from $.13$ to $.14$ for the autonomy items, and no item deletion produced a materially improved α for either scale.

Table 2 Internal Consistency of the Two Multi-Item Scales

Scale	Items	M	SD	α	Item-total r range
Perceived effectiveness (items 11–20)	10	3.53	0.37	.07	-.06 to .10
Learner autonomy (items 21–40)	20	3.52	0.24	-.09	-.13 to .14

Note. N = 150. M and SD refer to the composite (mean-of-items) score on the 1–5 response scale. α = Cronbach's alpha; item-total r = corrected item-total correlation.

These values indicate that responses to the items within each scale were essentially uncorrelated. An alpha near zero and especially a negative alpha is incompatible with the claim that the items jointly measure a single construct (Cortina, 1993). Several non-exclusive explanations are plausible: the instrument was not piloted; the items were administered in English without back-translation, so wording problems cannot be ruled out (Dörnyei & Taguchi, 2010); and the questionnaire contained no reverse-coded items or attention checks, so careless responding cannot be detected or excluded. The combination of near-uniform item means (all between 3.41 and 3.69) with near-zero inter-item correlations is in fact a pattern consistent with substantial random responding. Whatever the cause, the methodological consequence is unambiguous: the composite scores cannot be interpreted as trustworthy measures of perceived effectiveness or learner autonomy, and every scale-based result reported below must be read as provisional and primarily illustrative of the analytic procedure.

VII. PERCEIVED EFFECTIVENESS

At the descriptive level, the composite perceived-effectiveness score was moderately positive ($M = 3.53$, $SD = 0.37$ on the five-point scale). Item means were tightly clustered: the highest endorsements were for perceived overall proficiency improvement (item 20, $M = 3.62$) and preference for practising with AI rather than classmates (item 13, $M = 3.58$), while progress tracking received the lowest mean (item 17, $M = 3.45$). Given the reliability results, these figures are best read item-by-item as single indicators rather than as evidence about a coherent underlying attitude. One group difference reached statistical significance. A one-way ANOVA comparing composite effectiveness scores across the four duration-of-use groups was significant, $F(3, 146) = 3.68$, $p = .014$, $\eta^2 = .07$. Tukey HSD comparisons showed that respondents who had used AI tools for less than one

month reported higher perceived effectiveness ($M = 3.64$, $SD = 0.39$) than those who had used them for 6–12 months ($M = 3.41$, $SD = 0.38$), $p = .019$; no other pairwise contrast was significant. Perceived effectiveness did not differ by usage frequency, $F(2, 147) = 0.53$, $p = .590$, by primary purpose, $F(4, 145) = 0.59$, $p = .673$, or by gender, $t(148) = -0.05$, $p = .958$, $d = -0.01$.

VIII. LEARNER AUTONOMY AND ITS CORRELATES

The composite autonomy score was likewise moderately positive ($M = 3.52$, $SD = 0.24$), with item means between 3.41 (evaluating one's own performance) and 3.69 (identifying one's weaknesses). Table 3 presents the composite scores by group. None of the tested associations with autonomy was statistically significant. Autonomy did not differ by gender, $t(148) = -1.22$, $p = .224$, $d = -0.20$, by usage frequency, $F(2, 147) = 2.80$, $p = .064$, $\eta^2 = .04$, or by duration of use, $F(3, 146) = 1.35$, $p = .261$. Autonomy was uncorrelated with perceived effectiveness, $r(148) = -.06$, $p = .481$, and with session length, $r(148) = -.13$, $p = .128$.

Table 3 Scale Descriptives and Group Comparisons

Group	n	Effectiveness M (SD)	Autonomy M (SD)
Gender			
Female	76	3.53 (0.37)	3.55 (0.25)
Male	74	3.53 (0.37)	3.50 (0.23)
Frequency of use			
1–2 times/week	58	3.51 (0.40)	3.51 (0.20)
3–5 times/week	48	3.51 (0.37)	3.59 (0.28)
Daily	44	3.58 (0.32)	3.47 (0.23)
Duration of use			
Less than 1 month	39	3.64 (0.39)	3.50 (0.22)
1–6 months	38	3.61 (0.34)	3.58 (0.27)
6–12 months	44	3.41 (0.38)	3.48 (0.23)
More than 1 year	29	3.48 (0.30)	3.55 (0.24)
Total sample	150	3.53 (0.37)	3.52 (0.24)

Note. Composite scores on the 1–5 response scale. Of all group contrasts in this table, only the difference in

perceived effectiveness across duration-of-use groups was statistically significant (see text).

A multiple linear regression predicting autonomy from usage frequency, duration of use, session length, perceived effectiveness, and gender (Table 4) confirmed the null pattern: the model explained essentially no variance, $R^2 = .03$, $F(5, 144) = 0.86$, $p = .511$, and no individual predictor was significant. This stands in contrast to the theoretically motivated expectation, derived from connectivist and autonomy-focused accounts (Benson & Lamb, 2020; Siemens, 2005), that heavier and longer engagement with AI tools would be accompanied by higher self-reported autonomy.

Table 4 Multiple Regression Predicting Learner Autonomy Scores (N = 150)

Predictor	b	β	t	p
Usage frequency	-0.01	-.04	-0.48	.633
Duration of use	0.00	.01	0.12	.907
Session length (min)	-0.00	-.11	-1.36	.177
Perceived effectiveness	-0.03	-.05	-0.59	.555
Gender (male = 1)	-0.05	-.10	-1.17	.242

Note. $R^2 = .03$, $F(5, 144) = 0.86$, $p = .511$. Usage frequency coded 1–3 (1–2 times/week to daily); duration coded 1–4 (less than 1 month to more than 1 year). Unstandardised coefficients rounded to two decimals.

IX. DISCUSSION

The most secure findings of this study are the descriptive usage results, which rest on single factual items rather than on the problematic composite scales. Vietnamese undergraduates in this sample are multi-tool users who favour pronunciation-oriented applications, practise in half-hour sessions a few times a week, and are mostly recent adopters. The alignment between the leading tool (ELSA Speak) and the leading stated purpose (pronunciation) suggests purposeful rather than incidental adoption, and it is consistent with the documented salience of pronunciation difficulties for Vietnamese learners (Nguyen, 2021) and with the pedagogical potential of automatic speech recognition identified by McCrocklin (2016). For practitioners, the descriptive

picture indicates that students are already directing AI tools toward the skill area where classroom provision is weakest, largely without institutional guidance.

The one significant inferential result higher perceived effectiveness among the newest users than among those with 6–12 months of use is consistent with the novelty-and-decline pattern reported in the chatbot and mobile-learning literatures, in which initial enthusiasm fades as tools become routine (Fryer et al., 2017; Wang & Gunaban, 2023). The slight recovery among users of more than a year ($M = 3.48$) could be read as a survivorship effect, whereby only students who find the tools genuinely useful persist beyond the first year. Both interpretations, however, must remain tentative: the effect is modest ($\eta^2 = .07$), the comparison is cross-sectional rather than longitudinal, and the effectiveness scale on which it is based failed the reliability check.

The null associations between usage characteristics and learner autonomy admit at least three readings. Substantively, technology use may simply not translate into autonomy without deliberate pedagogical scaffolding, a position long argued in the autonomy literature (Benson & Lamb, 2020). Methodologically, the unreliable autonomy scale attenuates any true association toward zero, so the study cannot distinguish a genuine absence of relationship from a measurement artefact. Statistically, the coding of frequency and duration as coarse ordinal categories reduces power. The honest conclusion is that the question remains open.

The reliability failure itself is arguably the study's most instructive result. Alphas of .07 and -.09, combined with near-uniform item means and negligible inter-item correlations, most plausibly reflect some combination of unpiloted items, English-only administration to non-native speakers without back-translation, and undetected careless responding (Cortina, 1993; Dörnyei & Taguchi, 2010; Taber, 2018). The result illustrates, at the study's own expense, that enthusiasm for a research topic is no substitute for measurement rigour, and it cautions readers of the AI-in-education survey literature to attend closely to whether and how instrument reliability is reported. Beyond measurement, the study's limitations include the convenience sample, the cross-sectional design, exclusive reliance on self-report, and the fact that formal ethics review followed rather than preceded data collection.

X. SUGGESTIONS

1. For teachers and institutions: integrate AI tools deliberately rather than leaving adoption entirely to students. The usage data show that students already gravitate toward pronunciation practice with ASR-based applications; structured tasks that connect this out-of-class practice to classroom speaking assessment, together with explicit training in self-monitoring, would exploit the tools' scaffolding potential while addressing the possibility that unguided use does not by itself foster autonomy (Benson & Lamb, 2020; McCrocklin, 2016).

2. For programme designers: plan for the post-novelty phase. Perceived effectiveness in this sample was highest among users of less than one month and lowest at 6–12 months, echoing the interest decline documented experimentally by Fryer et al. (2017). Sustained-engagement features renewed goals, varied task types, periodic progress reviews should be built into any AI-supported programme rather than assuming that initial enthusiasm will persist.

3. For researchers: treat instrumentation as a first-order concern. Future questionnaire studies in this area should use validated scales translated into the respondents' first language with back-translation, pilot the instrument and report reliability before full administration, include reverse-coded items and attention checks to detect careless responding, and, where possible, complement self-report with objective measures such as standardised proficiency tests or in-app usage logs, ideally within longitudinal designs (Dörnyei, 2007; Dörnyei & Taguchi, 2010; Taber, 2018).

REFERENCES

- [1] P. Benson and T. Lamb, "Autonomy in the age of multilingualism," in *Autonomy in Language Education: Theory, Research and Practice*, M. Jiménez Raya and F. Vieira, Eds. London, U.K.: Routledge, 2020, pp. 74–88.
- [2] L. Chen, P. Chen, and Z. Lin, "Artificial intelligence in education: A review," *IEEE Access*, vol. 8, pp. 75264–75278, 2020, doi: 10.1109/ACCESS.2020.2988510.
- [3] J. Cohen, *Statistical Power Analysis for the Behavioral Sciences*, 2nd ed. Hillsdale, NJ, USA: Lawrence Erlbaum Associates, 1988.
- [4] J. M. Cortina, "What is coefficient alpha? An examination of theory and applications," *Journal of Applied Psychology*, vol. 78, no. 1, pp. 98–104, 1993.
- [5] J. W. Creswell and J. D. Creswell, *Research Design: Qualitative, Quantitative, and Mixed Methods Approaches*, 5th ed. Thousand Oaks, CA, USA: SAGE Publications, 2018.
- [6] Z. Dörnyei, *Research Methods in Applied Linguistics: Quantitative, Qualitative, and Mixed Methodologies*. Oxford, U.K.: Oxford University Press, 2007.
- [7] Z. Dörnyei and T. Taguchi, *Questionnaires in Second Language Research: Construction, Administration, and Processing*, 2nd ed. London, U.K.: Routledge, 2010.
- [8] L. K. Fryer and R. Carpenter, "Bots as language learning tools," *Language Learning & Technology*, vol. 10, no. 3, pp. 8–14, 2006.
- [9] L. K. Fryer, M. Ainley, A. Thompson, A. Gibson, and Z. Sherlock, "Stimulating and sustaining interest in a language course: An experimental comparison of chatbot and human task partners," *Computers in Human Behavior*, vol. 75, pp. 461–468, 2017.
- [10] R. Godwin-Jones, "Partnering with AI: Intelligent writing assistance and instructed language learning," *Language Learning & Technology*, vol. 26, no. 2, pp. 5–24, 2022.
- [11] J. G. S. Goldie, "Connectivism: A knowledge learning theory for the digital age?" *Medical Teacher*, vol. 38, no. 10, pp. 1064–1069, 2016.
- [12] L. Kohnke, B. L. Moorhouse, and D. Zou, "ChatGPT for language teaching and learning," *RELC Journal*, vol. 54, no. 2, pp. 537–550, 2023.
- [13] L. Kohnke, D. Zou, and B. L. Moorhouse, "Technostress and English language teaching in the age of generative AI," *Educational Technology & Society*, vol. 27, no. 2, pp. 306–320, 2024.
- [14] S. M. McCrocklin, "Pronunciation learner autonomy: The potential of automatic speech recognition," *System*, vol. 57, pp. 25–42, 2016.
- [15] L. T. H. Nguyen, "Teachers' perception of ICT integration in English language teaching at

- Vietnamese tertiary level,” *European Journal of Contemporary Education*, vol. 10, no. 3, pp. 697–710, 2021.
- [16] G. Siemens, “Connectivism: A learning theory for the digital age,” *International Journal of Instructional Technology and Distance Learning*, vol. 2, no. 1, pp. 3–10, 2005.
- [17] K. S. Taber, “The use of Cronbach's alpha when developing and reporting research instruments in science education,” *Research in Science Education*, vol. 48, no. 6, pp. 1273–1296, 2018, doi: 10.1007/s11165-016-9602-2.
- [18] L. S. Vygotsky, *Mind in Society: The Development of Higher Psychological Processes*. Cambridge, MA, USA: Harvard University Press, 1978.
- [19] X. Wang and M. G. Gunaban, “Effectiveness of mobile-assisted language learning in enhancing English proficiency,” *Journal of Contemporary Educational Research*, vol. 7, no. 11, pp. 140–146, 2023, doi: 10.26689/jcer.v7i11.5588.